

Binary Classification under Local Label Differential Privacy Using Randomized Response Mechanisms

Shirong Xu

*Department of Statistics and Data Science
University of California, Los Angeles*

shirong@stat.ucla.edu

Chendi Wang

*Department of Statistics and Data Science
University of Pennsylvania*

chendi@wharton.upenn.edu

Will Wei Sun

*Daniels School of Business
Purdue University*

sun244@purdue.edu

Guang Cheng

*Department of Statistics and Data Science
University of California, Los Angeles*

guangcheng@ucla.edu

Reviewed on OpenReview: <https://openreview.net/forum?id=uKCG0w9bGG>

Abstract

Label differential privacy is a popular branch of ϵ -differential privacy for protecting labels in training datasets with non-private features. In this paper, we study the generalization performance of a binary classifier trained on a dataset privatized under the label differential privacy achieved by the randomized response mechanism. Particularly, we establish minimax lower bounds for the excess risks of the deep neural network plug-in classifier, theoretically quantifying how privacy guarantee ϵ affects its generalization performance. Our theoretical result shows: (1) the randomized response mechanism slows down the convergence of excess risk by lessening the multiplicative constant term compared with the non-private case ($\epsilon = \infty$); (2) as ϵ decreases, the optimal structure of the neural network should be smaller for better generalization performance; (3) the convergence of its excess risk is guaranteed even if ϵ is adaptive to the size of training sample n at a rate slower than $O(n^{-1/2})$. Our theoretical results are validated by extensive simulated examples and two real applications.

1 Introduction

In the past decade, differential privacy (DP; Dwork, 2008) has emerged as a standard statistical framework to protect sensitive data before releasing it to an external party. The rationale behind DP is to ensure that information obtained by an external party is robust enough to the change of a single record in a dataset. Generally, the privacy protection via DP inevitably distorts the raw data and hence reduces data utility for downstream learning tasks (Alvim et al., 2012; Kairouz et al., 2016; Li et al., 2023a). To achieve better privacy-utility tradeoff, various research efforts have been devoted to analyzing the effect of DP on learning algorithms in the machine learning community (Ghazi et al., 2021; Bassily et al., 2022; Esfandiari et al., 2022). Depending on whether the data receiver is trusted or not, differential privacy can be categorized into two main classes in the literature, including central differential privacy (central DP; Erlingsson et al. 2019; Girgis et al. 2021; Wang et al. 2023) and local differential privacy (local DP; Wang et al. 2017; Arachchige et al. 2019). Central DP relies on a trusted curator to protect all data simultaneously, whereas local DP perturbs data on the users' side. Local DP becomes a more popular solution to privacy protection due to its

successful applications, including Google Chrome browser (Erlingsson et al., 2014) and macOS (Tang et al., 2017).

An important variant of differential privacy is label differential privacy (Label DP; Chaudhuri & Hsu, 2011), which is a relaxation of DP for some real-life scenarios where input features are assumed to be publicly available and labels are highly sensitive and should be protected. Label DP has been gaining increasing attention in recent years due to the emerging demands in some real applications. For example, in recommender systems, users’ ratings are sensitive for revealing users’ preferences that can be utilized for advertising purposes (McSherry & Mironov, 2009; Xin & Jaakkola, 2014). In online advertising, a user click behavior is usually treated as a sensitive label whereas the product description for the displayed advertisement is publicly available (McMahan et al., 2013; Chapelle et al., 2014). These real scenarios motivate various research efforts to develop mechanisms for achieving label differential privacy and understanding the fundamental tradeoff between data utility and privacy protection. In literature, label DP can be divided into two main classes depending on whether labels are protected in a local or central manner. In central label DP, privacy protection is guaranteed by ensuring the output of a randomized learning algorithm is robust to the change of a single label in the dataset (Chaudhuri & Hsu, 2011; Ghazi et al., 2021; Bassily et al., 2022; Ghazi et al., 2023). In local label DP, labels are altered at the users’ side before they are released to learning algorithms, ensuring that it is difficult to infer the true labels based on the released labels (Busa-Fekete et al., 2021; Cunningham et al., 2022).

In the literature, various efforts have been devoted to developing mechanisms to achieve label DP efficiently and analyze their essential privacy-utility tradeoffs in downstream learning tasks. The original way to achieve label DP is via randomized response mechanisms (Warner, 1965), which alters observed labels in a probabilistic manner. For binary labels, the randomized response mechanism flips labels onto the other side with a pre-determined probability (Nayak & Adeshiyan, 2009; Wang et al., 2016b; Busa-Fekete et al., 2021). Originally designed to safeguard individuals’ responses in surveys (Warner, 1965; Blair et al., 2015), the RR mechanism has found extensive data collection applications (Wang et al., 2016b), such as pairwise comparisons in ranking data (Li et al., 2023b) and edges in graph data (Hehir et al., 2022; Guo et al., 2023). Ghazi et al. (2021) proposed a multi-stage training algorithm called randomized response with prior, which flips training labels via a prior distribution learned by the trained model in the previous stage. Such a multi-stage training algorithm significantly improves the generalization performance of the trained model under the same privacy guarantee. Malek Esmaeili et al. (2021) proposed to apply Laplace noise addition to one-hot encodings of labels and utilized the iterative Bayesian inference to de-noise the outputs of the privacy-preserving mechanism. Bassily et al. (2022) developed a private learning algorithm under the central label DP and established a dimension-independent deviation margin bound for the generalization performance of several differentially private classifiers, showing that the margin guarantees that are independent of the input dimension. However, their developed learning algorithm relies on the partition of the hypothesis and is hence computationally inefficient. The current state-of-the-art approach, outlined in Ghazi et al. (2023), presents a variant of RR that incorporates additional information from the loss function to enhance model performance. Analyzing this variant can be challenging due to its construction involving the solution of an optimization problem without a closed-form representation.

