

Distributed Adaptive Nearest Neighbor Classifier: Algorithm and Theory

Ruiqi Liu^{1*}, Ganggang Xu² and Zuofeng Shang³

^{1*}Department of Mathematics and Statistics, Texas Tech University,
Lubbock, 79409, TX, USA.

²Department of Management Science, University of Miami, Coral
Gables, 33146, FL, USA.

³Department of Mathematical Sciences, New Jersey Institute of
Technology, City, 07102, NJ, USA.

*Corresponding author(s). E-mail(s): ruiqliu@ttu.edu;
Contributing authors: gangxu@bus.miami.edu; zshang@njit.edu;

Abstract

When data is of an extraordinarily large size or physically stored in different locations, the distributed nearest neighbor (NN) classifier is an attractive tool for classification. We propose a novel distributed adaptive NN classifier for which the number of nearest neighbors is a tuning parameter stochastically chosen by a data-driven criterion. An early stopping rule is proposed when searching for the optimal tuning parameter, which not only speeds up the computation but also improves the finite sample performance of the proposed algorithm. Convergence rate of excess risk of the distributed adaptive NN classifier is investigated under various sub-sample size compositions. In particular, we show that when the sub-sample sizes are sufficiently large, the proposed classifier achieves the nearly optimal convergence rate. Effectiveness of the proposed approach is demonstrated through simulation studies as well as an empirical application to a real-world dataset.

Keywords: Distributed Learning, Adaptive Procedure, Minimax Optimal, Binary Classification

1 Introduction

Nearest neighbor (NN) classifier is a simple but powerful tool for various applications such as text classification (Han et al., 2001, Jiang et al., 2012), query dependent ranking (Geng et al., 2008), and pattern recognition (Kowalski and Bender, 1972, Zheng et al., 2004, Xu et al., 2013). Consider $(Y_1, \mathbf{X}_1), \dots, (Y_N, \mathbf{X}_N)$ generated independently from an unknown probability distribution P , with $Y_i \in \{0, 1\}$ being the label and $\mathbf{X}_i$ being the corresponding d -dimensional feature vector for $i = 1, \dots, N$. The NN classifier predicts the label of a query point $\mathbf{x}$ based on labels of its neighboring observations. It is well-known that NN algorithm is sensitive to the scale of data as it relies on computing the distances. A popular procedure is to normalize each feature to $[0, 1]$. Without loss of generality, we assume that the feature space is $[0, 1]^d$ and that the Euclidean distance is used. This assumption was also used in Cai and Wei (2019). Given a new query point $\mathbf{x} \in [0, 1]^d$, denote $\mathbf{X}_{(i)}(\mathbf{x})$ as the i -th nearest point to $\mathbf{x}$ among $\mathbf{X}_1, \dots, \mathbf{X}_N$, and $Y_{(i)}(\mathbf{x})$ as the label associated with $\mathbf{X}_{(i)}(\mathbf{x})$. For a prespecified integer $1 \leq k \leq N$, the conditional probability $\eta(\mathbf{x}) := \mathbb{P}(Y = 1 | \mathbf{X} = \mathbf{x})$ can be approximated by the k -NN estimator $\eta_{NN,k}(\mathbf{x}) = \frac{1}{k} \sum_{i=1}^k Y_{(i)}(\mathbf{x})$ and the label associated with $\mathbf{x}$ is then predicted as $f_{NN,k}(\mathbf{x}) = \mathbb{I}(\eta_{NN,k}(\mathbf{x}) \geq 1/2)$, with $\mathbb{I}(\cdot)$ being the indicator function.

The performance of a binary classifier $f : [0, 1]^d \rightarrow \{0, 1\}$, which is trained using observed data $(Y_1, \mathbf{X}_1), \dots, (Y_N, \mathbf{X}_N)$, is commonly evaluated by the regret (or excess risk) defined as

$$\mathcal{R}(f) = \mathbb{P}(f(\mathbf{X}) \neq \mathbf{Y}) - \mathbb{P}(f^*(\mathbf{X}) \neq \mathbf{Y})$$

where $(Y, \mathbf{X}) \sim P$ is an independent copy of the training sample, $f^*(\mathbf{x}) = \mathbb{I}(\eta(\mathbf{x}) \geq 1/2)$ is the well-known Bayesian classifier, and the probability is with respect to the joint distribution of $(Y_1, \mathbf{X}_1), \dots, (Y_N, \mathbf{X}_N)$ and $(Y, \mathbf{X})$. A smaller regret indicates higher classification accuracy for a classifier f .

Notation: For deterministic positive sequences a_N and b_N , we denote $a_N \asymp$ (or $\asymp$) b_N if $a_N \asymp Cb_N$ for some $C > 0$ and sufficiently large N . If $a_N \asymp b_N$ and $a_N \asymp b_N$, we write $a_N \asymp b_N$. For any $a > 0$, we denote $\lceil a \rceil$ ($\lfloor a \rfloor$) as the smallest (largest) integer that is not less (greater) than a . We denote μ as the Lebesgue measure and $P_{\mathbf{X}}$ as the marginal distribution of $\mathbf{X}$ whose support is $\mathcal{X}$. For a set A , we use $|A|$ to denote its cardinality.

1.1 Related Work

The regret of the k -NN classifier has been shown to converge to 0 as $k \rightarrow \infty$ and $k/N \rightarrow 0$ in a general metric space with additional structural assumptions (Cover and Hart, 1967, Cerou and Guyader, 2006, Hanneke et al., 2021) and in the Euclidean space (Stone, 1977, Devroye et al., 1994). The convergence rate of the regret depends on properties of $\eta(\mathbf{x})$ and $P_{\mathbf{X}}$. Chaudhuri and Dasgupta (2014) established a nonasymptotic bound for the convergence rate, which achieves the minimax rate in the sense of Audibert and Tsybakov (2007) under some mild conditions. Gadat et al. (2016) further identified two sufficient and necessary conditions for the uniform consistency of

the k -NN classifier without rigid assumptions on the joint distribution of $(Y, \mathbf{X})$ and derived the corresponding optimal convergence rate. Samworth (2012) proposed an optimally weighted k -NN classifier based on a new asymptotic expansion of its regret.

When facing an extraordinarily large sample size, the k -NN classifier can be computationally intensive, especially when k is large. To address this issue, Qiao et al. (2019) and Duan et al. (2020) proposed two distributed k -NN classifiers, extending the work of Chaudhuri and Dasgupta (2014) and Samworth (2012), respectively. Their algorithms first divide the whole data into m equally-sized sub-samples, and for each sub-sample, a k -NN classifier is trained independently. The final prediction of a new query point is made by aggregating the m independently trained k -NN classifiers. Under suitable conditions, the regrets of both distributed k -NN classifiers were shown to achieve the optimal convergence rate. However, in many applications, the sub-samples may not have equal sample sizes, and to the best of our knowledge, there has yet been any existing work on distributed NN classifiers with unequal sized sub-samples.

Furthermore, the aforementioned theoretical results are based on the key assumption that the choice of k is pre-given and is deterministic. However, it is often desirable to have a data-driven choice of k for practical applications. There has been limited work on theoretical properties of the k -NN classifier with a data-driven choice of k in existing literature, with two notable exceptions, i.e., Cai and Wei (2019) and Bal-subramani et al. (2019). They independently proposed two adaptive procedures to stochastically choose k and established the convergence rates of the resulting adaptive NN classifiers under suitable conditions. However, while achieving improved classification accuracy, searching for an optimal k also significantly increases the computational burden for the adaptive NN classifier, making it desirable to consider a distributed adaptive NN classifier with favorable statistical properties when the sample size N is extraordinarily large. For applications where data are stored in different locations, a distributed adaptive NN classifier is also a natural and preferable choice.

1.2 Our Contribution

We propose a novel distributed adaptive NN classifier with a data-driven choice of k , which can be used to either speed up the computation when the data size is extraordinarily large or improve the classification accuracy when data are stored in different machines. Suppose that the whole data set is separately stored in m different locations, and each location has a sub-sample of size n_j , $j = 1, \dots, m$. The sub-sample sizes are allowed to be different from each other, in contrast to the existing divide-and-conquer framework (Qiao et al., 2019, Duan et al., 2020). Without loss of generality, we assume that $n_1 \geq n_2 \geq \dots \geq n_m$ and denote $N = n_1 + \dots + n_m$. Based on the j th sub-sample, a local k_j -NN classifier is constructed for a given query point $\mathbf{x}$ and an integer k_j , $j = 1, \dots, m$. The predicted label for $\mathbf{x}$ is then obtained by aggregating the m sub-sample NN classifiers with $k_1, \dots, k_m$ chosen by a data-driven criterion. See Section 2 for more details.

The computational efficiency of the proposed algorithm is achieved in two ways.

(1) Parallel computation. For a given k , the computational complexity of the standard k -NN classifier using the whole data is between $O(N)$ to $O(N \log(N))$ (Cormen et al., 2009), which needs to be carried out on a single machine. In comparison,

the computation of the distributed NN classifier can be easily paralleled, and each sub-sample only costs between $O(n)$ to $O(n \log(n))$ operations.

(2) Early stopping rule for k . The adaptive NN classifiers proposed in Cai and Wei (2019) and Balsubramani et al. (2019) search for an optimal k_1 by increasing k from 1 to N until a stopping rule is triggered. A straightforward extension of their approaches to the distributed setting is to search for k_j from 1 to n_j , $j = 1 \dots m$. However, we propose an early stopping rule for the choice of k_1 (which determines other k_j s), narrowing down the search range for k_1 to $1 \leq k_1 \leq n_1 N^{\frac{d}{2+d} \log(N)}$. As a result, the proposed algorithm significantly reduces the number of attempts needed to locate the optimal k_j s for the distributed adaptive NN classifier.

Our numerical studies show that such an early stopping rule for k_1 not only speeds up the computation but also yields superior finite sample performance for the proposed algorithm compared to the naive extension of Cai and Wei (2019) and Balsubramani et al. (2019). See Section 4.1 for more details.

From a theoretical point of view, our work extends the theory for distributed NN classifier with a fixed k (Qiao et al., 2019) to the more realistic distributed adaptive NN classifier based on unequal sub-sample sizes, whose k_j s are chosen by a data-driven procedure. Specifically, we derive the convergence rate of the regret of the proposed classifier and give sufficient conditions under which the convergence rate is optimal (up to logarithmic factors). Moreover, the convergence rate of the regret exhibits a phase transition characterized by sub-sample sizes. Finally, we wish to comment that the proof of adaptivity in the distributed framework relies on the uniform convergence in Lemma 5. This requires bounding the total model complexity (see Lemma 1) of all the local classifiers, which motivates the choices of k_j s in Algorithm 1.

The rest of this paper is structured as follows. Section 2 introduces the algorithm for the distributed adaptive NN classifier and Section 3 investigates its asymptotic properties. Section 4.1 carries out a set of simulation studies, and a real-world dataset is analyzed in Section 4.2. All technical proofs are provided in the Appendix.

2 Distributed Adaptive Nearest Neighbor Classifier

Suppose that the whole dataset, denoted as $\mathcal{Z} = (Y_1 \mathbf{X}_1) \dots (Y_N \mathbf{X}_N)$, are distributed across m machines. Each machine hosts a sub-sample of size n_j , denoted as $\mathcal{Z}_j = (Y_1^j \mathbf{X}_1^j) \dots (Y_{n_j}^j \mathbf{X}_{n_j}^j)$ for $j = 1 \dots m$. For the j th sub-sample, given an integer $k_j = 1 \dots n_j$, the j th local NN estimator of $\eta(\mathbf{x}) = \mathbb{P}(Y = 1 | \mathbf{X} = \mathbf{x})$ for a new query point $\mathbf{x} \in [0, 1]^d$ is defined as

$$\hat{\eta}_{k_j, j}(\mathbf{x}) = \frac{1}{k_j} \sum_{i=1}^{k_j} Y_{(i)}^j(\mathbf{x}) \quad j = 1 \dots m$$

where $Y_{(i)}^j(\mathbf{x})$ is the label associated with $\mathbf{X}_{(i)}^j(\mathbf{x})$, the i -th nearest neighbors of $\mathbf{x}$ among $\mathbf{X}_1^j, \dots, \mathbf{X}_{n_j}^j$. The proposed distributed NN classifier is subsequently defined as

$$f_{k_1:k_m}(\mathbf{x}) = \mathbb{I}(\mathbf{x}_{k_1:k_m}(\mathbf{x}) \geq 1/2) \quad \text{with} \quad \mathbf{x}_{k_1:k_m}(\mathbf{x}) = \frac{1}{\sum_{j=1}^m k_j} \sum_{j=1}^m k_j Y_{k_j}^j(\mathbf{x}) \quad (1)$$

where the integer sequence $k_1, \dots, k_m$ need to be chosen by some data-driven method.

The performance of the classifier (1) depends critically on the choice of $k_1, \dots, k_m$. The following Algorithm 1 is designed in the same spirit of Cai and Wei (2019) and Balsubramani et al. (2019).

Algorithm 1: Distributed Adaptive NN Classifier

Input: new query $\mathbf{x}$, training samples $\mathcal{Z}_j$, $j = 1, \dots, m$;
Initialization: set $k_1 = 0$;
while $k_1 < n_1 N^{-\frac{d}{2+d}} \log(N)$ **do**
 update $k_1 := k_1 + 1$;
 update $k_j := k_1 n_j / n_1$ and calculate $\mathbf{x}_{k_j}^j(\mathbf{x})$ for $j = 1, \dots, m$;
 calculate $\mathbf{x}_{k_1:k_m}(\mathbf{x})$ and $r_{k_1} = \frac{1}{2} \sum_{j=1}^m k_j \mathbf{x}_{k_1:k_m}(\mathbf{x}) \geq 1/2$;
 if $r_{k_1} > \frac{1}{(d+2) \log(N)}$ **or** $k_1 < n_1 N^{-\frac{d}{2+d}} \log(N)$ **then**
 set $k_j = k_1 n_j / n_1$ for $j = 2, \dots, m$;
 calculate $\mathbf{x}_{k_1:k_m}(\mathbf{x})$;
 exit loop;
 end if
end while
Output: classifier $f_{k_1:k_m}(\mathbf{x}) = \mathbb{I}(\mathbf{x}_{k_1:k_m}(\mathbf{x}) \geq 1/2)$.

Algorithm 1 assumes that each k_j is proportional to n_j for $j = 1, \dots, m$, and search for the optimal k_1 within the set $1 \leq k_1 \leq n_1 N^{-\frac{d}{2+d}} \log(N)$ such that $\mathbf{x}_{k_1:k_m}(\mathbf{x}) \geq 1/2$ based on the classifier (1) is strictly greater than $\frac{1}{(d+2) \log(N)} (2 \sum_{j=1}^m k_j)$. If no k_1 meets this criterion, we simply set $k_1 = n_1 N^{-\frac{d}{2+d}} \log(N)$. We comment that a naive extension of Cai and Wei (2019) and Balsubramani et al. (2019) to the distributed data setting would require searching for k_1 from 1 to n_1 . In this sense, the upper bound $n_1 N^{-\frac{d}{2+d}} \log(N)$ in Algorithm 1 serves an early stopping rule for the search of k_1 . Our simulation studies demonstrate that such an early stopping rule yields superior finite sample performance compared to the same algorithm but searches k_1 from 1 to n_1 .

An intuitive justification of Algorithm 1 is as follows. Denote $\mathcal{X} = \mathbf{X}_1, \dots, \mathbf{X}_N$. Under suitable conditions, one can show that $\mathbf{x}_{k_1:k_m}(\mathbf{x}) - \mathbb{E}(\mathbf{x}_{k_1:k_m}(\mathbf{x}) | \mathcal{X})$ is bounded by the sequence $\frac{1}{(d+2) \log(N)} (2 \sum_{j=1}^m k_j)$ uniformly for all $\mathbf{x}$ and $k_1, \dots, k_m$ with a high probability. The stopping rule designed in Algorithm 1 thus ensures that

$f_{k_1:k_m}(\mathbf{x}) \in \{-1, 2\}$ and $\mathbb{E}(f_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) \in \{-1, 2\}$ have the same sign with a high probability. Under suitable conditions, $\mathbb{E}(f_{k_1:k_m}(\mathbf{x}) | \mathcal{X})$ is a consistent estimator of $f(\mathbf{x})$, which further implies that the distributed adaptive NN classifier $f_{k_1:k_m}(\mathbf{x}) = \mathbb{I}(f_{k_1:k_m}(\mathbf{x}) \in \{-1, 2\})$ is asymptotically equivalent to the Bayesian classifier $f(\mathbf{x}) = \mathbb{I}(f(\mathbf{x}) \in \{-1, 2\})$.