A critical challenge in differential privacy is to understand the essential privacy-utility tradeoff that sheds light on the fundamental utility limit for a specific problem. For example, Wang & Xu (2019) studied the sparse linear regression problem under local label DP and establish the minimax risk for the estimation error under label DP. In this paper, we intend to study the generalization performance of binary classifiers satisfying ϵ -local label DP, aiming to theoretically understand how the generalization performance of differentially private classifiers are affected by the local label DP under the margin assumption (Tsybakov, 2004). To this end, we consider the local label DP via the randomized response mechanism due to its remarkable ability to incurring less utility loss (Wang et al., 2016b). Specifically, Wang et al. (2016b) demonstrated that, while adhering to the same privacy standard, the RR mechanism exhibits smaller expected mean square errors between released and actual values compared to the Laplace mechanism. Furthermore, the effectiveness of the RR mechanism surpasses that of the output perturbation approach, particularly in scenarios where the sensitivity of output functions is high. This result can be explained from the perspective of statistical hypothesis testing that the RR mechanism achieves the optimal tradeoff between type I and type II errors

under ϵ -label DP (Dong et al., 2022). Additionally, the RR mechanism has a succinct representation, which allows us to develop certain theories related to deep learning.

In literature, few attempts are made to theoretically quantify the generalization performance of classifiers under local label DP even though binary classification has already become an indispensable part of the machine learning community. An important characteristic distinguishing binary classification problem is that the convergence of the generalization performance depends on the behavior of data in the vicinity of the decision boundary, which is known as the margin assumption (Tsybakov, 2004). Therefore, our first contribution is that we theoretically quantify how local label DP alters the margin assumption, which allows us to bridge the connection between the local label DP and the generalization performance. Additionally, we mainly consider two scenarios for the function class of classifiers in this paper. First, we consider the large margin classifier with its hypothesis space being a parametric function class and the deep neural network plug-in classifier. For these two scenarios, we establish their upper bound and the minimax lower bound for their excess risks, which theoretically quantifies how ϵ affects the generalization performance. The implications of our theoretical results are three-fold. First, the Bayes classifier stays invariant to the randomized response mechanism with any small ϵ , which permits the possibility of learning the optimal classifier from the privatized dataset. Second, the local label DP achieved via the randomized response mechanism implicitly reduces the information for estimation. Specifically, we theoretically prove that the convergence rate of excess risk is slowed down with an additional multiplicative constant depending on ϵ . Third, based on our theoretical results, we show that the excess risk fails to converge when the ϵ is adaptive to the training sample size n at the order $O(n^{-1/2})$, which is independent of the margin assumption. To the best of our knowledge, no existing literature related to classification under DP or corrupted labels (Cannings et al., 2020; van Rooyen & Williamson, 2018) has investigated the effects of noise on neural network structures. Our theoretical results are supported by extensive simulations and two real applications.

The rest of this paper is organized as follows. After introducing some necessary notations in Section 1.1, Section 2 introduces the backgrounds of the binary classification problem, neural network, and the local label differential privacy. In Section 3, we introduce the framework of the differentially private learning under the label differential privacy and present theoretical results regarding the asymptotic behavior of the differentially private classifier. Section 4 quantifies how ϵ affects the generalization performance of the deep neural network plug-in classifier by establishing a minimax lower bound. Section 5 and Section 6 conduct a series of simulations and real applications to support our theoretical results. All technical proofs are provided in the Appendix.

1.1 Notation

For a vector $\mathbf{x} \in \mathbb{R}^p$, we denote its l_1 -norm and l_2 -norm as $\|\mathbf{x}\|_1 = \sum_{i=1}^p |x_i|$ and $\|\mathbf{x}\|_2 = (\sum_{i=1}^p |x_i|^2)^{1/2}$, respectively. For a function $f: \mathcal{X} \rightarrow \mathbb{R}$, we denote its L_p -norm with respect to the probability measure μ as $\|f\|_{L^p(\mu)} = (\int_{\mathcal{X}} |f(\mathbf{x})|^p d\mu(\mathbf{x}))^{1/p}$. For a real number a , we let $\lfloor a \rfloor$ denote the largest integer not larger than a . For a set S , we define $\mathcal{N}(\xi, S, \|\cdot\|)$ as the minimal number of ξ -balls needed to cover S under a generic metric $\|\cdot\|$. For two given sequences $\{A_n\}_{n \in \mathbb{N}}$ and $\{B_n\}_{n \in \mathbb{N}}$, we write $A_n \gtrsim B_n$ if there exists a constant $C > 0$ such that $A_n \geq CB_n$ for any $n \in \mathbb{N}$. Additionally, we write $A_n \asymp B_n$ if $A_n \gtrsim B_n$ and $A_n \lesssim B_n$.

2 Preliminaries

2.1 Binary Classification

The goal in binary classification is to learn a discriminant function f , which well characterizes the functional relationship between the feature vector $\mathbf{X} \in \mathcal{X}$ and its associated label $Y \in \{-1, 1\}$. To measure the quality of f , the 0-1 risk is usually employed,

$$R(f) = \mathbb{E}(I(f(\mathbf{X})Y < 0)) = \mathbb{P}(\text{sign}(f(\mathbf{X})) \neq Y),$$

where $I(\cdot)$ denotes the indicator function and the expectation is taken with respect to the joint distribution of $(\mathbf{X}, Y)$.

In addition to 0-1 loss, the false negative error (FNE) and false positive error (FPE) are frequently used to assess classifier performance when dealing with highly imbalanced datasets. Specifically, FNE measures the percentage of positive samples being classified as negative, whereas FPE measures the percentage of negative samples being classified as positive. For a margin classifier f , the expected false negative error (EFNE) and false positive error (EFPE) are written as

$$\begin{aligned}\text{EFNE}(f) &= \mathbb{E}\left(I(f(\mathbf{X}) < 0) | Y = 1\right) = \mathbb{P}\left(\text{sign}(f(\mathbf{X})) \neq Y | Y = 1\right), \\ \text{EFPE}(f) &= \mathbb{E}\left(I(f(\mathbf{X}) > 0) | Y = -1\right) = \mathbb{P}\left(\text{sign}(f(\mathbf{X})) \neq Y | Y = -1\right).\end{aligned}$$

Furthermore, the 0-1 risk $R(f)$ can be re-written as a weighted combination of EFNE and EFPE:

$$R(f) = \text{EFNE}(f)\mathbb{P}(Y = 1) + \text{EFPE}(f)\mathbb{P}(Y = -1).$$

Let $f^* = \inf_f R(f)$ denote the minimizer of $R(f)$, which refers to as the Bayes decision rule. Generally, f^* is obtained by minimizing $R(f)$ in a point-wise manner and given as $f^*(\mathbf{X}) = \text{sign}(\eta(\mathbf{X}) - 1/2)$ with $\eta(\mathbf{X}) = \mathbb{P}(Y = 1 | \mathbf{X})$. The minimal risk $R(f^*)$ can be written as $R(f^*) = \mathbb{E}_{\mathbf{X}}(\min\{\eta(\mathbf{X}), 1 - \eta(\mathbf{X})\})$. In practice, the underlying joint distribution on $(\mathbf{X}, Y)$ is unavailable, but a set of i.i.d. realizations $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$ is given. Therefore, it is a common practice to consider the estimation procedure based on minimizing the sample average of a surrogate loss, which is given as

$$\hat{R}_\phi(f) = \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i)y_i),$$

where $\phi(\cdot)$ is the surrogate loss function replacing the 0-1 loss since the 0-1 loss is computationally intractable (Arora et al., 1997).

Let $R_\phi(f) = \mathbb{E}(\phi(f(\mathbf{X})Y))$ denote the ϕ -risk. Given that ϕ is classification calibrated, the minimizer $f_\phi^* = \arg \min_f R_\phi(f)$ is consistent with the Bayes decision rule (Lin, 2004), i.e., $\text{sign}(f_\phi^*(\mathbf{x})) = \text{sign}(\eta(\mathbf{x}) - 1/2)$ for any $\mathbf{x} \in \mathcal{X}$. In literature, there are various classification-calibrated loss functions (Zhang, 2004; Bartlett et al., 2006), such as exponential loss, hinge loss, logistic loss, and ψ -loss (Shen et al., 2003).

2.2 Deep Neural Network and Function Class

Let $f(\mathbf{x}; \Theta)$ be an L -layer neural network with Rectified Linear Unit (ReLU) activation function, that is,

$$f(\mathbf{x}; \Theta) = \mathbf{A}_{L+1}(\mathbf{h}_L \circ \mathbf{h}_{L-1} \circ \cdots \circ \mathbf{h}_1(\mathbf{x})) + \mathbf{b}_{L+1},$$

where $\circ$ denotes function composition, $\mathbf{h}_l(\mathbf{x}) = \sigma(\mathbf{A}_l \mathbf{x} + \mathbf{b}_l)$ denotes the l -th layer, and $\Theta = \{(\mathbf{A}_l, \mathbf{b}_l)\}_{l=1, \dots, L+1}$ denotes all the parameters. Here $\mathbf{A}_l \in \mathbb{R}^{p_l \times p_{l-1}}$ is the weight matrix, $\mathbf{b}_l \in \mathbb{R}^{p_l}$ is the bias term, p_l is the number of neurons in the l -th layer, and $\sigma(x) = \max\{0, x\}$ is the ReLU function. To characterize the network architecture of f , we denote the number of layers in Θ as $\Upsilon(\Theta)$, the maximum number of nodes as $\Delta(\Theta)$, the number of non-zero parameters as $\|\Theta\|_0$, the largest absolute value in Θ as $\|\Theta\|_\infty$. For a given n , we denote by $\mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)$ a class of neural networks, which is defined as

$$\begin{aligned}\mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n) &= \{f(\mathbf{x}; \Theta) : \Upsilon(\Theta) \leq L_n, \Delta(\Theta) \leq N_n, \\ &\quad \|\Theta\|_0 \leq P_n, \|\Theta\|_\infty \leq B_n, \|f\|_\infty \leq V_n, \}.\end{aligned}$$

Let $\beta > 0$ be a degree of smoothness, then the Hölder space is defined as

$$\mathcal{H}(\beta, \mathcal{X}) = \{f \in \mathcal{C}^{[\beta]}(\mathcal{X}) : \|f\|_{\mathcal{H}(\beta, \mathcal{X})} < \infty\},$$

where $\mathcal{C}^{[\beta]}(\mathcal{X})$ the class of $[\beta]$ times continuously differentiable functions on the open set $\mathcal{X}$ and the Hölder norm $\|f\|_{\mathcal{H}(\beta, \mathcal{X})}$ is given as

$$\|f\|_{\mathcal{H}(\beta, \mathcal{X})} = \max_{\mathbf{m}: \|\mathbf{m}\|_1 \leq [\beta]} \sup_{\mathbf{x} \in \mathcal{X}} |\partial^{\mathbf{m}} f(\mathbf{x})| + \max_{\mathbf{m}: \|\mathbf{m}\|_1 = [\beta]} \sup_{\mathbf{x}, \mathbf{y} \in \mathcal{X}, \mathbf{x} \neq \mathbf{y}} \frac{|\partial^{\mathbf{m}} f(\mathbf{x}) - \partial^{\mathbf{m}} f(\mathbf{y})|}{\|\mathbf{x} - \mathbf{y}\|_2^{\beta - [\beta]}},$$

where $\partial^{\mathbf{m}} f(\mathbf{x}) = \frac{\partial^{m_1+\dots+m_p}}{\partial x_1^{m_1} \dots \partial x_p^{m_p}} f(\mathbf{x})$ denotes the partial derivative of order $\mathbf{m}$ with respect to $\mathbf{x}$ and $\mathbf{m} = (m_1, \dots, m_p) \in \mathbb{N}_0^p$ is a multi-index with $\mathbb{N}_0 = \mathbb{N} \cup \{0\}$. Further, we let $\mathcal{H}(\beta, \mathcal{X}, M) = \{f \in \mathcal{H}(\beta, \mathcal{X}) : \|f\|_{\mathcal{H}(\beta, \mathcal{X})} \leq M\}$ be a closed ball of radius M in $\mathcal{H}(\beta, \mathcal{X})$.