3 Asymptotic Properties

3.1 Technical Assumptions

To investigate the asymptotic properties of the proposed adaptive distributed NN classifier obtained from Algorithm 1, several technical assumptions are needed.

Assumption A1. (Strong Density) For some constants $c, r > 0$, it holds that (a) $\mathbb{P}(B(\mathbf{x}, r)) \geq c$ for all $0 < r < r$ and $\mathbf{x}$; and (b) $c < \frac{dP_{\mathbf{X}}(\mathbf{x})}{d} < c^{-1}$ for all $\mathbf{x}$.

Assumption A2. (Smoothness) There exist constants $\alpha \in (0, 1]$ and $C > 0$ such that $\|\mathbf{x}_1 - \mathbf{x}_2\| \leq C \|\mathbf{x}_1 - \mathbf{x}_2\|^\alpha$ holds for all $\mathbf{x}_1, \mathbf{x}_2$.

Assumption A3. (Marginal Assumption) For some constants $\beta \in [0, d]$ and $C > 0$ and all $t \in (0, 1/2]$, the inequality $\mathbb{P}(\|\mathbf{X}\|_2 < t) \leq C t^\beta$ holds.

Assumption A1 is the so-called strong density assumption (Audibert and Tsybakov, 2007) that imposes two conditions on the distribution of the feature vector $\mathbf{X}$. In particular, A1(a) requires that the support does not contain any isolate points and A1(b) assumes the probability density of $\mathbf{X}$ is bounded above and below in its support, as commonly required in the literature (e.g., Huang, 1998, 2003). Assumption A2 is the uniform Lipschitz condition imposed on the conditional probability $f(\mathbf{x})$, and similar conditions were imposed in Cai and Wei (2019), Audibert and Tsybakov (2007), Gadat et al. (2016). Assumption A3 is a popular condition in classification problems (e.g., see Audibert and Tsybakov, 2007, Gadat et al., 2016), which characterizes the strength of the signal $\|\mathbf{X}\|_2$. With a larger β , $\mathbf{X}$ is near the decision boundary $\{-1, 2\}$ with a lower probability, leading to an easier classification problem.

3.2 Theoretical Results in General Setting

In this section, we first present some theoretical results on the distributed adaptive NN classifier in a general setting where sub-sample sizes (i.e., n_j 's) are allowed to be different. The following theorem gives an upper bound of the regret of the proposed classifier in Algorithm 1.

Theorem 1. Under Assumptions A1-A3 and $\min_{j=1:m} n_j \geq N^1$ for some $\gamma < \frac{2}{2+d}$. It follows that

$$\mathcal{R}(f_{k_1:k_m}) \leq [N \log(N)]^{\frac{(1+\gamma)}{2+d}}$$

The proof is given in the Appendix.

Theorem 1 establishes the convergence rate of the proposed classifier when sub-sample sizes are not too small, i.e., $\min_j n_j \geq N^1$ for some $\gamma < 2/(2+d)$. We remark that this convergence rate coincides with the minimax lower bound given in Audibert and Tsybakov (2007) up to a logarithm factor. The additional $\log(N)$ term is

the price to pay for the adaptive choice of tuning parameters $k_1, \dots, k_m$, as commonly seen in the literature (e.g., see [Lepskii, 1991](#), [Lepski and Spokoiny, 1997](#)).

To shed more lights on this issue, we consider a distributed NN classifier with a non-stochastic choice of tuning parameter satisfying $k_j \leq n_j N^{-\frac{d}{2+d}}$, $j = 1, \dots, m$, which is essentially an extension of [Qiao et al. \(2019\)](#) which only considered the case $n_1 = \dots = n_m$. The following theorem gives an upper bound of the regret of the resulting distributed NN classifier given in (1).

Theorem 2. (Non-adaptive k_j 's) Suppose that Assumptions A1-A3 hold and $\min_{j=1, \dots, m} n_j \geq N^{\frac{1}{2+d}}$ for some $\frac{1}{2} < \frac{2}{2+d}$. Then if $k_j \leq n_j N^{-\frac{d}{2+d}}$ for $j = 1, \dots, m$, it follows that

$$\mathcal{R}(f_{k_1:k_m}) \leq N^{-\frac{(1+\frac{1}{2+d})}{2+d}}$$

The proof is given in the Appendix.

Theorem 2 asserts that if k_j 's are not chosen by a data-driven method, the minimax lower bound of the regret ([Audibert and Tsybakov, 2007](#)) is achieved by the distributed NN classifier provided that $k_j = C_j(n_j N^{-\frac{d}{2+d}})$ for some constant $C_j > 0$, $j = 1, \dots, m$. Although Theorem 2 is of limited practical interest since it is difficult to determine the values of C_j 's and $\frac{1}{2}$ for a given data set, it indeed motivates us to propose the early stopping threshold $n_1 N^{-\frac{d}{2+d}} \log(N)$ when searching for the optimal k_1 in Algorithm 1, which resulted in an extra $\log(N)$ term in its regret convergence rate as suggested by Theorem 1.

Even though the convergence rates in Theorems 1 and 2 look similar, their proofs rely on completely different techniques. Since $k_1, \dots, k_m$ are deterministic in Theorem 2, the regret of $f_{k_1:k_m}$ can be established through calculating its bias and variance. However, when $k_1, \dots, k_m$ are data-driven, the regret of $f_{k_1:k_m}$ requires more sophisticated analysis. One major difficulty, for instance, is quantifying the model complexity, which relies on the following lemma.

Lemma 1. Given observations $\mathbf{X}_1^1, \dots, \mathbf{X}_{n_1}^1, \dots, \mathbf{X}_1^m, \dots, \mathbf{X}_{n_m}^m$, for $k_j = 1, \dots, n_j$ with $j = 1, \dots, m$, we define sets

$$\begin{aligned} \mathcal{A}_{k_j:j}(\mathbf{x}) &:= \mathbf{X}_{(1)}^j(\mathbf{x}), \dots, \mathbf{X}_{(k_j)}^j(\mathbf{x}) \\ \mathcal{B} &:= \mathcal{B}(k_1, \dots, k_m) = \{\mathcal{A}_{k_1:1}(\mathbf{x}), \dots, \mathcal{A}_{k_m:m}(\mathbf{x}) : \mathbf{x} \in [0, 1]^d\} \end{aligned}$$

Then the cardinality of $\mathcal{B}$ is bounded by dN^d .

The proof is given in the Appendix.

Lemma 1 counts the number of sets of the form $\mathcal{A}_{k_1:1}(\mathbf{x}), \dots, \mathcal{A}_{k_m:m}(\mathbf{x})$ when $\mathbf{x}$ is running over $[0, 1]^d$. It shows that this number is upper bounded by dN^d . This is a generalization of Lemma 3 in [Jiang \(2019\)](#) from $m = 1$ to $m > 1$. The selection of $k_1, \dots, k_m$ can be viewed as a model selection problem with $n_1, \dots, n_m$ candidate models, and the complexity of each model is measured by $\mathcal{B}(k_1, \dots, k_m)$. The proof of Theorem 1 requires controlling the complexity of all the candidate models. If we do not specify any constraints on the k_j 's and allow for all the combinations of $k_1, \dots, k_m$, then the complexity of all the candidates models can be evaluated by the following:

$$\sum_{k_m=1}^{n_m} \sum_{k_1=1}^{n_1} \mathcal{B}(k_1, \dots, k_m) \leq dN^d \sum_{k_1=1}^{n_1} \sum_{k_m=1}^{n_m}$$

which is relatively large. As a matter of fact, if we impose a restriction that $k_j = k_1 n_j / n_1$ for all $j = 1, \dots, m$, then we only need to conduct model selection among n_1 models, and the corresponding complexity can be bounded by

$$\sum_{k_1=1}^{n_1} \sum_{k_2=1}^{k_1 n_2 / n_1} \dots \sum_{k_m=1}^{k_1 n_m / n_1} \mathcal{B}(k_1, \dots, k_m) \leq d N^d / n_1$$

This reduced complexity plays an important role in deriving the near optimal rate in Theorem 1, and it also motivates the choice $k_j = k_1 n_j / n_1$ in Algorithm 1.

3.3 Theoretical Results with $n_1 = \dots = n_m$

Theorem 1 is limited to the case when $\min_{1 \leq j \leq m} n_j \geq N^{1/(2+d)}$ for some $\alpha < 2/(2+d)$, where it asserts that the optimal convergence rate (up to a factor of $\log(N)$) of the regret can be achieved by the proposed classifier. However, theoretical properties of the proposed classifier are unclear when $\min_{1 \leq j \leq m} n_j \leq N^{1/(2+d)}$ only holds for some

$\alpha < 2/(2+d)$. While it is difficult to study in general, we manage to provide a partial answer by considering the special case $n_1 = \dots = n_m$, which has been widely studied under the so-called divide-and-conquer framework (Qiao et al., 2019, Duan et al., 2020, Zhang et al., 2015, Shang and Cheng, 2017, Xu et al., 2018, Shang et al., 2019, Xu et al., 2019).

Theorem 3. Suppose that Assumptions A1-A3 hold and that $n_1 = \dots = n_m = n \geq N^{1/(2+d)}$ for some $\alpha \in [0, 1)$, then it holds that (a) if $\alpha < \frac{2}{2+d}$, then $\mathcal{R}(f_{k_1:k_m}) \leq [N \log(N)]^{\frac{(1+\alpha)}{2+d}}$; (b) if $\alpha \geq \frac{2}{2+d}$, then $\mathcal{R}(f_{k_1:k_m}) \leq [\log(N)]^{1-\alpha} [N \log(N)]^{\frac{(1-\alpha)(1+\alpha)}{d}}$ for some $\beta > 0$.

The proof is given in the Appendix.

Theorem 3 characterizes the asymptotic behavior of the proposed classifier in two scenarios. When $\alpha < 2/(2+d)$, part (a) is a special case of Theorem 1, where the regret convergence rate is free of α and is nearly optimal up to a logarithm factor (Audibert and Tsybakov, 2007). However, when $\alpha \geq 2/(2+d)$, each sub-sample has a smaller sample size, and the resulting convergence rate of the regret becomes $[\log(N)]^{1-\alpha} [N \log(N)]^{\frac{(1-\alpha)(1+\alpha)}{d}}$ for some constant $\beta > 0$, which slows down when α increases. In contrast, the convergence rate in part (a) remains the same as α changes.

It is unclear whether the convergence rate given in part (b) is optimal since existing literature on distributed NN classifier has mainly focused on the case with $\alpha < 2/(2+d)$ (e.g. Qiao et al., 2019). However, we can show that the convergence rate in part (b) is closely related to that of the distributed 1-NN classifier, as given in the following theorem.

Theorem 4. Suppose that Assumptions A1-A3 hold and that $n_1 = \dots = n_m = n \geq N^{1/(2+d)}$ for some $\alpha \in [0, 1)$. Then if $\alpha \geq 2/(2+d)$ and fixing $k_1 = \dots = k_m = 1$, it holds that $\mathcal{R}(f_{k_1:k_m}) \leq [\log(N)]^{1-\alpha} [N \log(N)]^{\frac{(1-\alpha)(1+\alpha)}{d}}$ for some $\beta > 0$.

The proof is given in the Appendix.

Theorem 4 shows that the distributed 1-NN classifier can achieve the same convergence rate as the proposed adaptive NN classifier when $\alpha \geq 2/(2+d)$. This makes intuitive sense because when α is large, the aggregated classifier (1) averages

over a large number of NN classifiers built on sub-samples (i.e., $m = N/n \rightarrow N$) and the overall variability of the resulting aggregated NN classifier can be significantly smaller than its prediction bias, which is of the same magnitude of individual NN classifiers from sub-samples. Consequently, to improve the prediction accuracy of the aggregated NN classifier, it is desirable to use the 1-NN classifier for each sub-sample, which has the smallest prediction bias among NN classifiers for a given sample size.

The similarity between Theorem 3 part (b) and Theorem 4 suggests that when $2 \leq (2 + d)$, the proposed classifier behave similarly to the distributed 1-NN classifier. This conjecture is supported by our simulation studies in Section 4.1 not only in the case where $n_1 = \dots = n_m$ but also in the case where sub-sample sizes are not equal. However, the distributed 1-NN classifier performs much worse than the proposed classifier when d is small. One advantage of the proposed classifier is that it can automatically adjust to both scenarios without the knowledge of the true value of d .

4 Numerical Results

4.1 Simulation Studies

In this section, we evaluate the finite sample performance of the proposed algorithm. The following marginal distributions of $\mathbf{X}$ will be considered.

- (a) $\mathbf{X} \sim g_1(\mathbf{x})$: $\mathbf{X} = (X_1, X_2, X_3) \in [0, 1]^3$ with $X_1 = R \cos(\theta_1) \cos(\theta_2)$, $X_2 = R \cos(\theta_1) \sin(\theta_2)$, and $X_3 = R \sin(\theta_1)$. Here $\theta_1, \theta_2 \sim \text{Unif}(0, 2\pi)$, and $R \sim \text{Unif}(0, 1)$ are three independent uniform random variables.
 - (b) $\mathbf{X} \sim g_2(\mathbf{x})$: $\mathbf{X} = (X_1, X_2, X_3) \in [0, 1]^3$ is generated by a similar process as (a) except $R \sim 0.5\text{Beta}(5, 1) + 0.5\text{Beta}(1, 6)$ follows a Beta mixture distribution.
- Given $\mathbf{X} = \mathbf{x}$, the conditional probability function is $p(\mathbf{z} | \mathbf{x}) = h(\mathbf{z} | \mathbf{x})$, where

$$h(z) = \begin{cases} 0.8 & \text{if } 0 \leq z \leq 0.3 \\ 6z + 2.6 & \text{if } 0.3 < z \leq 0.4 \\ 0.2 & \text{if } 0.4 < z \leq 0.7 \\ 2.6z - 1.62 & \text{if } 0.7 < z \leq 0.8 \\ 0.46 & \text{if } 0.8 < z \leq 1 \end{cases}$$

The total sample size is set as $N = 60000$, and the data are randomly divided into $m = N/n$, $n = 0.01, \dots, 0.8$, sub-samples by the following two approaches:

I. Equally Splitting: The N observations are split into m datasets with (roughly) equal sample size.

II. Unequal Splitting: The N observations are split into m datasets, and the sample sizes $(n_1, \dots, n_m)$ follow a multinomial distribution with probabilities $(p_1, \dots, p_m)$ for $s = 1, \dots, m$.

For comparison purpose, we consider the following classifiers:

DAES: The proposed distributed adaptive NN classifier in Algorithm 1 with an early stopping bound $n_1 N^{\frac{d}{2+d}}$;

DA: Modified Algorithm 1, where the early stopping bound is replaced by n_1 ;

DK The distributed NN classifier (Qiao et al., 2019) with $k_j = \lfloor n_j N^{\frac{d}{2+d}} \rfloor$, $j = 1, \dots, m$.

D1: The distributed 1-NN classifier by setting $k_1 = \dots = k_m = 1$ in (1).

For DK in the unequal splitting case, we use $k_j = \lfloor n_j N^{\frac{d}{2+d}} \rfloor$ for $j = 1, \dots, m$, as suggested by our Theorem 2. Such a choice reduces to $k = \lfloor n N^{\frac{d}{2+d}} \rfloor$ when $n_1 = \dots = n_m = n$, which is the choice adopted by Qiao et al. (2019). For each simulation run, the above four classifiers are trained using m sub-samples to predict the label of a new feature $\mathbf{x}$ randomly generated from the marginal distribution of $\mathbf{X}$. To evaluate the classification accuracy, we treat the Bayesian classifier $f^*(\mathbf{x}) = \mathbb{I}(\langle \mathbf{x}, \mathbf{g}_2 \rangle \geq 0.5)$ as the golden rule and calculate the percentage of times a classifier gives the same prediction as the Bayesian classifier. The average computation times (measured in second and taking log) of DA and DAES with different α are also recorded. To investigate the role of the early stopping rule, we also compare the numbers of neighbor (k_1) chosen by DA and DAES. Summary statistics based on 200 simulation runs are reported in Figures 1-6.

First, Figures 1 and 2 suggest that the proposed DAES classifier has a better overall performance than DK. In particular, their classification accuracies are practically the same for $\alpha = 0.5$, while the proposed DAES performs significantly better than DK when $\alpha = 0.3$ and $\mathbf{X} \sim g_2(\mathbf{x})$, which demonstrates the benefits of searching for an optimal k using a data-driven Algorithm 1.