The Hölder assumption serves as a common assumption for studying the generalization capabilities of neural networks for approximating functions possessing a certain degree of smoothness or regularity (Audibert & Tsybakov, 2007; Kim et al., 2021). This assumption is useful in developing a tighter generalization bounds. However, it is essential to acknowledge that the Hölder assumption presupposes a level of smoothness in the underlying function. In reality, many real-world problems involve functions that deviate from this idealized smoothness. When the actual smoothness of the function does not align with the assumption, algorithms reliant on it may fall short of the anticipated generalization performance.

2.3 Local Label Differential Privacy

Label differential privacy (Label DP; Chaudhuri & Hsu, 2011) is proposed as a relaxation of differential privacy (Dwork, 2008), aiming to protect the privacy of labels in the dataset, whereas training features are non-sensitive and hence publicly available. An effective approach to label DP is the randomized response mechanism (Busa-Fekete et al., 2021). As its name suggests, it protects the privacy of labels via some local randomized response mechanisms under the framework of local differential privacy.

Definition 1. (*Label Differential Privacy; Ghazi et al., 2021*) A randomized training algorithm $\mathcal{A}$ taking as input a dataset is said to be (ϵ, δ) -label differentially private if for any two datasets $\mathcal{D}$ and $\mathcal{D}'$ that differ in the label of a single example, and for any subset S of the outputs of $\mathcal{A}$, it is the case that $\mathbb{P}(\mathcal{A}(\mathcal{D}) \in S) \leq \epsilon \mathbb{P}(\mathcal{A}(\mathcal{D}') \in S) + \delta$, then $\mathcal{A}$ is said to be ϵ -label differentially private (ϵ -label DP).

A direct way to achieve ϵ -Label DP in binary classification is to employ the randomized response mechanism (Warner, 1965; Nayak & Adeshiyan, 2009; Wang et al., 2016b; Karwa et al., 2017). The main idea of the binary randomized response mechanism is to flip observed labels independently with a fixed probability. Specifically, let $\mathcal{A}_\theta$ denote the randomized response mechanism parametrized by θ . For an input label Y , we define

$$\mathcal{A}_\theta(Y) = \begin{cases} Y, & \text{with probability } \theta, \\ -Y, & \text{with probability } 1 - \theta, \end{cases}$$

where $\theta > 1/2$ denotes the probability that the value of Y stays unchanged. It is straightforward to verify that the randomized response mechanism satisfies ϵ -label DP with $\epsilon = \log(\theta/(1 - \theta))$ (Ghazi et al., 2021; Busa-Fekete et al., 2021). In this paper, we denote the ϵ -label DP achieved through the randomized response mechanism as ϵ -local label DP.

3 Differentially Private Learning

3.1 Effect of Locally-Label Differential Privacy

Under the setting of local differential privacy, users do not trust the data curator and hence privatize their sensitive data via some randomized mechanisms locally before releasing them to central servers. We let $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$ denote the original dataset containing n i.i.d. realizations of the random pair $(\mathbf{X}, Y)$. As illustrated in Figure 1, users' labels are privatized by a locally differentially private protocol $\mathcal{A}_\theta$, and then the untrusted curator receives a privatized dataset, which we denote as $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, \tilde{y}_i)\}_{i=1}^n$.

A natural question is what the quantitative relation between ϵ -label DP and the discrepancy between the distributions of $\mathcal{D}$ and $\tilde{\mathcal{D}}$ is. This is useful to analyze how privacy parameter ϵ deteriorates the utility of data for downstream learning tasks.

Notice that $\tilde{y}_i$'s are generated locally and independently, therefore the randomized response mechanism only alters the conditional distribution of Y given $\mathbf{X}$, whereas the marginal distribution of $\mathbf{X}$ stays unchanged. Hence, we can assume $\tilde{\mathcal{D}}$ is a set of i.i.d. realizations of $(\mathbf{X}, \tilde{Y})$, and it is straightforward to verify that the

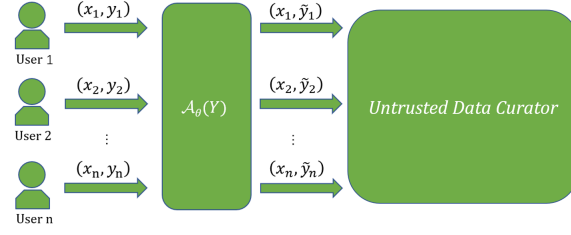


Figure 1: The framework of local label differential privacy.

conditional distribution of $\tilde{Y}$ given $\mathbf{X} = \mathbf{x}$ can be written as

$$\tilde{\eta}(\mathbf{x}) = \mathbb{P}(\tilde{Y} = 1 | \mathbf{X} = \mathbf{x}) = \theta\eta(\mathbf{x}) + (1 - \theta)(1 - \eta(\mathbf{x})).$$

Clearly, $\mathcal{A}_\theta$ amplifies the uncertainty of the observed labels by shrinking $\tilde{\eta}(\mathbf{x})$ towards $1/2$. Particularly, $\tilde{\eta}(\mathbf{X}) = 1/2$ almost surely when $\theta = 1/2$. In this case, the privatized dataset $\tilde{\mathcal{D}}$ conveys no information for learning the decision function.

The following Lemma quantifies how the randomized response mechanism alters the conditional distribution of Y and $\tilde{Y}$ given the feature $\mathbf{x}$ via the Kullback-Leibler divergence (KL) divergence between these two Bernoulli distributions. Specifically, the inequalities in (1) explicates how the divergence changes with privacy guarantee ϵ .

Lemma 1. *Suppose the randomized response mechanism $\tilde{Y} = \mathcal{A}_\theta(Y)$ satisfies ϵ -label DP with $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$, then for any $\mathbf{x} \in \mathcal{X}$ it holds that*

$$\mathcal{L}(\epsilon, \mathbf{x}) \leq D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) \leq \mathcal{U}(\epsilon, \mathbf{x}), \quad (1)$$

where $D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}})$ denotes the Kullback-Leibler divergence (KL) divergence between $\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}$ and $\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}$, $\mathcal{L}(\epsilon, \mathbf{x}) = 2(2\eta(\mathbf{x}) - 1)^2(1 + \exp(\epsilon))^{-2}$, and $\mathcal{U}(\epsilon, \mathbf{x}) = \min\{U_1(\epsilon, \mathbf{x}), U_2(\epsilon, \mathbf{x})\}$ with $U_1(\epsilon, \mathbf{x}) = (2\eta(\mathbf{x}) - 1)^2\eta^{-1}(\mathbf{x})(1 - \eta(\mathbf{x}))^{-1}(1 + \exp(\epsilon))^{-2}$ and $U_2(\epsilon, \mathbf{x}) = (2\eta(\mathbf{x}) - 1)^2\exp(-\epsilon)$.

In Lemma 1, the lower bound $\mathcal{L}(\epsilon, \mathbf{x})$ and the upper bound $\mathcal{U}(\epsilon, \mathbf{x})$ share a factor $(1 + \exp(-\epsilon))^2$, indicating that $D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}})$ decreases exponentially with respect to ϵ .

3.2 Differentially Private Classifier

On the side of the untrusted curator, inference tasks vary according to the purpose of data collector, such as estimation of population statistics (Joseph et al., 2018; Yan et al., 2019) and supervised learning (Ghazi et al., 2021; Esfandiari et al., 2022). In this paper, we suppose that the computation based on $\tilde{\mathcal{D}}$ implemented by the untrusted curator is formulated as the following regularized empirical risk minimization task,

$$\min_{f \in \mathcal{F}} L_n(f) = \min_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i)\tilde{y}_i) + \lambda_n J(f), \quad (2)$$

where ϕ is a surrogate loss, λ_n is a tuning parameter, $J(\cdot)$ is a penalty term, and $\mathcal{F}$ is a pre-specified hypothesis space.

We denote by $\tilde{R}(f) = \mathbb{E}[\text{sign}(f(\mathbf{X})) \neq \tilde{Y}]$ the risk with expectation taken with respect to the joint distribution of $(\mathbf{X}, \tilde{Y})$ and let $\tilde{f}^* = \arg \min_f \tilde{R}(f)$ denote the Bayes decision rule under the distribution of $(\mathbf{X}, \tilde{Y})$. The excess risk of f under the distributions of $(\mathbf{X}, Y)$ and $(\mathbf{X}, \tilde{Y})$ is denoted as $D(f, f^*) = R(f) - R(f^*)$ and $\tilde{D}(f, \tilde{f}^*) = \tilde{R}(f) - \tilde{R}(\tilde{f}^*)$, respectively.

Lemma 2. *If $\tilde{Y} = \mathcal{A}_\theta(Y)$ with $\theta > 1/2$, then $f^*(\mathbf{x}) = \tilde{f}^*(\mathbf{x})$ for any $\mathbf{x} \in \mathcal{X}$ and $\tilde{D}(f, \tilde{f}^*) = (2\theta - 1)D(f, f^*)$ for any f .*

Lemma 2 shows that the Bayes decision rule stays invariant under the randomized response mechanism. It is clear to see that $\tilde{D}(f, f^*)$ diminishes as θ gets close to $1/2$. Particularly, $\tilde{D}(f, f^*)$ is equal to 0 when $\theta = 1/2$. This is as expected since $\theta = 1/2$ implies that $\tilde{\eta}(\mathbf{X}) = 1/2$ almost surely, where all classifiers deteriorates in performance simultaneously. This result demonstrates that the optimal classifiers for the underlying distributions of $\mathcal{D}$ and $\tilde{\mathcal{D}}$ are identical. Basically, Lemma 2 reveals an inherent debiasing process in the development of the generalization performance of the differentially private classifier $\tilde{f}$. To be more specific, we can deduce the excess risk of $\tilde{f}$ by the relation $\tilde{R}(\tilde{f}) - \tilde{R}(f^*) = \frac{e^\epsilon - 1}{e^\epsilon + 1}(R(\tilde{f}) - R(f^*))$ under ϵ -local label DP.

Let $f_\phi^* = \arg \min_f R_\phi(f)$ and $\tilde{f}_\phi^* = \arg \min_f \tilde{R}_\phi(f)$ be the minimizers of ϕ -risk under the joint distributions of $(\mathbf{X}, Y)$ and $(\mathbf{X}, \tilde{Y})$, respectively. For illustration, we only consider hinge loss $\phi(x) = \max\{1 - x, 0\}$ in this paper, and similar results can be easily obtained for other loss functions by the comparison theorem (Bartlett et al., 2006; Bartlett & Wegkamp, 2008). With $\phi(\cdot)$ being the hinge loss, we have $\tilde{f}_\phi^* = f_\phi^* = f^* = \tilde{f}^*$ (Bartlett et al., 2006; Lecué, 2007). With a slight abuse of notation, we use f^* to refer to these four functions simultaneously in the sequel.

3.3 Consistency under the Randomized Response Mechanism

In this section, we establish the asymptotic behavior of the classifier trained on privatized datasets. Specifically, we theoretically quantify how the randomized response mechanism affects the convergence rate of excess risk under the low-noise assumption (Lecué, 2007).