A second observation from Figures 1 and 2 is that the DAES classifier appears to be consistently inferior to the DA classifier. This highlights the importance of imposing an early stopping bound $\lfloor n_1 N^{\frac{d}{2+d}} \rfloor$ when searching for the optimal k_1 . This can be explained by the fact that searching for k_1 from 1 to n_1 may introduce too much uncertainty in the choice of k_1 (as well as other k_j 's), which may, in turn, results in greater variability for the final aggregated NN classifier. This explanation also can be supported by Figures 3-6. For example, Figure 3 shows that the k_1 chosen by DA is generally larger than that chosen by DAES. When $\alpha = 0$, DA could choose a k_1 larger than 10000 given $N = 60000$, which may increase a lot of uncertainty for classification.

Third, Figures 1 and 2 also indicate that the proposed DAES classifier performs similarly to the D1 classifier when α is large, supporting our theoretical findings in Theorems 3-4. However, the D1 classifier performs much worse than the DAES classifier when α is small, demonstrating the advantage of the proposed DAES classifier due to its adaptivity in choosing an optimal k_1 (as well as other k_j 's).

Finally, Figures 1 and 2 show that for each different marginal distributions of $\mathbf{X}$, the DAES classifier outperforms the DA classifier in terms of computation time, both of which are U-shaped functions with respect to α and attain the minimal when α is around 0.5. For small α , a large proportion of the computation time is spent on choosing $k_1 \sim k_m$. However, when α is large, the main computational cost is to aggregate the sub-samples, resulting in increased run time as α continues to increase. All numerical studies are conducted via High Performance Computing Center at Texas Tech University.

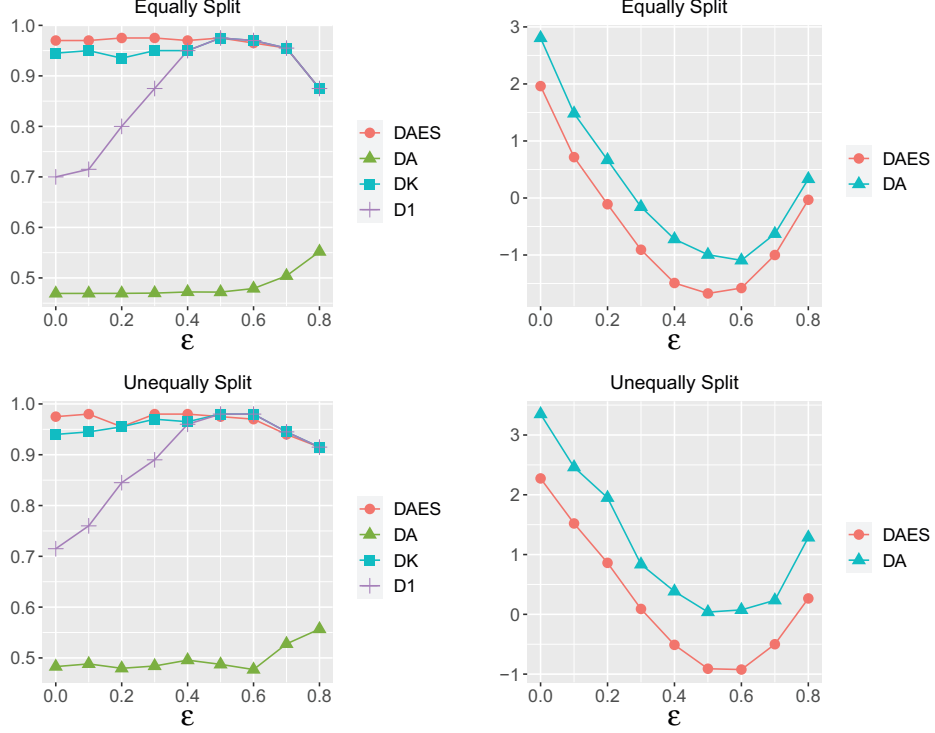


Fig. 1 Classification accuracy and computation time with $\mathbf{X} \sim g_1(\mathbf{x})$.

4.2 A Real Data Analysis

In this section, we apply the four classifiers in Section 4.1 to the adult income dataset from UCI Machine Learning Repository (Dua and Graff, 2017). The goal is predict whether a person makes over 50K a year. After removing missing values, we retain 32561 observations and use age, final weight, education, capital gain, capital loss and weekly working hours as the feature vector. The whole data is divided into a training dataset with 26049 observations (about 80%) and a testing dataset with sample size 6512 (about 20%). We use the same settings in Section 4.1 to evaluate the prediction error of the testing dataset. The results are summarized in Figure 7. Overall, our proposed algorithm DAES has the best performance under various choices of ϵ . Moreover, compared with DA, our estimator DAES significantly speeds up the computation when $\epsilon \leq 0.5$.

5 Conclusion

In this work, we study the binary classification problem in the big data setting, and propose a distributed adaptive NN classifier with the tuning parameter being selected by a data-driven criterion. Under mild conditions, we prove the proposed classifier

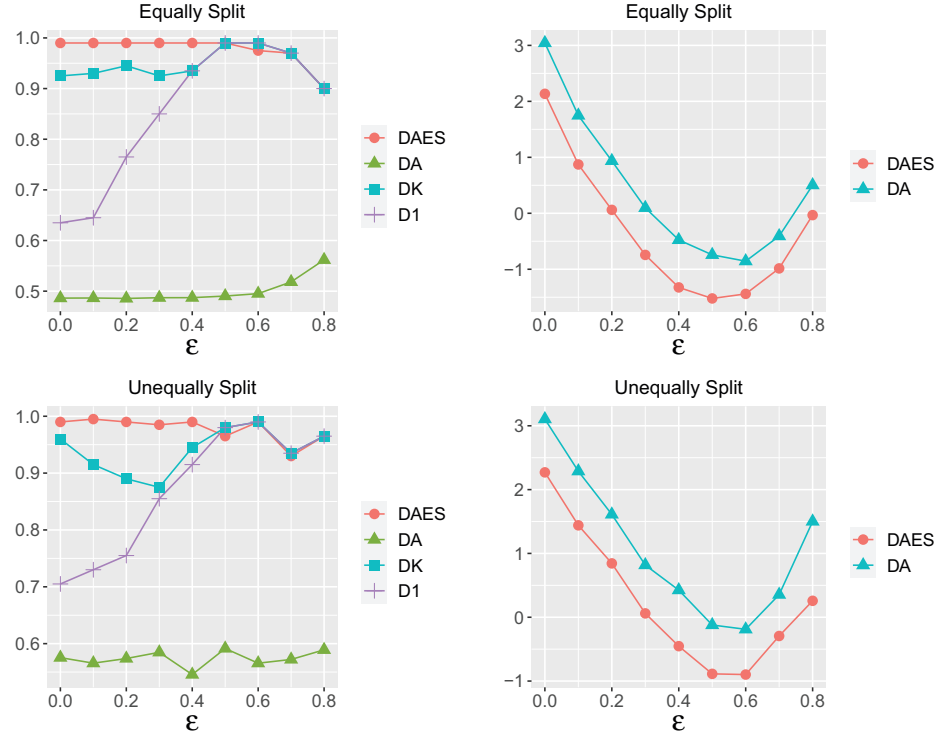


Fig. 2 Classification accuracy and computation time with $\mathbf{X} \sim g_2(\mathbf{x})$.

can achieve the minimax optimal rate of excess risk. Numerical results demonstrate its effectiveness and efficiency.

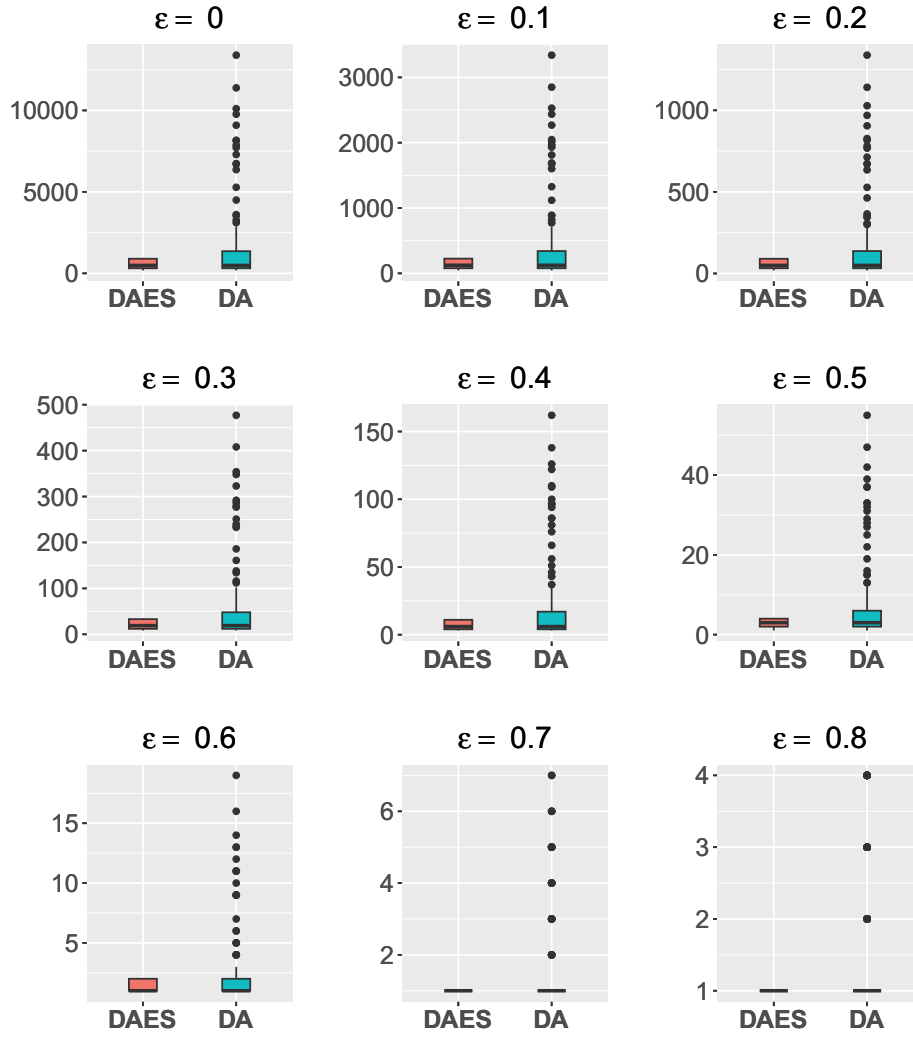


Fig. 3 Selected k_1 with $\mathbf{X} \sim g_1(\mathbf{x})$ and equally split sub-samples.

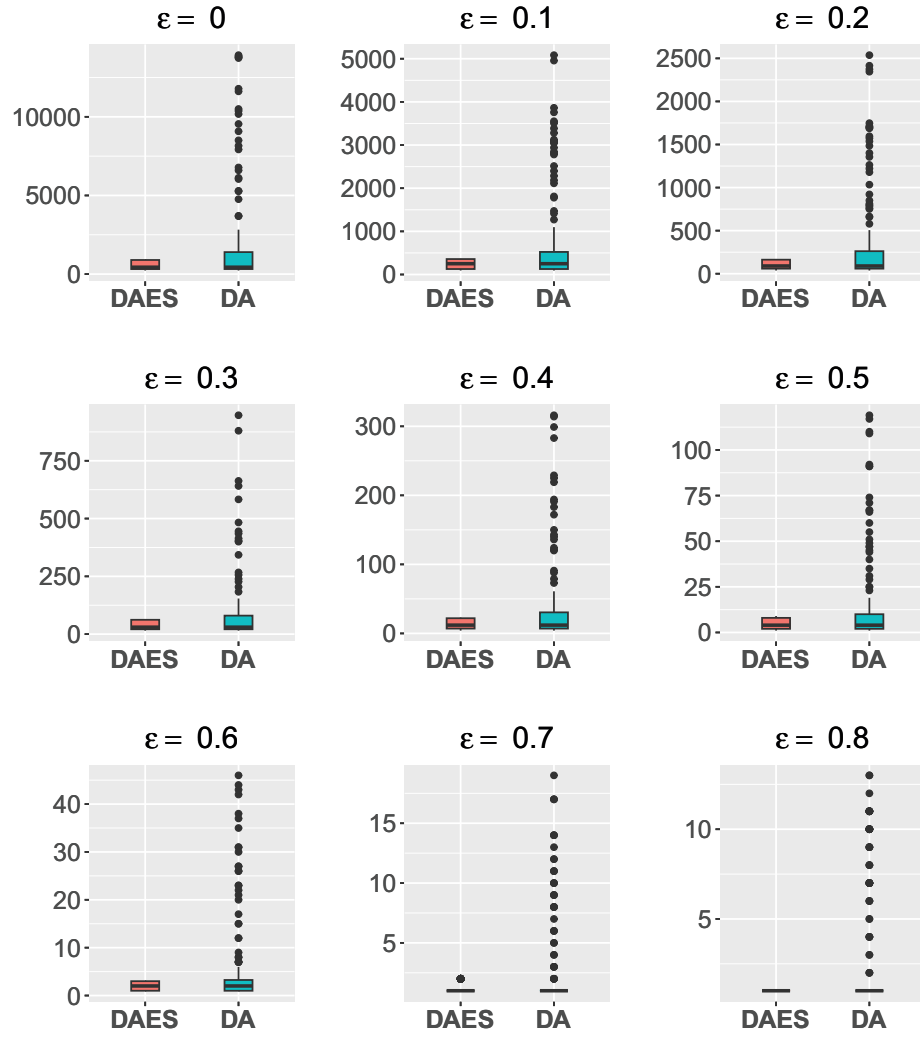


Fig. 4 Selected k_1 with $\mathbf{X} \sim g_1(\mathbf{x})$ and unequally split sub-samples.

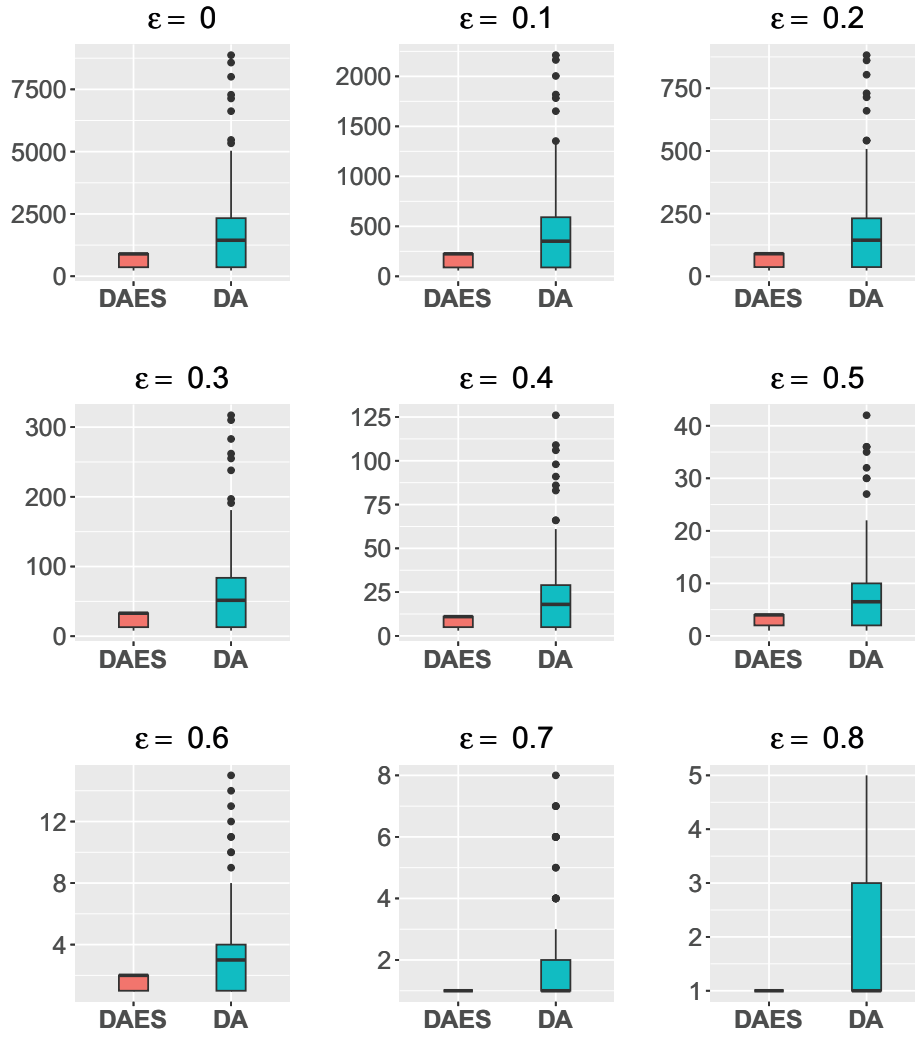


Fig. 5 Selected k_1 with $\mathbf{X} \sim g_2(\mathbf{x})$ and equally split sub-samples.

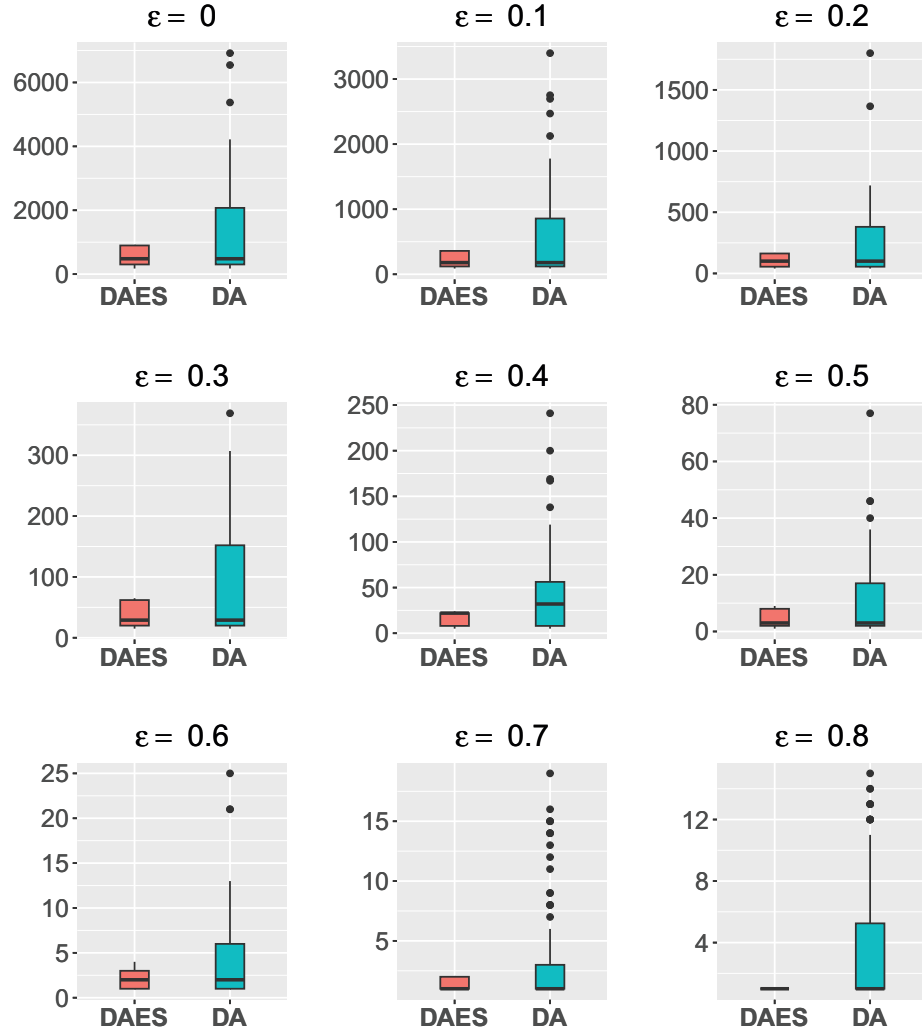


Fig. 6 Selected k_1 with $\mathbf{X} \sim g_2(\mathbf{x})$ and unequally split sub-samples.

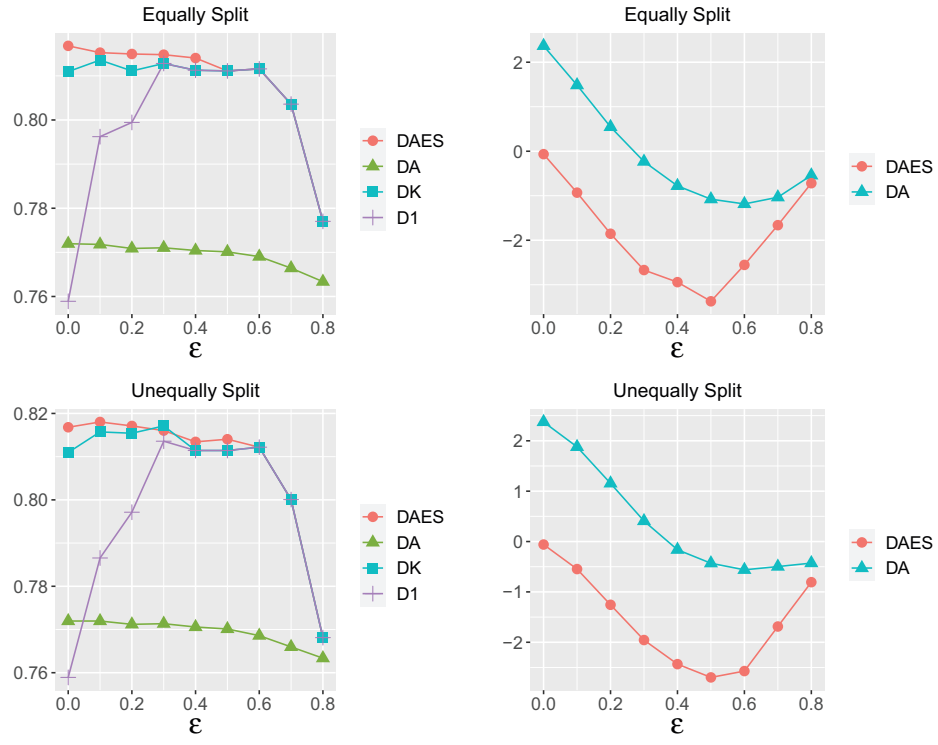


Fig. 7 Classification accuracy and computation time for adult income dataset.

Acknowledgments. The authors would like to express sincere appreciation to Editor Dr. Ricardo Henao and the two anonymous reviewers for their valuable and insightful comments.

Appendix A Mathematical Proofs

In this Appendix, we provide the mathematical proofs of the theorems and relevant lemmas.

We denote $\mathcal{X} = \{\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_N\}$ as the collection of all covariates. For $k_j = 1, \dots, n_j$ with $j = 1, \dots, m$ and $0 < a < 1$, we define events

$$E_j(k_j - a) = \{X_{(k_j)}^j(\mathbf{x}) - \mathbf{x} \leq C_D \frac{k_j}{n_j^{1-a}}\} \text{ for all } \mathbf{x}$$

and

$$E_P(k_1 : k_m - a) = \bigcap_{j=1}^m E_j(k_j - a)$$

Sometime we may write $E_P(k_1 : k_m - a)$ as E_P if there is no confusion in the context. By Lemma 2 below, it follows that

$$\mathbb{P}(E_P(k_1 : k_m - a)) \geq 1 - C_D \prod_{j=1}^m \frac{n_j^{1-a}}{k_j} \exp(-n_j^a k_j^{-6}) \quad (\text{A1})$$

A1 Preliminary Lemmas

Lemma 2. *There exist $C_D > 0$ such that for all $a \in [0, 1)$, $k_j = 1, \dots, n_j$ and $j = 1, \dots, m$, the following holds with probability at least $1 - C_D \frac{n_j^{1-a}}{k_j} e^{-n_j^a k_j^{-6}}$:*

$$X_{(k_j)}^j(\mathbf{x}) - \mathbf{x} \leq C_D \frac{k_j}{n_j^{1-a}} \text{ for all } \mathbf{x}$$

Moreover, with probability at least $1 - C_D \frac{n_j}{k_j} e^{-k_j^{-6}}$, it also holds that

$$X_{(k_j)}^j(\mathbf{x}) - \mathbf{x} \geq \frac{1}{C_D} \frac{k_j}{n_j} \text{ for all } \mathbf{x}$$

Proof. The proofs of the upper bound and lower bound are almost the same. In the following, we prove the upper bound. For simplicity, we will omit the index j .

Let $B(\mathbf{x}, r)$ be the ball centered at $\mathbf{x}$ with radius r . By Assumption A1, therefore we have

$$\mathbb{P}(\mathbf{X} \in B(\mathbf{x}, r)) = \int_{B(\mathbf{x}, r)} \frac{dP_{\mathbf{X}}(\mathbf{x})}{d}(\mathbf{x}) d\mathbf{x} \leq c \cdot (B(\mathbf{x}, r))$$

$$c^2 (B(\mathbf{x}, r)) = c^2 (B(0, 1))r^d$$

For simplicity, we denote $c = c^2 (B(0, 1))$, so $\mathbb{P}(\mathbf{X} \in B(\mathbf{x}, r)) = cr^d$. Let $r = c_r(k n^{1-a})^{\frac{1}{d}}$ for some $0 < a < 1$ and $c_r = (2/c)^{\frac{1}{d}}$. Moreover, we define $S(\mathbf{x}) = \sum_{i=1}^n \mathbb{I}(\mathbf{X}_i^j \in B(\mathbf{x}, r))$ and $W \sim \text{Binomial}(n, cr^d)$. Hence, Bernstein's inequality implies that

$$\begin{aligned} \mathbb{P}(S(\mathbf{x}) < k) &= \mathbb{P}(W < k) = \mathbb{P}(W - \mathbb{E}(W) < k - \mathbb{E}(W)) \\ &= \mathbb{P}(W - \mathbb{E}(W) < k - cnr^d) \\ &= \mathbb{P}(W - \mathbb{E}(W) < k - cc_r^d n^a k) \\ &= \mathbb{P}(W - \mathbb{E}(W) < k - 2n^a k) \end{aligned}$$

$$\mathbb{P}(W - \mathbb{E}(W) < -n^a k)$$

$$\exp\left(-\frac{3n^a k}{14}\right) \leq \exp(-n^a k/6)$$

where we use the fact that $a < 0$. Let $\mathcal{B}$ be a finite set such that $\mathbf{x} \in \mathcal{B} \Rightarrow B(\mathbf{x}, r)$, and we can verify $|\mathcal{B}| \leq Cr^{-d}$ for some $C > 0$. As a consequence, we have

$$\mathbb{P}(\mathbf{x} \in \mathcal{B} \mid S(\mathbf{x}) < k) \leq Cr^{-d} \exp(-n^a k/6) \leq Cc_r^d \frac{n^{1-a}}{k} \exp(-n^a k/6)$$

For any $\mathbf{x}$, there is a $\mathbf{x} \in \mathcal{B}$ such that $\|\mathbf{x} - \mathbf{x}\| \leq 2r$. Under the event $E_2 = \{\mathbf{x} \in \mathcal{B} \mid S(\mathbf{x}) < k\}$, there are at least k covariates among $\mathbf{X}_1^j, \dots, \mathbf{X}_n^j$ in the ball $B(\mathbf{x}, r)$, and thus there are at least k covariates among $\mathbf{X}_1^j, \dots, \mathbf{X}_n^j$ in the ball $B(\mathbf{x}, 2r)$. Hence, we have

$$\mathbb{P}(\mathbf{x} \in \mathcal{B} \mid \mathbf{X}_{(k)}^j(\mathbf{x}) \leq \mathbf{x} - 2r) \leq \mathbb{P}(E_2) \leq 1 - Cc_r^d \frac{n^{1-a}}{k} \exp(-n^a k/6)$$

□

Lemma 3. Fixing $k_1 = 1 \leq n_1$ and setting $k_j = \lfloor k_1 n_j / n_1 \rfloor$ with $j = 1, \dots, m$, there exist $c_b, C_b > 0$ free of k_j such that

$$\begin{aligned} \mathbb{E}(\mathbf{1}_{\{k_j \leq j\}}(\mathbf{x}) \mid \mathcal{X}) &\leq \frac{1}{2} - c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 1 \\ \mathbb{E}(\mathbf{1}_{\{k_j \leq j\}}(\mathbf{x}) \mid \mathcal{X}) &\leq \frac{1}{2} - c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 0 \end{aligned}$$

holds for all $\mathbf{x}$ with $\|\mathbf{x}\| \leq C_b \|\mathbf{X}_{(k_j)}^j(\mathbf{x}) - \mathbf{x}\|$. Moreover, if $k_1 n_j \leq n_1$ for all $j = 1, \dots, m$, then the following statements hold on event $E_P(k_1 : k_m = 0)$:

$$\begin{aligned} \mathbb{E}(\mathbf{1}_{\{k_1 : k_m\}}(\mathbf{x}) \mid \mathcal{X}) &\leq \frac{1}{2} - c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 1 \\ \mathbb{E}(\mathbf{1}_{\{k_1 : k_m\}}(\mathbf{x}) \mid \mathcal{X}) &\leq \frac{1}{2} - c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 0; \end{aligned}$$

for all $\mathbf{x}$ with $\|\mathbf{x}\| \leq C_b(k_1 - n_1)^{\frac{1}{2}}$. In addition, if $k_1 = \dots = k_m = k$ and $n_1 = \dots = n_m = n$, then for any $a \in [0, 1]$, the following statements hold on event $E_P(k_1 : k_m - a)$:

$$\begin{aligned}\mathbb{E}(\|k_1:k_m(\mathbf{x})\|_{\mathcal{X}}) &= \frac{1}{2} C_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 1 \\ \mathbb{E}(\|k_1:k_m(\mathbf{x})\|_{\mathcal{X}}) &= \frac{1}{2} C_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 0;\end{aligned}$$

for all $\mathbf{x}$ with $\|\mathbf{x}\| \leq C_b(k - n^{\frac{1}{2}})^{\frac{1}{2}}$.

Proof. Since $\mathbb{E}(Y_{(i)}^j(\mathbf{x}) | \mathcal{X}) = \langle \mathbf{X}_{(i)}^j(\mathbf{x}), \mathbf{x} \rangle$, by Assumption A2, we show that

$$\begin{aligned}\mathbb{E}(\|k_j(\mathbf{x})\|_{\mathcal{X}}) &= \frac{1}{k_j} \sum_{i=1}^{k_j} [\mathbb{E}(Y_{(i)}^j(\mathbf{x}) | \mathcal{X})] \\ &= \frac{1}{k_j} \sum_{i=1}^{k_j} \langle \mathbf{X}_{(i)}^j(\mathbf{x}), \mathbf{x} \rangle \\ &= \frac{C}{k_j} \sum_{i=1}^{k_j} \langle \mathbf{X}_{(i)}^j(\mathbf{x}), \mathbf{x} \rangle \\ &= C \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle\end{aligned}$$

Therefore, choosing $C_b = 2C$ and if $\|\mathbf{x}\| \leq C_b \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle = 2C \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle$ and $f(\mathbf{x}) = 1$, then we have

$$\begin{aligned}\mathbb{E}(\|k_j(\mathbf{x})\|_{\mathcal{X}}) &= \frac{1}{2} C_b(\mathbf{x}) = \frac{1}{2} \mathbb{E}(\|k_j(\mathbf{x})\|_{\mathcal{X}}) \\ &= C \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle = \frac{1}{2} C_b(\mathbf{x})\end{aligned}$$

So the statement will hold for $C_b = 1/2$ and $C_b = 2C$. Similarly, we can prove the case when $\|\mathbf{x}\| \leq C_b \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle$ and $f(\mathbf{x}) = 0$. Consequently, on event $E_P(k_1 : k_m - 0)$, we have

$$\begin{aligned}\mathbb{E}(\|k_1:k_m(\mathbf{x})\|_{\mathcal{X}}) &= \frac{1}{\sum_{j=1}^m k_j} \sum_{j=1}^m \sum_{i=1}^{k_j} [\mathbb{E}(Y_{(i)}^j(\mathbf{x}) | \mathcal{X})] \\ &= \frac{1}{\sum_{j=1}^m k_j} \sum_{j=1}^m \sum_{i=1}^{k_j} \langle \mathbf{X}_{(i)}^j(\mathbf{x}), \mathbf{x} \rangle \\ &= \frac{C}{\sum_{j=1}^m k_j} \sum_{j=1}^m \sum_{i=1}^{k_j} \langle \mathbf{X}_{(i)}^j(\mathbf{x}), \mathbf{x} \rangle \\ &= \frac{C}{\sum_{j=1}^m k_j} \sum_{j=1}^m k_j \langle \mathbf{X}_{(k_j)}^j(\mathbf{x}), \mathbf{x} \rangle\end{aligned}$$

$$\frac{C - C_D}{\prod_{j=1}^m k_j} \prod_{j=1}^m k_j \frac{k_j}{n_j}^\alpha$$

$$\frac{C - C_D}{\prod_{j=1}^m k_j} \prod_{j=1}^m k_j \frac{2k_1}{n_1}^\alpha = 2^\alpha C - C_D \frac{k_1}{n_1}^\alpha$$

where the condition $k_j = k_1 n_j / n_1 = 2k_1 n_j / n_1$ is used. For $C_b = 2^{1+\alpha} C - C_D$ and $\mathbf{x}$ such that $\phi(\mathbf{x}) = C_b(k_1 / n_1)^\alpha$, $f(\mathbf{x}) = 1$, it holds that

$$\mathbb{E}(\prod_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) = \frac{1}{2} \quad \phi(\mathbf{x}) = \frac{1}{2} \quad \mathbb{E}(\prod_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) = \phi(\mathbf{x})$$

$$\phi(\mathbf{x}) = 2^\alpha C - C_D \frac{k_1}{n_1}^\alpha$$

$$C_b \frac{k_1}{n_1}^\alpha = \frac{C_b}{2} \frac{k_1}{n_1}^\alpha = \frac{C_b}{2} \frac{k_1}{n_1}^\alpha$$

Finally, choosing $c_b = C_b / 2$, we complete the proof of the second statement. Similarly, we can prove the case when $\phi(\mathbf{x}) = C_b(k_1 / n_1)^\alpha$ and $f(\mathbf{x}) = 0$.