We suppose that the randomized response mechanism satisfies the ϵ -label DP, which indicates that $\epsilon = \log(\theta/(1 - \theta))$. Furthermore, we denote that $\tilde{f}_n = \arg \min_{f \in \mathcal{F}} L_n(f)$ and $H(f, f^*) = R_\phi(f) - R_\phi(f^*)$. In classification problems, $H(f, f^*)$ is an important metric that admits the decomposition into the estimation error and the approximation error (Bartlett et al., 2006),

$$H(f, f^*) = H(f, f_{\mathcal{F}}^*) + H(f_{\mathcal{F}}^*, f^*),$$

where $f_{\mathcal{F}}^* = \arg \min_{f \in \mathcal{F}} R_\phi(f)$. Here $H(f, f_{\mathcal{F}}^*)$ is the estimation error and $H(f_{\mathcal{F}}^*, f^*)$ is the approximation error. The estimation error depends on the learning algorithm in finding $f_{\mathcal{F}}^*$ based on a dataset with finite samples, where the searching difficulty is unavoidably affected by the complexity of $\mathcal{F}$ as the approximation error. Generally, the richness of $\mathcal{F}$ can be measured by VC dimension (Blumer et al., 1989; Vapnik & Chervonenkis, 2015), Rademacher complexity (Bartlett & Mendelson, 2002; Smale & Zhou, 2003), and metric entropy methods (Zhou, 2002; Shalev-Shwartz & Ben-David, 2014).

Assumption 1. (*Low-noise assumption*) *There exists a constant $c > 0$ and $0 < \gamma \leq +\infty$ such that $\mathbb{P}(|2\eta(\mathbf{X}) - 1| \leq t) \leq ct^\gamma$, for any $t \in [0, 1)$.*

Assumption 1 is known as the low-noise assumption in the binary classification (Lecué, 2007; Shen et al., 2003; Bartlett et al., 2006), which characterizes the behavior of $2\eta(\mathbf{x}) - 1$ around the decision boundary $\eta(\mathbf{x}) = 1/2$. Particularly, the case with $\gamma = +\infty$ and $c = 1$ implies that the labels y_i 's are deterministic in that $\eta(\mathbf{X})$ takes values in $\{0, 1\}$ almost surely, resulting in the fastest convergence rate of the estimation error.

Lemma 3. *Denote that $\kappa_\epsilon = (e^\epsilon - 1)/(e^\epsilon + 1)$. Under Assumption 1, for any $\theta \in (1/2, 1]$, it holds that*

- (1) $\mathbb{P}(|2\tilde{\eta}(\mathbf{X}) - 1| \leq t) \leq c\kappa_\epsilon^{-\gamma}t^\gamma$, for any $t \in [0, 1)$,
- (2) $\tilde{R}_\phi(f) - \tilde{R}_\phi(f^*) \geq \kappa_\epsilon(4c)^{-1/\gamma} \left(\mathbb{E}[|f(\mathbf{X}) - f^*(\mathbf{X})|] \right)^{\frac{\gamma+1}{\gamma}}$ for any $f \in \mathcal{F}$,

Lemma 3 presents some insights regarding the influence of the randomized response mechanism on the low-noise assumption and the margin relation (Lecué, 2007), showing that the low-noise structure and margin relation are both invariant to the randomized response mechanism in the sense that only the multiplicative terms are enlarged by the effect of the privacy guarantee ϵ .

Assumption 2. We assume that $\mathcal{F}$ is properly chosen satisfying that $\|f\|_\infty \leq 1$ for any $f \in \mathcal{F}$ and

$$\log \mathcal{N}(\xi, \mathcal{F}, \|\cdot\|_{L^2(\mu)}) \asymp \mathcal{V}_1(\Theta) \log(1 + \xi^{-1} \mathcal{V}_2(\Theta)), \text{ and } VC(\mathcal{G}_{\mathcal{F}}) \asymp \mathcal{V}_1(\Theta)$$

where $\mathcal{G}_{\mathcal{F}} = \{\{\mathbf{x} : \text{sign}(f(\mathbf{x})) = 1\} : f \in \mathcal{F}\}$, $VC(\mathcal{G}_{\mathcal{F}})$ denotes the VC dimension of $\mathcal{G}_{\mathcal{F}}$, μ denotes the marginal distribution of $\mathbf{X}$, Θ denotes the parameters of f , and $\mathcal{V}_1(\Theta)$ and $\mathcal{V}_2(\Theta)$ are some functions depending on Θ .

Assumption 2 characterizes the complexity of function class $\mathcal{F}$ through the metric entropy (Zhou, 2002; Bousquet et al., 2003; Lei et al., 2016), where $\mathcal{V}_1(\Theta)$ and $\mathcal{V}_2(\Theta)$ are quantities increasing with the size of Θ . Assumption 2 generally holds for function classes of parametric models (Wang et al., 2016a; Xu et al., 2021), most notably for deep neural networks (Bartlett et al., 2019; Schmidt-Hieber, 2020). Additionally, Assumption 2 also holds for those VC classes with VC dimensions increasing with the size of Θ (Wellner et al., 2013; Bartlett et al., 2019; Lee et al., 1994).

Theorem 1. Under Assumptions 1 and 2, for any minimizer $\tilde{f}_n$ of (2), there exist some positive constants A_1 and A_2 such that

$$A_2 \left\{ \left(\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2} \right)^{\frac{\gamma+1}{\gamma+2}} + \tau_n \right\} \leq \sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{f}_n) - R(f^*)] \leq A_1 \left\{ \left(\frac{\mathcal{V}_1(\Theta) \log(n)}{n\kappa_\epsilon^2} \right)^{\frac{\gamma+1}{\gamma+2}} + s_n \right\},$$

where $\mathcal{P}_\gamma$ be a class of distributions of $(\mathbf{X}, Y)$ satisfying Assumption 1, $s_n = \sup_{\pi \in \mathcal{P}_\gamma} \inf_{f \in \mathcal{F}} H(f, f^*)$, and $\tau_n = \sup_{\pi \in \mathcal{P}_\gamma} \inf_{f \in \mathcal{F}} D(f, f^*)$.