The proof of the third statement is similar to the second one. Hence, we omit it. $\square$

Lemma 4. Let c_b and C_b be the constants in Lemma 3. Fixing $k_1 = 1$ and $n_1 = n$ and setting $k_j = k_1 n_j / n_1$ with $j = 1, \dots, m$, if $k_1 n_j = n_1$ for all $j = 1, \dots, m$, then the following holds on event $E_P(k_1 : k_m, 0)$:

$$\mathbb{P}(f_{k_1:k_m}(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) = \exp \left(-2c_b^2 \prod_{j=1}^m k_j^{-2}(\mathbf{x}) \right)$$

for all $\mathbf{x}$ with $\phi(\mathbf{x}) = C_b(k_1 / n_1)^\alpha$. In addition, if $k_1 = \dots = k_m = k$ and $n_1 = \dots = n_m = n$, then for any $a \in [0, 1]$, the following statements hold on event $E_P(k_1 : k_m, a)$:

$$\mathbb{P}(f_{k_1:k_m}(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) = \exp \left(-2c_b^2 m k^{-2}(\mathbf{x}) \right)$$

for all $\mathbf{x}$ with $\phi(\mathbf{x}) = C_b(k / n^{1-\alpha})^\alpha$.

Proof. Suppose $f(\mathbf{x}) = 1$ and $\phi(\mathbf{x}) = C_b(k_1 / n_1)^\alpha$. By Lemma 3, under event $E_P(k_1 : k_m, 0)$, we have

$$\mathbb{E}(\prod_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) = \frac{1}{2} \quad \phi(\mathbf{x}) = c_b \quad \phi(\mathbf{x}) \tag{A1}$$

Furthermore, we observe the that

$$f_{k_1:k_m}(\mathbf{x}) = \frac{1}{m} \prod_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x})$$

and $Y_{(1)}^1(\mathbf{x}), \dots, Y_{(k_1)}^1(\mathbf{x}), \dots, Y_{(1)}^m(\mathbf{x}), \dots, Y_{(k_m)}^m(\mathbf{x})$ are independent conditional on $\mathcal{X}$. Hence, it follows from Hoeffding's inequality and (A1) that

$$\begin{aligned} & \mathbb{P}(f_{k_1:k_m}(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) \\ &= \mathbb{P}(f_{k_1:k_m}(\mathbf{x}) = 0 | \mathcal{X}) \\ &= \mathbb{P}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) < \frac{1}{2} \mid \mathcal{X}\right) \\ &= \mathbb{P}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) - \mathbb{E}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) \mid \mathcal{X}\right) < \frac{1}{2} - \mathbb{E}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) \mid \mathcal{X}\right) \mid \mathcal{X}\right) \\ &\leq \mathbb{P}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) - \mathbb{E}\left(\sum_{j=1}^m k_j^{k_j} Y_{(i)}^j(\mathbf{x}) \mid \mathcal{X}\right) < -c_b(\mathbf{x}) \mid \mathcal{X}\right) \\ &\leq \exp\left(-\frac{2c_b^2(\mathbf{x})}{\sum_{j=1}^m k_j^{k_j}}\right) \end{aligned}$$

Using similar argument, we can prove the case when $f(\mathbf{x}) = C_b(k_1, \dots, k_n)^{-\alpha}$ and $f(\mathbf{x}) = 0$. $\square$

A2 Proof of Lemma 1

Let $\mathcal{H}$ be the partition of $[0, 1]^d$ induced by $\binom{n}{2}$ hyperplanes defined as the perpendicular bisectors of each pair of points $(\mathbf{X}_s^j, \mathbf{X}_p^j)$ for $1 \leq s < p \leq n$ and $j = 1, \dots, m$ (see Figure 3 for the case with $m = 2, k_1 = 3, k_2 = 2$). If $\mathbf{x}$ and $\mathbf{x}'$ are in the same partition, then $A_{k_j, j}(\mathbf{x}) = A_{k_j, j}(\mathbf{x}')$ for all $j = 1, \dots, m$ (see Figures 1 and 2). As a consequence, the cardinality $|\mathcal{B}| = |\mathcal{H}|$. Now consider $\mathcal{H}'$ to be the partition of $[0, 1]^d$ induced by $\binom{N}{2}$ hyperplanes defined as the perpendicular bisectors of each pair of points $(\mathbf{X}, \mathbf{X}')$ with $\mathbf{X} = \mathbf{X}'$. Then $\mathcal{H}'$ is a refined partition of $\mathcal{H}$, thus $|\mathcal{H}'| \geq |\mathcal{H}|$. Now by Lemma 3 in Jiang (2019), we have $|\mathcal{H}'| \leq dN^d$.

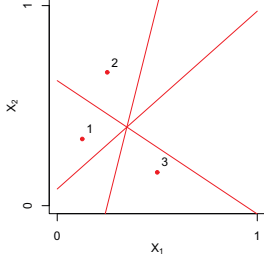


Fig. 1 The partition determining the possible sets of $A_{3,1}(\mathbf{x})$ for three points.

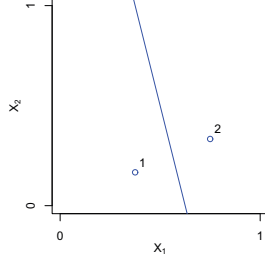


Fig. 2 The partition determining the possible sets of $A_{2,2}(\mathbf{x})$ for two points.

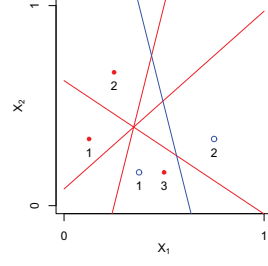


Fig. 3 The partition determining the possible sets of $A_{3,1}(\mathbf{x}) \times A_{2,2}(\mathbf{x})$.

A3 Proof of Theorem 1

In this Section, let us define set

$$\Gamma_m = \{(k_1, \dots, k_m) : k_j = \lceil k_1 n_j / n_1 \rceil, k_1 = 1, \dots, n_1, j = 1, \dots, m\}.$$

and quantity

$$\delta = C_\delta [N / \log(N)]^{-\frac{\beta}{2\beta+d}}$$

for some large enough constant C_δ . For $j = 1, \dots, m$, we denote the following random quantities:

$$k_j^{\text{opt}}(\mathbf{x}) = \max \left\{ k : \|\mathbf{X}_{(k)}^j(\mathbf{x}) - \mathbf{x}\| \leq (C_b^{-1} \delta)^{\frac{1}{\beta}} \right\}.$$

For simplicity, we may write k_j^{opt} as $k_j^{\text{opt}}(\mathbf{x})$ during the proof, if there is no confusion in the context. Define event E_A that (A3) holds for all $\mathbf{x} \in [0, 1]^d$ and all $(k_1, \dots, k_m) \in \Gamma_m$. Then by Lemma 5, we have

$$\mathbb{P}(E_A) \geq 1 - dN^{-1} \quad (\text{A2})$$

Lemma 5. *For any $\epsilon > 0$, with probability at least $1 - \tau$, the following holds:*

$$|\hat{\eta}_{k_1:k_m}(\mathbf{x}) - \mathbb{E}(\hat{\eta}_{k_1:k_m}(\mathbf{x})|\mathcal{X})| \leq \sqrt{\frac{(d+1)\log(N) - \log(\tau/d)}{2 \sum_{j=1}^m k_j}},$$

for all $\mathbf{x} \in [0, 1]^d$ and all $(k_1, \dots, k_m) \in \Gamma_m$. As a consequence, choosing $\tau = dN^{-1}$, the following holds with probability at least $1 - dN^{-1}$:

$$|\hat{\eta}_{k_1:k_m}(\mathbf{x}) - \mathbb{E}(\hat{\eta}_{k_1:k_m}(\mathbf{x})|\mathcal{X})| \leq \sqrt{\frac{(d+2)\log(N)}{2 \sum_{j=1}^m k_m}}, \quad (\text{A3})$$

for all $\mathbf{x} \in [0, 1]^d$ and all $(k_1, \dots, k_m) \in \Gamma_m$.

Proof. Notice that $k_{1:k_m}(\mathbf{x}) = \frac{1}{\prod_{j=1}^m k_j} \prod_{j=1}^m Y_{(j)}^j(\mathbf{x})$, and $Y_{(1)}^1(\mathbf{x}) \dots Y_{(k_1)}^1(\mathbf{x}) Y_{(1)}^m(\mathbf{x}) \dots Y_{(k_m)}^m(\mathbf{x})$ are independent conditional on $\mathcal{X}$. Therefore, Hoeffding's inequality implies that

$$\mathbb{P}(k_{1:k_m}(\mathbf{x}) - \mathbb{E}(k_{1:k_m}(\mathbf{x}) | \mathcal{X}) > t | \mathcal{X}) \leq \exp(-2t^2 \prod_{j=1}^m k_j)$$

Conditioning on $\mathcal{X}$, for fixed $(k_1, \dots, k_m) \in \mathbb{N}^m$, when $\mathbf{x}$ is running over $[0, 1]^d$, then by Lemma 1, there are at most dN^d different choices of $Y_{(1)}^1(\mathbf{x}) \dots Y_{(k_1)}^1(\mathbf{x}) Y_{(1)}^m(\mathbf{x}) \dots Y_{(k_m)}^m(\mathbf{x})$. Therefore, it follows that

$$\mathbb{P}(\mathbf{x} \in [0, 1]^d \text{ such that } k_{1:k_m}(\mathbf{x}) - \mathbb{E}(k_{1:k_m}(\mathbf{x}) | \mathcal{X}) > t | \mathcal{X}) \leq dN^d \exp(-2t^2 \prod_{j=1}^m k_j)$$

which further implies that

$$\mathbb{P}((k_1, \dots, k_m) \in \mathbb{N}^m \text{ such that } k_{1:k_m}(\mathbf{x}) - \mathbb{E}(k_{1:k_m}(\mathbf{x}) | \mathcal{X}) > t | \mathcal{X}) \leq dn_1 N^d \exp(-2t^2 \prod_{j=1}^m k_j) \leq d \exp(-2t^2 \prod_{j=1}^m k_j + (d+1)\log(N))$$

Plug in $t = \frac{(d+1)\log(N) - \log(d)}{2 \prod_{j=1}^m k_j}$ into above inequality and take expectation, we complete the proof. $\square$

Lemma 6. If $f(\mathbf{x}) = 1$ and $k_j = k_j^{\text{opt}}(\mathbf{x})$, then it holds that

$$\begin{aligned} \mathbb{E}(k_j(\mathbf{x}) | \mathcal{X}) &= \frac{1}{2} c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 1 \\ \mathbb{E}(k_j(\mathbf{x}) | \mathcal{X}) &= \frac{1}{2} c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 0 \end{aligned}$$

As a consequence, if $f(\mathbf{x}) = 1$ and $k_j = k_j^{\text{opt}}(\mathbf{x})$ for all $j = 1, \dots, m$, then the following holds:

$$\begin{aligned} \mathbb{E}(k_{1:k_m}(\mathbf{x}) | \mathcal{X}) &= \frac{1}{2} c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 1 \\ \mathbb{E}(k_{1:k_m}(\mathbf{x}) | \mathcal{X}) &= \frac{1}{2} c_b(\mathbf{x}) \quad \text{if } f(\mathbf{x}) = 0 \end{aligned}$$

Proof. For $k_j = k_j^{\text{opt}}$, we have

$$X_{(k)}^j(\mathbf{x}) = \mathbf{x} \cdot X_{(k_j^{\text{opt}})}^j(\mathbf{x}) = \mathbf{x} \cdot (C_b^{-1})^\perp$$

which further implies that $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = C_b \mathbf{X}_{(k)}^j(\mathbf{x}) - \mathbf{x}$. Applying Lemma 3, we complete the proof of first statement. The second statement follows from the definition that $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = \prod_{j=1}^m k_j - k_j(\mathbf{x}) \left(\prod_{j=1}^m k_j \right)$. $\square$

Lemma 7. *Under event E_A , if $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = \mathbf{x}$ and $k_j(\mathbf{x}) = k_j^{\text{opt}}(\mathbf{x})$ for all $j = 1, \dots, m$, then $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = \mathbf{x}$.*

Proof. By definition of $k_1, \dots, k_m$, we have

$$\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbf{x} \geq \frac{(d+2)\log(N)}{2 \prod_{j=1}^m k_j}$$

On event E_A , it follows that

$$\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbb{E}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) \geq \frac{(d+2)\log(N)}{2 \prod_{j=1}^m k_j}$$

Combining above, we conclude that

$$\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbf{x} \geq \mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbb{E}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) - \mathbf{x}$$

which further implies that

$$\text{sign}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbf{x}) = \text{sign}(\mathbb{E}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) | \mathcal{X}) - \mathbf{x})$$

for all $\mathbf{x}$ with $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = \mathbf{x}$ on event E_A . Finally, by Lemma 6 and above equation, on event E_A , if $\mathbf{f}_{k_1:k_m}(\mathbf{x}) = \mathbf{x}$ and $k_j(\mathbf{x}) = k_j^{\text{opt}}(\mathbf{x})$ for all $j = 1, \dots, m$, then

$$\text{sign}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbf{x}) = \begin{cases} 1 & \text{if } \mathbf{f}_{k_1:k_m}(\mathbf{x}) = 1 \\ -1 & \text{if } \mathbf{f}_{k_1:k_m}(\mathbf{x}) = 0 \end{cases}$$

which completes the proof by noticing that $\mathbf{f}_{k_1:k_m} = \mathbb{I}(\mathbf{f}_{k_1:k_m}(\mathbf{x}) - \mathbf{x} \geq 0)$. $\square$

We are ready to prove Theorem 1. Let us define the deterministic integers

$$k_j = \frac{1}{2} n_j C_D^d C_b^{\frac{d}{2}} \quad k = n_j C_D^d C_b^{\frac{d}{2}}$$

and events

$$E_j = \mathbf{X}_{(k_j)}^j(\mathbf{x}) = \mathbf{x} \quad C_D \frac{k_j}{n_j}^{\frac{1}{d}} \text{ for all } \mathbf{x} \quad \text{and } j = 1, \dots, m$$

and

$$E = \{ \mathbf{X}_{(k_j)}^j(\mathbf{x}) \leq \mathbf{x} \leq \frac{1}{C_D} \frac{k_j}{n_j}^{\frac{1}{d}} \text{ for all } \mathbf{x} \text{ and } j = 1, \dots, m \}$$

Since $\frac{1}{C_D} \frac{k_j}{n_j}^{\frac{1}{d}} > \frac{1}{2} \frac{d}{d+1}$, so $n_j^{\frac{d}{d+1}} = C^{\frac{d}{d+1}} n_j N^{-\frac{d}{2+d}} [\log(N)]^{\frac{d}{2+d}}$ $C^{\frac{d}{d+1}} N^{1 - \frac{d}{2+d}} [\log(N)]^{\frac{d}{2+d}}$ is diverging. Without loss of generality, we may assume $k_j \geq 1$ and $k_j \leq 1$. On event E , it follows from the definition of k_j^{opt} that

$$\begin{aligned} \mathbf{X}_{(k_j)}^j(\mathbf{x}) \leq \mathbf{x} &\leq C_D \frac{k_j}{n_j}^{\frac{1}{d}} = C_D \frac{n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}}}{n_j} \\ &= (C_b^{-1})^{\frac{1}{d}} < X_{(k_j^{\text{opt}}+1)}^j(\mathbf{x}) \leq \mathbf{x} \end{aligned}$$

which further implies that

$$k_j^{\text{opt}}(\mathbf{x}) \leq k_j \leq \frac{1}{2} n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}} \text{ for all } \mathbf{x} \text{ and } j = 1, \dots, m \quad (\text{A4})$$

Moreover, on event E , it also holds that

$$\begin{aligned} \mathbf{X}_{(k_j)}^j \leq \mathbf{x} &\leq \frac{1}{C_D} \frac{k_j}{n_j}^{\frac{1}{d}} = \frac{1}{C_D} \frac{n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}}}{n_j} \\ &= (C_b^{-1})^{\frac{1}{d}} = X_{(k_j^{\text{opt}})}^j \leq \mathbf{x} \end{aligned}$$

where the last equation follows from the definition of k_j^{opt} . Above inequality implies that

$$k_j^{\text{opt}}(\mathbf{x}) \leq k_j \leq n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}} \quad (\text{A5})$$

for all $\mathbf{x}$ and $j = 1, \dots, m$. Combining (A4) and (A5), we conclude that on event $E = E_1 \cap \dots \cap E_m$, the following holds:

$$\frac{1}{2} n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}} \leq k_j^{\text{opt}}(\mathbf{x}) \leq n_j C_D^d C_b^{\frac{d}{d+1}}^{\frac{1}{d}} \quad (\text{A6})$$

for all $\mathbf{x}$ and $j = 1, \dots, m$. If we define $k_1^{\min} = \min_{j=1, \dots, m} k_j^{\text{opt}} n_j$ and $k_j^{\min} = k_1^{\min} n_j / n_1$, then we can show that

$$k_j^{\min} = \frac{k_1^{\text{opt}} n_j}{n_1} = \frac{k_j^{\text{opt}} n_1 n_j}{n_j n_1} = k_j^{\text{opt}} = k_j^{\text{opt}} \quad (\text{A7})$$

Hence, by (A6), it holds on event $E = E$ that

$$\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}} \leq k_j^{\min}(\mathbf{x}) \leq 2n_j C_D^d C_b^{\frac{d}{2}} \quad (\text{A8})$$

for all $\mathbf{x}$ and $j = 1, \dots, m$. By definition, it follows that $k_j^{\min}(\mathbf{x}) \leq k_j^{\text{opt}}$. So under event $E = E$, for all $\mathbf{x}$ with $f(\mathbf{x}) = 1$, Lemma 6 and (A8) together imply that

$$\begin{aligned} & \overline{\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}}}_{j=1}^m \leq \mathbb{E}(\overline{k_1^{\min}, k_1^{\min}}(\mathbf{x}) | \mathcal{X}) \leq \frac{1}{2} \\ & \overline{\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}}}_{j=1}^m \leq \overline{c_b^{\frac{d}{2}}}(\mathbf{x}) \\ & \overline{\frac{1}{4}c_b^2 C_D^d C_b^{\frac{d}{2}}} \leq N^{\frac{d}{2}+2} \\ & = \overline{\frac{1}{4}c_b^2 C_D^d C_b^{\frac{d}{2}} C^{\frac{2+d}{2}}} \leq \overline{N \log(N)}^1 \\ & \leq 3 \frac{(d+2)\log(N)}{2} \end{aligned}$$

where the last inequality follows if we choose C large. By above inequality, on event $E_A = E = E$, for all $\mathbf{x}$ with $f(\mathbf{x}) = 1$, it follows that

$$\begin{aligned} & \overline{\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}}}_{j=1}^m \leq \mathbb{E}(\overline{k_1^{\min}, k_1^{\min}}(\mathbf{x}) | \mathcal{X}) \leq \frac{1}{2} \\ & \overline{\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}}}_{j=1}^m \leq \mathbb{E}(\overline{k_1^{\min}, k_m^{\min}}(\mathbf{x}) | \mathcal{X}) \leq \frac{1}{2} \\ & \overline{\frac{1}{4}n_j C_D^d C_b^{\frac{d}{2}}}_{j=1}^m \leq \mathbb{E}(\overline{k_1^{\min}, k_m^{\min}}(\mathbf{x}) | \mathcal{X}) \\ & \leq 3 \frac{(d+2)\log(N)}{2} \leq \frac{(d+2)\log(N)}{2} \\ & > \frac{(d+2)\log(N)}{2} \end{aligned}$$

Similarly, on event $E_A \cap E \cap E^c$, for all $\mathbf{x}$ with $(\mathbf{x}) \leq \frac{1}{2}$ and $f(\mathbf{x}) = 0$, we can show that

$$\overline{\prod_{j=1}^m k_j^{\min}} \mathbb{E}(\prod_{j=1}^m k_j^{\min} | \mathcal{X}) \leq \frac{1}{2} < \frac{(d+2)\log(N)}{2}$$

Therefore, we prove that on event $E_A \cap E \cap E^c$, the following holds:

$$\overline{\prod_{j=1}^m k_j^{\min}} \prod_{j=1}^m k_j^{\min} > \frac{(d+2)\log(N)}{2} \quad (\text{A9})$$

for all $\mathbf{x}$ with $(\mathbf{x}) \leq \frac{1}{2}$. Now by (A8), on event $E_A \cap E \cap E^c$, we have

$$\begin{aligned} k_1^{\min}(\mathbf{x}) &= n_1 C_D^d C_b^{\frac{d}{2+d}} \\ &= C_D^d C_b^{\frac{d}{2+d}} C^{\frac{d}{2+d}} n_1 N^{\frac{d}{2+d}} \log(N) \\ &= C_D^d C_b^{\frac{d}{2+d}} C^{\frac{d}{2+d}} n_1 N^{\frac{d}{2+d}} [\log(N)]^{\frac{d}{2+d}} \\ &= n_1 N^{\frac{d}{2+d}} \log(N) \end{aligned} \quad (\text{A10})$$

for all $\mathbf{x}$ with $(\mathbf{x}) \leq \frac{1}{2}$ and sufficiently large N . In the following, we will calculate the probability of $E_A \cap E \cap E^c$. Since $d > 1$, it follows that

$$\frac{(1+\frac{1}{d})}{2+\frac{1}{d}} = \frac{1}{2+\frac{1}{d}} < \frac{1}{2} < 1$$

Using the inequality above and (A2), it follows that

$$\mathbb{P}(E_A) \leq 1 - dN^{-1} \leq 1 - dC^{1+\frac{1}{d}} N^{-\frac{(1+\frac{1}{d})}{2+\frac{1}{d}}} [\log(N)]^{\frac{(1+\frac{1}{d})}{2+\frac{1}{d}}} = 1 - d^{-1+\frac{1}{d}}$$

for sufficiently large N . Moreover, by Lemma 2, (A4) and (A5), it follows that

$$\begin{aligned} \mathbb{P}(E^c) &= 1 - C_D \prod_{j=1}^m \frac{n_j}{k_j} \exp(-k_j) \\ &= 1 - C_D N \exp\left(-\frac{n_j}{12C_D^d C_b^{\frac{d}{2+d}}}\right) \\ &= 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2+d}}} N^{\frac{d}{2+d}} \log(N)\right) \\ &= 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2+d}}} N^{\frac{d}{2+d}} [\log(N)]^{\frac{d}{2+d}}\right) \end{aligned}$$

$$1 - \frac{1}{1+d}$$

where we use the fact that $1 - \frac{1}{1+d} > d/(2+d)$. Similarly, we can show that $\mathbb{P}(E) \geq 1 - \frac{1}{1+d}$. Combining above, we show that

$$\mathbb{P}(E_A \cap E \cap E^c) \geq 1 - \mathbb{P}(E_A) - \mathbb{P}(E) - \mathbb{P}(E^c) \geq 1 - \frac{1}{1+d} \quad (\text{A11})$$

By (A7), (A9) and the definition of $k_1(\mathbf{x}) \dots k_m(\mathbf{x})$, the following holds on event $E_A \cap E \cap E^c$:

$$k_j(\mathbf{x}) = k_j^{\min}(\mathbf{x}) = k_j^{\text{opt}}(\mathbf{x}) \quad (\text{A12})$$

for all $\mathbf{x}$ with $(\mathbf{x}) \in \mathcal{X}$ and $j = 1 \dots m$. By (A10), (A12) and Lemma 7, we can see, it holds on event $E_A \cap E \cap E^c$ that:

$$f_{k_1:k_m}(\mathbf{x}) = f(\mathbf{x}) \quad \text{and} \quad k_1(\mathbf{x}) \leq n_1 N^{-\frac{d}{2+d}} \log(N) \quad (\text{A13})$$

for all $\mathbf{x}$ with $(\mathbf{x}) \in \mathcal{X}$. By (A13), on event $E_A \cap E \cap E^c$, $f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})$ implies $(x) < \epsilon$. As a consequence of (A11) and Assumption A3, we have

$$\begin{aligned} & \mathcal{R}(f_{k_1:k_m}) \\ &= \mathbb{E}(\mathbb{I}(\mathbf{X}) \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}))) \\ &= \mathbb{E}(\mathbb{I}(\mathbf{X}) \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) \cap E_A \cap E \cap E^c)) + \mathbb{P}((E_A \cap E \cap E^c)^c) \\ &= \mathbb{E}(\mathbb{I}(\mathbf{X}) \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) \cap E_A \cap E \cap E^c \cap (\mathbf{X}) < \epsilon)) + \mathbb{P}((E_A \cap E \cap E^c)^c) \\ &= \mathbb{P}((\mathbf{X}) < \epsilon) + \frac{1}{1+d} \end{aligned} \quad (\text{A14})$$

A4 Proof of Theorem 2

By the conditions given, we have $k_1 \dots n_1 N^{\frac{d}{2+d}} = N^{1 - \frac{d}{2+d}} = N^{\frac{2}{2+d}}$ and $k_1 n_j \dots n_1 = n_j N^{\frac{d}{2+d}} = N^{1 - \frac{d}{2+d}} = N^{\frac{2}{2+d}}$. Without loss of generality, we assume $k_1 n_j \dots n_1 = 1$ for all $j = 1 \dots m$. Let C large enough constant such that $\dots = C (k_1 \dots n_1)^{\frac{d}{2+d}} = C_b (k_1 \dots n_1)^{\frac{d}{2+d}}$. We further define sets $A_0 = \{\mathbf{x} : \|\mathbf{x}\| \leq 1\}$ and $A_j = \{\mathbf{x} : 2^{j-1} < \|\mathbf{x}\| \leq 2^j\}$ for $j = 1$. For simplicity, let us write $E_P(k_1 : k_m, 0)$ as E_P and $\prod_{j=1}^m k_j \dots m$ as k .

If $j = 0$, then Assumption A3 shows that

$$\mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_0 \setminus E_P) \leq 2 \mathbb{P}(\mathbf{X} \in A_0) \leq 2C^{-1}$$

If $j = 1$, Assumption A3 and Lemma 4 imply that

$$\begin{aligned} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \\ \leq 2^{j+1} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \mathbb{P}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) | \mathbf{X} \in \mathcal{X}) \\ \leq 2^{j+1} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \exp(-2c_b^2 m k^{-2}(\mathbf{x})) \\ \leq 2^{j+1} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \exp(-2c_b^2 m k 4^{j-1-2}) \end{aligned}$$

By conditions given, it follows that

$$mk = \prod_{j=1}^m k_j \prod_{j=1}^m \frac{k_1 n_j}{n_1} \prod_{j=1}^m n_j N^{\frac{d}{2+d}} = N^{\frac{2}{2+d}}$$

and $\dots = C (k_1 \dots n_1)^{\frac{d}{2+d}} = C N^{\frac{d}{2+d}}$. Therefore, it follows that

$$\begin{aligned} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \\ \leq 2^{j+1} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \exp\left(-\frac{1}{2} c_b^2 C^2 4^j N^{\frac{2}{2+d}} N^{-\frac{2}{2+d}}\right) \\ \leq 2^{j+1} \mathbb{E} \|\mathbf{X}\|^{-1} \mathbb{I}(\mathbf{X} \in A_j \setminus E_P) \exp\left(-\frac{1}{2} c_b^2 C^2 4^j\right) \\ \leq 2^{j+1} \mathbb{P}(\mathbf{X} \in A_j) \exp\left(-\frac{1}{2} c_b^2 C^2 4^j\right) \\ \leq 2^{j+1} C (2^j)^{-\frac{d}{2+d}} \exp\left(-\frac{1}{2} c_b^2 C^2 4^j\right) \end{aligned}$$

$$2^{j+1} C (2^j)^{-1} \exp \left(-\frac{1}{2} c_b^2 C^2 4^j \right)$$

$$2C^{-1+} 2^{j(1+)} \exp \left(-\frac{1}{2} c_b^2 C^2 4^j \right)$$

Taking summation and using the fact that $\sum_{j=1}^{\infty} b^j \exp(-c4^j) < \infty$ for all $b, c > 0$, we have

$$\mathbb{E} \sum_{j=1}^{\infty} 2^j (\mathbf{X})^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_j \cap E_P^c)$$

$$2C^{-1+} \sum_{j=1}^{\infty} 2^{j(1+)} \exp \left(-\frac{1}{2} c_b^2 C^2 4^j \right) = 1 +$$

Combining the bounds above, it follows that

$$\begin{aligned} \mathcal{R}(f_{k_1:k_m}) &= \mathbb{E} \sum_{j=1}^{\infty} 2^j (\mathbf{X})^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X})) \\ &\leq \mathbb{E} \sum_{j=1}^{\infty} 2^j (\mathbf{X})^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) \mid E_P^c) \\ &\quad + \mathbb{E} \sum_{j=1}^{\infty} 2^j (\mathbf{X})^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) \mid \mathbf{X} \in A_0 \cap E_P) \\ &\quad + \sum_{j=1}^{\infty} \mathbb{E} \sum_{j=1}^{\infty} 2^j (\mathbf{X})^{-1} \mathbb{I}(f_{k_1:k_m}(\mathbf{X}) = f(\mathbf{X}) \mid \mathbf{X} \in A_j \cap E_P) \\ &\leq \mathbb{P}(E_P^c) + 1 + \\ &\leq \mathbb{P}(E_P^c) + N^{-\frac{(1+)}{2+d}} \end{aligned}$$

By (A1) and the fact that $k_j \leq k_1 n_j \leq n_1 \leq n_j N^{\frac{d}{2+d}} \leq N^{1-\frac{d}{2+d}} = N^{\frac{2}{2+d}}$, it yields that

$$\mathbb{P}(E_P^c) \leq C_D \sum_{j=1}^m \frac{n_j}{k_j} \exp(-k_j/6) \leq N^{-\frac{(1+)}{d}}$$

Combining the above two inequalities, we complete the proof.

A5 Proof of Theorem 3

In this section, we consider the equal-size sub-samples. For simplicity, let us assume $n_1 = \dots = n_m = n = N^{1-}$ and $m = N - n = N^+$. Notice that $k_1 = \dots = k_m$ in this setting, so we rewrite $f_{k_1:k_m}$ as f_k and $f_{k_1:k_m}$ as f_k if $k_1 = \dots = k_m = k$. Moreover, we also define $k = k_1 = \dots = k_m$ to be the data-driven quantities in Algorithm 1. Since Theorem 1 studies the case with $\frac{2}{2+d} < \frac{2}{2+d}$, we will focus on the case with $\frac{2}{2+d}$. Abusing notation, let us define quantities

$$v := v(a) = \frac{(1-a)(1-a)}{d} \quad \text{if} \quad \frac{2}{2+d} < a < 1$$

and

$$= C [N \log(N)]^{-v}$$

for some large enough constant C . For $j = 1 \dots m$, we denote the following random quantities:

$$k_j^{\text{opt}}(\mathbf{x}) = \max_{k: \mathbf{X}_{(k)}^j(\mathbf{x}) \leq \mathbf{x}} (C_b^{-1})^{\frac{1}{d}}$$

For simplicity, we may write v as $v(a)$ and k_j^{opt} as $k_j^{\text{opt}}(\mathbf{x})$ during the proof, if there is no confusion in the context. Clearly, Lemmas 6 and 7 are still valid under the new v and $k_j^{\text{opt}}(\mathbf{x})$.

Lemma 8. Suppose $\frac{2}{2+d} < \frac{2}{2+d}$, then there exists a constant $c > 0$ such that the following holds

$$\mathcal{R}(f_k) \leq N \log(N)^{\frac{(1-a)}{d}} [\log(N)]$$

for some $\epsilon > 0$ depending on d .