Theorem 1 quantifies the asymptotic behavior of $\tilde{f}_n$ by establishing its upper and lower bounds, which explicitly demonstrates the quantitative relation between the effect of privacy guarantee ϵ and the excess risk. The upper bound is proven through a uniform concentration inequality under the framework of empirical risk minimization analysis. The estimator $\tilde{f}_n$ is derived from a pre-specified function class $\mathcal{F}$ that may not include the true underlying function f^* . Consequently, s_n represents the approximation error, quantifying the capability of the optimal function within $\mathcal{F}$ to approximate f^* . The accuracy of estimating the optimal function in $\mathcal{F}$ for approximating f^* depends on the complexity of $\mathcal{F}$ (Assumption 2) and the size of the training set n . A larger $\mathcal{F}$ may reduce s_n but can increase the estimation error. Thus, achieving the best convergence rate involves striking the right balance between these two sources of error. The proof of the lower bound is similar to Theorem 2 in Lecué (2007), mainly applying the Assouad's lemma (Yu, 1997) to an analytic subset of $\mathcal{P}_\gamma$. Further details regarding the proof are provided in the Appendix.

It is worth noting that the upper bound matches the lower bound except for a logarithmic factor when the approximation error term $s_n \lesssim \left(\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2} \right)^{\frac{\gamma+1}{\gamma+2}}$. This shows that the randomized response mechanism slows down the convergence rate by enlarging the multiplicative constant. Moreover, based on Theorem 1, we can obtain the optimal convergence rate of the excess risk in classification problem under the low-noise assumption (Lecué, 2007) by setting $\epsilon = \infty$ and $|\mathcal{F}| < \infty$.

Lemma 4. Denote that $\mathcal{M} = \max\{EFNE(f) - EFNE(f^*), EFPE(f) - EFPE(f^*)\}$. Under Assumption 1, for any margin classifier f , it holds that

$$R(f) - R(f^*) \leq \mathcal{M} \leq \frac{1}{2 \min\{\mathbb{P}(Y=1), \mathbb{P}(Y=-1)\}} \left(2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*) \right), \quad (3)$$

where c and γ are as defined in Assumption 1. Particularly, if Assumption 1 holds with $\gamma = \infty$, for any margin classifier f , (3) becomes

$$R(f) - R(f^*) \leq \mathcal{M} \leq \frac{3}{2 \min\{\mathbb{P}(Y=1), \mathbb{P}(Y=-1)\}} (R(f) - R(f^*)).$$

Lemma 4 establishes a crucial relationship between excess risk and the maximum of excess EFNE and EFPE for any classifier f , which also includes $\tilde{f}_n$ as a special case. This connection enables us to demonstrate the convergence of $EFNE(\tilde{f}_n)$ and $EFPE(\tilde{f}_n)$ based on that of the excess risk. Remarkably, this finding

implies that the differentially private classifier $\tilde{f}_n$ will exhibit similar false negative and false positive rates to those of the Bayes classifier as the sample size n tends to infinity. In addition, increasing the degree of data imbalance will enlarge the upper bound in (3), indicating that data imbalance slows down the convergence rates of EFNE($\tilde{f}_n$) and EFPE($\tilde{f}_n$). Particularly, under the low-noise assumption with $\gamma = \infty$, which implies that samples are separable, both excess EFNE and EFPE of $\tilde{f}_n$ exhibit the same convergence rate as the excess risk regardless of the degree of class imbalance. Furthermore, this result indicates that the privacy guarantee ϵ has a similar impact on the excess EFNE and EFPE as it does on the excess risk.

4 Deep Learning with Local Label Differential Privacy

A typical class of models that are popularly considered in the domain of differential privacy is deep neural network (Ghazi et al., 2021; Yuan et al., 2021) due to its success in various applications in the past decade. Unlike Section 3.3 considering the estimation and approximation errors separately, we establish theoretical results regarding the convergence rate of the excess risk of the deep neural network plug-in classifier trained from $\tilde{\mathcal{D}}$, which is obtained by making a tradeoff between the estimation and approximation errors of deep neural networks (Schmidt-Hieber, 2020). Our theoretical results not only quantify how the optimal structure of the deep neural network changes with the privacy parameter ϵ , but also derive the optimal privacy guarantee we can achieve for the deep neural network classifier.

Remind that $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, \tilde{y}_i)\}_{i=1}^n$ is the privatized dataset with $\tilde{y}_i = \mathcal{A}_\theta(y_i)$ and $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$. The deep neural network is solved as

$$\tilde{f}_{nn} = \arg \min_{f \in \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)} \frac{1}{n} \sum_{i=1}^n (f(\mathbf{x}_i) - \tilde{z}_i)^2, \quad (4)$$

where $\tilde{z}_i = (\tilde{y}_i + 1)/2$ and $\mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)$ is a class of multilayer perceptrons defined in Section 2.2. The plug-in classifier based on $\tilde{f}_{nn}$ can be obtained as $\tilde{s}_{nn} = \text{sign}(\tilde{f}_{nn} - 1/2)$. To quantify the asymptotic behavior of $\tilde{s}_{nn}$, we further assume that the support of $\mathbf{x}$ is $[0, 1]^p$, which is a common assumption for deep neural networks (Yarotsky, 2017; Nakada & Imaizumi, 2020)