Proof. For fixed (a) , we define the deterministic integers

$$k = \frac{1}{2} n^{1-a} C_D^d C_b^{\frac{d}{d-1}} \quad k = n C_D^d C_b^{\frac{d}{d-1}}$$

and events

$$E(a) = \mathbf{X}_{(k)}^j(\mathbf{x}) \leq \mathbf{x} \quad C_D \frac{k}{n^{1-a}}^{\frac{1}{d}} \text{ for all } \mathbf{x} \quad \text{and } j = 1 \dots m$$

and

$$E = \mathbf{X}_{(k)}^j(\mathbf{x}) \leq \mathbf{x} \quad \frac{1}{C_D} \frac{k}{n}^{\frac{1}{d}} \text{ for all } \mathbf{x} \quad \text{and } j = 1 \dots m$$

By the definition of v , for any (a) , we can verify that $(1-a)(1-a) < v/d$ and $n^{1-a} = C^{\frac{d}{d-1}} n^{1-a} N^{\frac{v-d}{d-1}} [\log(N)]^{\frac{v-d}{d-1}} = C^{\frac{d}{d-1}} N^{(1-a)(1-a) \frac{v-d}{d-1}} [\log(N)]^{\frac{v-d}{d-1}}$ is diverging. Consequently, we may assume $k \rightarrow 1$ and $k \rightarrow 1$. On event $E(a)$, it follows from the definition of k_j^{opt} that

$$\begin{aligned} \mathbf{X}_{(k)}^j(\mathbf{x}) \leq \mathbf{x} &\quad C_D \frac{k}{n^{1-a}}^{\frac{1}{d}} < C_D \frac{n^{1-a} C_D^d C_b^{\frac{d}{d-1}}}{n^{1-a}}^{\frac{1}{d}} \\ &\quad (C_b^{-1})^{\frac{1}{d}} < X_{(k_j^{\text{opt}}+1)}^j(\mathbf{x}) \leq \mathbf{x} \end{aligned}$$

which further implies that

$$k_j^{\text{opt}}(\mathbf{x}) \leq k + \frac{1}{2}n^{1-a}C_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \quad (\text{A15})$$

for all $\mathbf{x}$ and $j = 1, \dots, m$. Moreover, on event E , it also holds that

$$\mathbf{X}_{(k_j)}^j(\mathbf{x}) \leq \frac{1}{C_D} \frac{k}{n} \frac{1}{n^{\frac{1}{d}}} (C_b^{-1})^{\frac{1}{d}} X_{(k_j^{\text{opt}})}^j(\mathbf{x})$$

where the last equation follows from the definition of k_j^{opt} . Above inequality implies that

$$k_j^{\text{opt}}(\mathbf{x}) \leq k + nC_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \quad (\text{A16})$$

for all $\mathbf{x}$ and $j = 1, \dots, m$. Combining (A15) and (A16), we conclude that on event E (a) E , the following holds:

$$\frac{1}{2}n^{1-a}C_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \leq k_j^{\text{opt}}(\mathbf{x}) \leq nC_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \quad \text{for all } \mathbf{x} \text{ and } j = 1, \dots, m \quad (\text{A17})$$

Define $k^{\min} = \min_{1 \leq j \leq m} k_j^{\text{opt}} \leq k_m^{\text{opt}}$. Hence, by (A17), it holds on event E (a) E that

$$\frac{1}{2}n^{1-a}C_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \leq k^{\min}(\mathbf{x}) \leq nC_D^dC_b^{\frac{d}{2}-\frac{d}{2}} \quad \text{for all } \mathbf{x} \quad (\text{A18})$$

Since $k^{\min}(\mathbf{x}) \leq k_j^{\text{opt}}$, so under event E (a) E , for all $\mathbf{x}$ with $\mathbf{f}(\mathbf{x}) = 1$, Lemma 6 and (A18) together imply that

$$\begin{aligned} & \overline{mk^{\min}} \leq \mathbb{E}(\overline{k^{\min}(\mathbf{x})}) \leq \frac{1}{2} \\ & \overline{\frac{1}{2}mn^{1-a}C_D^dC_b^{\frac{d}{2}-\frac{d}{2}}c_b(\mathbf{x})} \\ & \overline{\frac{1}{2}c_b^2C_D^dC_b^{\frac{d}{2}-\frac{d}{2}}} \leq \overline{m^aN^{1-a}N^{\frac{d}{2}+2}} \\ & = \overline{\frac{1}{2}c_b^2C_D^dC_b^{\frac{d}{2}-\frac{d}{2}}C^{\frac{2+d}{2}}} \leq \overline{N^aN^{1-a}N \log(N)^{\frac{v(2+d)}{2}}} \end{aligned}$$

Since by definition, $v = (1-a)(1-a-d)$ if $2 \leq (2+d)$ $0 < a < 1$, it holds that

$$\frac{v(2+d)}{2} = \frac{(1-a)(1-a)(2+d)}{d} \leq 1-a$$

Combining the above two inequalities, we have

$$\begin{aligned}
& \overline{mk^{\min}} \mathbb{E}(k^{\min}(\mathbf{x}) | \mathcal{X}) \leq \frac{1}{2} \\
& \frac{\frac{1}{2}c_b^2 C_D^d C_b^{\frac{d}{2+d}} C^{\frac{2+d}{2}}}{N^a N^{1-a} \log(N)} \leq \frac{1}{2} \quad \text{if } a < \frac{2}{2+d} \quad a = 1 \\
& \frac{\frac{1}{2}c_b^2 C_D^d C_b^{\frac{d}{2+d}} C^{\frac{2+d}{2}}}{N^a N^{1-a} \log(N)} \leq \frac{1}{2} \quad \text{if } a > \frac{2}{2+d} \quad a > 0 \\
& 3 \frac{(d+2)\log(N)}{2}
\end{aligned}$$

where the last inequality follows if we choose C large. By above inequality, on event $E_A = E(a) = E$, for all $\mathbf{x}$ with $k^{\min}(\mathbf{x}) > \frac{1}{2}$ and $f(\mathbf{x}) = 1$, it follows that

$$\begin{aligned}
& \overline{mk^{\min}} k^{\min}(\mathbf{x}) \leq \frac{1}{2} \\
& \overline{mk^{\min}} \mathbb{E}(k^{\min}(\mathbf{x}) | \mathcal{X}) \leq \frac{1}{2} \quad \overline{mk^{\min}} k^{\min}(\mathbf{x}) \leq \mathbb{E}(k^{\min}(\mathbf{x}) | \mathcal{X}) \\
& 3 \frac{(d+2)\log(N)}{2} \leq \frac{(d+2)\log(N)}{2} \\
& > \frac{(d+2)\log(N)}{2}
\end{aligned}$$

Similarly, on event $E_A = E(a) = E$, for all $\mathbf{x}$ with $k^{\min}(\mathbf{x}) > \frac{1}{2}$ and $f(\mathbf{x}) = 0$, we can show that

$$\overline{mk^{\min}} k^{\min}(\mathbf{x}) \leq \frac{1}{2} < \frac{(d+2)\log(N)}{2}$$

Therefore, we prove that on event $E_A = E(a) = E$, the following holds:

$$\overline{mk^{\min}} k^{\min}(\mathbf{x}) \leq \frac{1}{2} > \frac{(d+2)\log(N)}{2} \text{ for all } \mathbf{x} \text{ with } k^{\min}(\mathbf{x}) > \frac{1}{2} \quad (\text{A19})$$

Now by (A18) and the definition of v , on event $E_A = E(a) = E$, we have

$$\begin{aligned}
k^{\min}(\mathbf{x}) &= n C_D^d C_b^{\frac{d}{2+d}} C^{\frac{2+d}{2}} \\
&= C_D^d C_b^{\frac{d}{2+d}} C^{\frac{2+d}{2}} n^{\frac{v-d}{2}} N^{\log(N)} \\
&= C_D^d C_b^{\frac{d}{2+d}} C^{\frac{2+d}{2}} n N^{(1-\frac{v}{2})(1-a)\log(N)} \quad \text{if } \frac{2}{2+d} < a < 1
\end{aligned}$$

(A20)

for all $\mathbf{x}$ with $\|\mathbf{x}\| \leq \frac{1}{2}$. In the following, we will calculate the probability of $E_A \cap E(a) \cap E$. Since $\frac{1}{2} \leq d$ by Assumption A3, it follows that

$$\begin{aligned} v(1 + \frac{1}{2}) &= \frac{(1 - \frac{1}{2})(1 - a)(1 + \frac{1}{2})}{d} - \frac{(1 - a)(1 + \frac{1}{2})}{2 + d} - \frac{(1 + \frac{1}{2})}{2 + d} \\ &= \frac{1}{2 + d} - \frac{1 + d}{2 + d} < 1 \end{aligned}$$

Using the inequality above and (A2), it follows that

$$\mathbb{P}(E_A \cap E(a) \cap E) \leq dN^{-1} \leq dC^{1+} N^{-v(1+\frac{1}{2})} [\log(N)]^{v(1+\frac{1}{2})} = 1 - d^{-1+}$$

for sufficiently large N . Moreover, by Lemma 2 and the definition of $E(a) \cap E$, it follows that

$$\begin{aligned} \mathbb{P}(E(a)) &\leq 1 - C_D \frac{mn^{1-a}}{k} \exp(-n^a k^{-6}) \\ &\leq 1 - C_D N \exp\left(-\frac{n^{\frac{d}{2}}}{12C_D^d C_b^{\frac{d}{2}}}\right) \\ &= 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2}}} N^{1-\frac{v}{2}} \log(N)^{\frac{v}{2}}\right) \\ &= 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2}}} N^{(1-a)a} [\log(N)]^{(1-a)(1-a)}\right) \end{aligned}$$

and

$$\begin{aligned} \mathbb{P}(E) &\leq 1 - C_D \frac{mn}{k} \exp(-k^{-6}) \\ &\leq 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2}}} n^{\frac{d}{2}}\right) \\ &\leq 1 - C_D N \exp\left(-\frac{1}{12C_D^d C_b^{\frac{d}{2}}} N^{(1-a)a} [\log(N)]^{(1-a)(1-a)}\right) \end{aligned}$$

Combining the above three inequalities, we show that

$$\begin{aligned} \mathbb{P}(E_A \cap E(a) \cap E) &\leq 1 - \mathbb{P}(E_A) - \mathbb{P}(E(a)) - \mathbb{P}(E) \\ &\leq 1 - d^{-1+} - CN \exp\left(-\frac{1}{C} N^{(1-a)a} [\log(N)]^{(1-a)(1-a)}\right) \\ &\leq 1 - d^{-1+} - CN \exp\left(-\frac{1}{C} N^{(1-a)a}\right) \end{aligned}$$

where C is some constant greater than $C_D + 12C_D^d C_b^{\frac{d}{2+d}} + 12C_D^d C_b^{\frac{d}{2}}$. Without loss of generality, we assume $C > 6$. if we choose $a = \frac{\log(2C \log(N))}{(1-\frac{1}{2+d})\log(N)}$, then we have

$$\begin{aligned} N^a &= [2C \log(N)]^{\frac{1}{1-\frac{1}{2+d}}} \\ &= C^{-\frac{(1-\frac{1}{2+d})(1-a)}{d}} N^{\log(N)} C^{-\frac{(1-\frac{1}{2+d})}{d}} [2C \log(N)]^{\frac{1}{d}} \end{aligned} \quad (\text{A21})$$

and the probability can be bounded by

$$\mathbb{P}(E_A \cap E(a) \cap E) \leq 1 - C N^{-1} \geq 1 - (C+1)^{-1} \quad (\text{A22})$$

First, let us consider the case $\frac{2}{2+d} < \frac{2}{2+d}$. By (A19), (A20) and the definition of $k(\mathbf{x})$, since $\frac{2}{2+d} < \frac{2}{2+d}$ and $0 < a = \frac{\log(2C \log(N))}{(1-\frac{1}{2+d})\log(N)} \leq 1 - \frac{d}{(2+d)(1-\frac{1}{2+d})}$ for large N , we have

$$k(\mathbf{x}) \leq k^{\min}(\mathbf{x}) \leq n N^{-\frac{d}{2+d} \log(N)} \quad \text{and} \quad k(\mathbf{x}) \leq k^{\min}(\mathbf{x}) \leq k_j^{\text{opt}}(\mathbf{x})$$

holds for all $j = 1, \dots, m$ and all $\mathbf{x}$ with $(\mathbf{x})$ on event $E_A \cap E(a) \cap E$. Now, applying Lemma 7, we can see, it holds on event $E_A \cap E(a) \cap E$ that:

$$f_k(\mathbf{x}) = f(\mathbf{x}) \quad \text{for all } \mathbf{x} \text{ with } (\mathbf{x})$$

By (A22), (A21) and the above equation, we can complete the proof using the same argument as (A14).

Second, let us assume $\frac{2}{2+d} = \frac{2}{2+d}$. By (A20) and the definition of $k(\mathbf{x})$, the following holds on event $E_A \cap E(a) \cap E$:

$$k(\mathbf{x}) \leq k^{\min}(\mathbf{x}) \leq C_D^d C_b^{\frac{d}{2+d}} C^{\frac{d}{2+d}} n N^{-(1-\frac{1}{2+d})(1-a) \log(N)} = C_D^d C_b^{\frac{d}{2+d}} C^{\frac{d}{2+d}} N^{(1-\frac{1}{2+d})a} \log(N) := v_N$$

for all $\mathbf{x}$ with $(\mathbf{x})$, which further leads to

$$C_b N^{\frac{(1-\frac{1}{2+d})(1-a)}{d}} \leq C_b [k(\mathbf{x}) N^{(1-\frac{1}{2+d})(1-a)}]^{\frac{1}{d}} \leq C_b v_N^{\frac{1}{d}} N^{\frac{(1-\frac{1}{2+d})(1-a)}{d}} \quad (\text{A23})$$

Let us define classifier

$$f(\mathbf{x}) = \begin{cases} f_k(\mathbf{x}) & \text{if } f_k(\mathbf{x}) = f(\mathbf{x}) \text{ for all } 1 \leq k \leq v_N; \\ f(\mathbf{x}) & \text{elsewhere} \end{cases}$$

By the definition of $f(\mathbf{x})$, it follows that

$$\mathbb{P}(f(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) = \mathbb{P}(1 \leq k \leq v_N \text{ such that } f_k(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X})$$

$$\mathbb{P}(f_k(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) = \frac{v_N}{k=1}$$

By the above inequality, (A23) and Lemma 4, we conclude the following hold on event $\bigcap_{k=1}^{v_N} E_P(k : k-a) \cap E_A \cap E(a) \cap E$:

$$\mathbb{P}(f_k(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) = \mathbb{P}(f_k(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) \quad \mathbb{P}(f(\mathbf{x}) = f(\mathbf{x}) | \mathcal{X}) \\ v_N \exp \left(-2c_b^2 m^{-2}(\mathbf{x}) \right) \quad (\text{A24})$$

for all $\mathbf{x}$ with $\|\mathbf{x}\| \leq \max_{k=1}^{v_N} C_b v_N^{\frac{1}{d}} N^{\frac{(1-a)}{d}}$. For simplicity, let us denote $\beta = \max_{k=1}^{v_N} C_b v_N^{\frac{1}{d}} N^{\frac{(1-a)}{d}}$, $E = \bigcap_{k=1}^{v_N} E_P(k : k-a) \cap E_A \cap E(a) \cap E$, $A_0 = \{\mathbf{x} : \|\mathbf{x}\| \leq \beta\}$, and $A_j = \{\mathbf{x} : 2^{j-1} \beta < \|\mathbf{x}\| \leq 2^j \beta\}$. If $j = 0$, then Assumption A3 shows that

$$\mathbb{E} \left[\sum_{\mathbf{x} \in A_0} \mathbb{I}(f_k(\mathbf{x}) = f(\mathbf{x})) \mathbb{I}(\mathbf{x} \in A_0 \cap E) \right] \leq 2^j P(\mathbf{x} \in A_0) \leq 2C^{-1+}$$

If $j \geq 1$, (A24) implies that

$$\mathbb{E} \left[\sum_{\mathbf{x} \in A_j} \mathbb{I}(f_k(\mathbf{x}) = f(\mathbf{x})) \mathbb{I}(\mathbf{x} \in A_j \cap E) \right] \\ \leq 2^{j+1} \mathbb{E} \left[\mathbb{I}(\mathbf{x} \in A_j \cap E) \mathbb{P}(f_k(\mathbf{x}) = f(\mathbf{x}) | \mathbf{x}, \mathcal{X}) \right] \\ \leq 2^{j+1} v_N \mathbb{E} \left[\mathbb{I}(\mathbf{x} \in A_j \cap E) \exp \left(-2c_b^2 m^{-2}(\mathbf{x}) \right) \right] \\ \leq 2^{j+1} v_N \mathbb{E} \left[\mathbb{I}(\mathbf{x} \in A_j \cap E) \exp \left(-2c_b^2 m 4^{j-1} \right) \right]$$

Since $\beta = \max_{k=1}^{v_N} C_b v_N^{\frac{1}{d}} N^{\frac{(1-a)}{d}} = C_D^d C_b^{\frac{d+}{d}} C^{\frac{d}{d}} N^{\frac{(d+)a}{2+d}} \log(N)$ and $m = N^{-\frac{2}{2+d}}$, so $m^{-2} \geq 1$ for all $a > 0$. As a consequence of Assumption A3, it follows that

$$\mathbb{E} \left[\sum_{j=1}^{\infty} \sum_{\mathbf{x} \in A_j} \mathbb{I}(f_k(\mathbf{x}) = f(\mathbf{x})) \mathbb{I}(\mathbf{x} \in A_j \cap E) \right] \\ \leq v_N \sum_{j=1}^{\infty} 2^{j+1} \exp \left(-2c_b^2 4^{j-1} \right) \mathbb{P}(A_j) \\ \leq 2^{1+} v_N \sum_{j=1}^{\infty} 2^{j(1+)} \exp \left(-\frac{1}{2} c_b^2 4^j \right) \leq C^{-1+} v_N$$

where we use the fact that $\sum_{j=1}^{\infty} b^j \exp(-c4^j) < \infty$ for all $b, c > 0$, and $C > 0$ is a constant free of N . Combining the bounds above with (A1) and (A22), it follows that

$$\begin{aligned} \mathcal{R}(f_k) &= \mathbb{E} \sum_{\mathbf{X}} \frac{1}{N} \mathbb{I}(f_k(\mathbf{X}) = f(\mathbf{X})) \\ &\leq \mathbb{E} \sum_{\mathbf{X}} \frac{1}{N} \mathbb{I}(f_k(\mathbf{X}) = f(\mathbf{X}) | E^c) \\ &\quad + \mathbb{E} \sum_{\mathbf{X}} \frac{1}{N} \mathbb{I}(f_k(\mathbf{X}) = f(\mathbf{X}) | \mathbf{X} \in A_0 \cup E_P) \\ &\quad + \sum_{j=1}^{\infty} \mathbb{E} \sum_{\mathbf{X}} \frac{1}{N} \mathbb{I}(f_k(\mathbf{X}) = f(\mathbf{X}) | \mathbf{X} \in A_j \cup E) \\ &\leq \mathbb{P}(E^c) + C^{1+\frac{1}{d}} v_N \\ &\quad + \sum_{k=1}^{\infty} \mathbb{P}(E_P^c(k, a)) + \mathbb{P}((E_A^c \cup E \cup A_0 \cup E)^c) + C^{1+\frac{1}{d}} v_N \\ &\leq C_D v_N N \exp(-n^a/6) + (1+C)^{1+\frac{1}{d}} + C^{1+\frac{1}{d}} v_N \end{aligned}$$

Notice that $a = \frac{\log(2C \log(N))}{(1-\frac{1}{d}) \log(N)}$ with $C > 6$, we have

$$\begin{aligned} v_N &= C_D^d C_b^{\frac{d}{d-1}} C^{\frac{d}{d-1}} N^{(1-\frac{1}{d})a} \log(N) = 2C C_D^d C_b^{\frac{d}{d-1}} C^{\frac{d}{d-1}} [\log(N)]^2 \\ &= C_b v_N^{\frac{d}{d-1}} N^{\frac{(1-\frac{1}{d})(1-a)}{d}} N^{\frac{(1-\frac{1}{d})}{d}} [\log(N)]^{\frac{3}{d}} \\ \exp(-n^a/6) &= \exp(-[2C \log(N)]/6) \leq N^{-2} \end{aligned}$$

Since $\frac{1}{d} = \frac{2}{2+d}$, we conclude that

$$\mathcal{R}(f_k) \leq N \log(N)^{\frac{(1-\frac{1}{d})(1+\frac{1}{d})}{d}} [\log(N)]^{\frac{3}{d}} \text{ for some } \gamma > 0$$

□

Lemma 9. Under Assumptions A1-A3, if $\frac{1}{d} \leq (2 + d)$ and $k = 1$, then

$$\mathcal{R}(f_k) \leq N^{\frac{(1-\frac{1}{d})(1+\frac{1}{d})}{d}} [\log(N)]$$

for some $\gamma > 0$ depending on d and $\frac{1}{d}$.

Proof. Let us define $\tilde{v}_N = C_b [k N^{(1-\frac{1}{d})(1-a)}]^\frac{d}{d-1}$, where C_b is the constant in Lemma 3, and we allow a is depending on N . We further define sets $A_0 = \{\mathbf{x} : (\mathbf{x})_1 \geq 2\}$ and $A_j = \{\mathbf{x} : 2^{j-1} < (\mathbf{x})_1 \leq 2^j\}$ for $j \geq 1$. For simplicity, we write $E_P(k_1 : k_m, a) = E_P$ and in the equal sub-sample size setting, we have $k_1 = \dots = k_m = k = 1$. If $j = 0$, then Assumption A3 shows that

$$\mathbb{E} \sum_{\mathbf{X}} \frac{1}{N} \mathbb{I}(f_k(\mathbf{X}) = f(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_0 \cup E_P) \leq 2 \mathbb{P}(\mathbf{X} \in A_0) \leq 2C^{1+\frac{1}{d}}$$

If $j \geq 1$, Assumption A3 and Lemma 4 imply that

$$\begin{aligned}
& \mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \\
& 2^{j+1} \mathbb{E} \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \mathbb{P}(f_k(\mathbf{X}) = f^*(\mathbf{X}) | \mathbf{X} \in \mathcal{X}) \\
& 2^{j+1} \mathbb{E} \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \exp(-2c_b^2 m k^{-2}(\mathbf{x})) \\
& 2^{j+1} \mathbb{E} \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \exp(-2c_b^2 m 4^{j-1}) \quad (A25)
\end{aligned}$$

Since $\gamma = C_b[k N^{(1-\alpha)(1-a)}]^\alpha$ and $\alpha \in (2, (2+d))$, any $a \in (0, 1)$ satisfies $1-a > \frac{d \log(m)}{2(1-\alpha) \log(N)} = \frac{d}{2(1-\alpha)}$. Therefore, we can pick any $a \in (0, 1)$ and it follows that $m \leq N^{\frac{2(1-\alpha)(1-a)}{d}}$. By the choice with $k=1$, we have $\gamma = C_b N^{-\frac{(1-\alpha)(1-a)}{d}}$ and

$$\begin{aligned}
(A25) \quad & 2^{j+1} \mathbb{E} \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \exp(-\frac{1}{2} c_b^2 C_b^2 4^j m [N^{-(1-\alpha)(1-a)}]^\alpha) \\
& 2^{j+1} \mathbb{E} \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \exp(-\frac{1}{2} c_b^2 C_b^2 4^j) \\
& = 2C^{-1+\alpha} 2^{j(1+\alpha)} \exp(-\frac{1}{2} c_b^2 C_b^2 4^j)
\end{aligned}$$

Taking summation and using the fact that $\sum_{j=1}^\infty b^j \exp(-c 4^j) < \infty$ for all $b, c > 0$, we have

$$\mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X})) \mathbb{I}(\mathbf{X} \in A_j \cap E_P) \leq C^{-1+\alpha}$$

where C is some constant free of N and a . Combining the bounds above, it follows that

$$\begin{aligned}
\mathcal{R}(f_k) &= \mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X})) \\
&= \mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X}) | E_P^c) \\
&\quad + \mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X}) | \mathbf{X} \in A_0 \cap E_P) \\
&\quad + \sum_{j=1}^\infty \mathbb{E} \sum_{k=1}^2 \mathbb{I}(f_k(\mathbf{X}) = f^*(\mathbf{X}) | \mathbf{X} \in A_j \cap E_P) \\
&= \mathbb{P}(E_P^c) + C^{-1+\alpha} \\
&= \mathbb{P}(E_P^c) + C C_b^{1+\alpha} N^{-\frac{(1-\alpha)(1-a)(1+\alpha)}{d}}
\end{aligned}$$

Let $a = \frac{s \log(\log(N))}{\log(N)}$ for some $s > 0$ and notice that $N^{\frac{s \log(\log(N))}{\log(N)}} = e^{s \log(\log(N))} = [\log(N)]^s$. As a consequence, it follows that

$$N^{\frac{(1-\epsilon)(1+\epsilon)}{d}} = N^{\frac{(1-\epsilon)(1+\epsilon)}{d}} [\log(N)]^{\frac{s(1-\epsilon)(1+\epsilon)}{d}}$$

By (A1), since $2 \leq (2+d)$ and $k = 1$, if we choose s such that $2^{s(1-\epsilon)} = 48$, then

$$\begin{aligned} \mathbb{P}(E_P^c) &\leq C_D m n^{1-a} \exp(-n^a/6) \\ &\leq C_D N \exp(-n^a/6) \\ &= C_D N \exp(-N^{(1-\epsilon)a}/6) \\ &= C_D N \exp(-[\log(N)]^{s(1-\epsilon)}/6) \\ &= C_D N \exp(-\log(N)[\log(N)]^{s(1-\epsilon)-1}/6) \\ &\leq C_D N \exp(-\log(N)2^{s(1-\epsilon)-1}/6) \\ &\leq C_D N \exp(-4\log(N)) \leq C_D N^{-3} \end{aligned}$$

Combining the above three inequalities and noticing that $\frac{(1-\epsilon)(1+\epsilon)}{d} = \frac{(1+\epsilon)}{d} = \frac{1+\epsilon}{d} - \frac{1+d}{d} \leq 2$ by Assumption A3, we complete the proof with $\epsilon = \frac{s(1-\epsilon)(1+\epsilon)}{d}$ and $s = \frac{\log(48)}{(1-\epsilon)\log(2)}$. $\square$

Based on the above lemmas, we are ready to prove Theorem 3.

If $\epsilon < \frac{2}{2+d}$, it follows from Theorem 1.

If $\epsilon > \frac{2}{2+d}$, then $\frac{2}{2+d}$ for any $\epsilon \in (0, 1]$. As a consequence, we have $nN^{\frac{d}{2+d}} \log(N) = N^{1-\frac{d}{2+d}} \log(N) < 1$ and $k = 1$. The desired result follows from Lemma 9.

If $\frac{2}{2+d} \leq \epsilon \leq \frac{2}{2+d}$, the convergence rate follows from Lemma 8.

A6 Proof of Theorem 4

Theorem 4 follows directly from Lemma 9.

References

- Han, E.-H.S., Karypis, G., Kumar, V.: Text categorization using weight adjusted k-nearest neighbor classification. In: Pacific Asia Conference on Knowledge Discovery and Data Mining, pp. 53–65 (2001). Springer
- Jiang, S., Pang, G., Wu, M., Kuang, L.: An improved k-nearest-neighbor algorithm for text categorization. *Expert Systems with Applications* **39**(1), 1503–1509 (2012)
- Geng, X., Liu, T.-Y., Qin, T., Arnold, A., Li, H., Shum, H.-Y.: Query dependent ranking using k-nearest neighbor. In: Proceedings of the 31st Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, pp. 115–122 (2008)
- Kowalski, B.R., Bender, C.: k-nearest neighbor classification rule (pattern recognition) applied to nuclear magnetic resonance spectral interpretation. *Analytical Chemistry* **44**(8), 1405–1411 (1972)
- Zheng, W., Zhao, L., Zou, C.: Locally nearest neighbor classifiers for pattern classification. *Pattern recognition* **37**(6), 1307–1309 (2004)
- Xu, Y., Zhu, Q., Chen, Y., Pan, J.-S.: An improvement to the nearest neighbor classifier and face recognition experiments. *Int. J. Innov. Comput. Inf. Control* **9**(2), 543–554 (2013)
- Cai, T.T., Wei, H.: Transfer learning for nonparametric classification: Minimax rate and adaptive classifier. *Annals of Statistics* (2019). To appear
- Cover, T., Hart, P.: Nearest neighbor pattern classification. *IEEE transactions on information theory* **13**(1), 21–27 (1967)
- Cerou, F., Guyader, A.: Nearest neighbor classification in infinite dimension. *ESAIM: Probability and Statistics* **10**, 340–355 (2006) <https://doi.org/10.1051/ps:2006014>
- Hanneke, S., Kontorovich, A., Sabato, S., Weiss, R.: Universal Bayes consistency in metric spaces. *The Annals of Statistics* **49**(4), 2129–2150 (2021) <https://doi.org/10.1214/20-AOS2029>
- Stone, C.J.: Consistent nonparametric regression. *Annals of Statistics* **5**(4), 595–620 (1977)
- Devroye, L., Györfi, L., Krzyżak, A., Lugosi, G.: On the strong universal consistency of nearest neighbor regression function estimates. *Annals of Statistics* **22**(3), 1371–1385 (1994)
- Chaudhuri, K., Dasgupta, S.: Rates of convergence for nearest neighbor classification. In: *Advances in Neural Information Processing Systems*, pp. 3437–3445 (2014)

- Audibert, J.-Y., Tsybakov, A.B.: Fast learning rates for plug-in classifiers. *Annals of Statistics* **35**(2), 608–633 (2007)
- Gadat, S., Klein, T., Marteau, C.: Classification in general finite dimensional spaces with the k-nearest neighbor rule. *Annals of Statistics* **44**(3), 982–1009 (2016) <https://doi.org/10.1214/15-AOS1395>
- Samworth, R.J.: Optimal weighted nearest neighbour classifiers. *Annals of Statistics* **40**(5), 2733–2763 (2012) <https://doi.org/10.1214/12-AOS1049>
- Qiao, X., Duan, J., Cheng, G.: Rates of convergence for large-scale nearest neighbor classification. In: *Advances in Neural Information Processing Systems*, pp. 10769–10780 (2019)
- Duan, J., Qiao, X., Cheng, G.: Statistical guarantees of distributed nearest neighbor classification. *Advances in Neural Information Processing Systems* **33** (2020)
- Balsubramani, A., Dasgupta, S., Moran, S.: An adaptive nearest neighbor rule for classification. In: *Advances in Neural Information Processing Systems*, pp. 7579–7588 (2019)
- Cormen, T.H., Leiserson, C.E., Rivest, R.L., Stein, C.: *Introduction to Algorithms*. MIT press, ??? (2009)
- Huang, J.: Projection estimation in multiple regression with application to functional anova models. *Annals of Statistics* **26**(1), 242–272 (1998)
- Huang, J.: Local asymptotics for polynomial spline regression. *Annals of Statistics* **31**(5), 1600–1635 (2003)
- Lepskii, O.: On a problem of adaptive estimation in gaussian white noise. *Theory of Probability & Its Applications* **35**(3), 454–466 (1991)
- Lepski, O.V., Spokoiny, V.G.: Optimal pointwise adaptive methods in nonparametric estimation. *Annals of Statistics*, 2512–2546 (1997)
- Jiang, H.: Non-asymptotic uniform rates of consistency for k-nn regression. In: *AAAI Conference on Artificial Intelligence*, vol. 33, pp. 3999–4006 (2019)
- Zhang, Y., Duchi, J., Wainwright, M.: Divide and conquer kernel ridge regression: A distributed algorithm with minimax optimal rates. *Journal of Machine Learning Research* **16**(1), 3299–3340 (2015)
- Shang, Z., Cheng, G.: Computational limits of a distributed algorithm for smoothing spline. *Journal of Machine Learning Research* **18**(1), 3809–3845 (2017)
- Xu, G., Shang, Z., Cheng, G.: Optimal tuning for divide-and-conquer kernel ridge regression with massive data. In: *International Conference on Machine Learning*,

pp. 5483–5491 (2018). PMLR

Shang, Z., Hao, B., Cheng, G.: Nonparametric bayesian aggregation for massive data. *Journal of Machine Learning Research* **20**(140), 1–81 (2019)

Xu, G., Shang, Z., Cheng, G.: Distributed generalized cross-validation for divide-and-conquer kernel ridge regression and its asymptotic optimality. *Journal of Computational and Graphical Statistics* **28**(4), 891–908 (2019) <https://doi.org/10.1080/10618600.2019.1586714> <https://doi.org/10.1080/10618600.2019.1586714>

Dua, D., Graff, C.: UCI Machine Learning Repository (2017). <http://archive.ics.uci.edu/ml>