Theorem 2. *Let $\mathcal{P}_{\gamma, \beta}$ be a class of probability measures on $\mathcal{X} \times \{-1, 1\}$ satisfying Assumption 1 and $\eta(\mathbf{X}) \in \mathcal{H}(\beta, [0, 1]^p, M)$. For any minimizer $\tilde{f}_{nn}$ in (4) with $L_n \asymp \log(\kappa_\epsilon n / \log(n))$, $N_n \asymp (\kappa_\epsilon n / \log(n))^{\frac{2p}{2\beta+p}}$, $B_n = 1$, and $P_n \asymp N_n \log(\kappa_\epsilon n / \log(n))$, we have*

$$\left(\frac{1}{n\kappa_\epsilon^2} \right)^{\frac{\beta(\gamma+1)}{\beta(\gamma+2)+p}} \lesssim \sup_{\pi \in \mathcal{P}_{\gamma, \beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] \lesssim \left(\frac{\log n}{n\kappa_\epsilon^2} \right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}}. \quad (5)$$

Particularly, $\sup_{\pi \in \mathcal{P}_{\gamma, \beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] = o(1)$ given that $\epsilon \gtrsim n^{-1/2+\zeta}$ for any $\zeta > 0$.

In Theorem 2, we quantify the asymptotic behavior of the excess risk of $\tilde{s}_{nn}$ by providing upper and lower bounds for $\sup_{\pi \in \mathcal{P}_{\gamma, \beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R^*]$. Similar to Theorem 1, the proof of the lower bound of (5) relies on the Assouad's lemma. A significant distinction of the upper bound in (5) from that of Theorem 1 is explicating the approximation error s_n with respect to the structure of neural network and making the optimal tradeoff between estimation and approximation errors to achieve the fastest convergence rate. It should be noted that if $\epsilon = \infty$ which refers to the non-private case, the upper and lower bounds in (5) match with existing theoretical results established in Audibert & Tsybakov (2007). Moreover, Theorem 2 goes a step further by precisely characterizing the impact of ϵ on the convergence of $\tilde{s}_{nn}$ to the Bayes decision rule. Additionally, it specifies how the optimal neural network's structure contracts as ϵ decreases, crucial for achieving the fastest convergence rate. Specifically, attaining this rapid convergence necessitates reducing the maximum number of hidden units at an order of $O(\epsilon^{2p/(2\beta+p)})$ when compared to the non-private case. Furthermore, we leverage Theorem 2 to derive the fastest adaptive rate of ϵ under the consistency of $\tilde{s}_{nn}$. Specifically, we find that $\epsilon \gtrsim n^{-1/2+\zeta}$ for any $\zeta > 0$ to achieve this desired consistency rate. This result represents a crucial step towards understanding the interplay between privacy and performance in our framework.

5 Simulated Experiments

This section aims to validate our theoretical results through extensive simulated examples. Specifically, we show that the excess risk of a differentially private classifier converges to 0 for any fixed ϵ , whereas the convergence is not achievable as long as ϵ is adaptive to the size of the training dataset with some properly chosen orders as shown in Theorems 1 and 2.

5.1 Support Vector Machine

This simulation experimentally analyzes the effect of label DP on the SVM classifier. The generation of simulated datasets is as follows. First, we set the regression function in classification as $\eta(\mathbf{x}) = 1/(1 + \exp(-\beta_0^T \mathbf{x}))$, where $\beta_0 \in \mathbb{R}^p$ and $\mathbf{x}$ are both p -dimensional vectors generated via $\beta_{0i}, x_i \sim \text{Unif}(-1, 1)$ for $i = 1, \dots, p$. For each feature $\mathbf{x}$, its label y is chosen from $\{1, -1\}$ with probabilities $\eta(\mathbf{x})$ and $1 - \eta(\mathbf{x})$, respectively. Repeating the above process n times, we obtain a non-private training dataset $\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$. Subsequently, we apply the randomized response mechanism to generate $\tilde{y}_i$ as $\tilde{y}_i = \mathcal{A}_\theta(y_i)$, where $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$ and ϵ is the privacy guarantee. Therefore, the obtained privatized training dataset $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, \tilde{y}_i)\}_{i=1}^n$ satisfies ϵ -Label DP. Based on $\tilde{\mathcal{D}}$, we obtain the SVM classifier (Cortes & Vapnik, 1995),

$$\tilde{f}_n = \arg \min_{\beta} \frac{1}{n} \sum_{i=1}^n (1 - \tilde{y}_i \beta^T \mathbf{x}_i)_+ + \lambda \|\beta\|_2^2,$$

where $(x)_+ = \max\{x, 0\}$. Next, we evaluate the performance of $\tilde{f}_n$ in terms of the empirical excess risk and the classification error,

$$\begin{aligned} \hat{E}(\tilde{f}_n) &= \frac{1}{n_{\text{test}}} \sum_{i=1}^{n_{\text{test}}} I\left(\text{sign}(\tilde{f}_n(\mathbf{x}'_i)) \neq \text{sign}(\eta(\mathbf{x}'_i) - 1/2)\right) |2\eta(\mathbf{x}'_i) - 1|, \\ \text{CE}(\tilde{f}_n) &= \frac{1}{n_{\text{test}}} \sum_{i=1}^{n_{\text{test}}} I\left(\text{sign}(\tilde{f}_n(\mathbf{x}'_i)) \neq \text{sign}(\eta(\mathbf{x}'_i) - 1/2)\right). \end{aligned}$$

where $\mathbf{x}'_i$'s are testing samples generated in the same way as $\mathbf{x}_i$'s.

Scenario I. In the first scenario, we aim to verify that $\hat{E}(\tilde{f}_n)$ will converge to 0 as sample size n increases when the privacy parameter is a fixed constant. To this end, we consider cases $(n, \epsilon) = \{100 \times 2^i, i = 0, 1, \dots, 8\} \times \{1, 2, 3, 4, \infty\}$.

Scenario II. In the second scenario, we explore the asymptotic behavior of $\hat{E}(\tilde{f}_n)$ with ϵ adaptive to the sample size n . Specifically, we set $\epsilon = 5n^{-\zeta}$ and consider cases $(n, \zeta) = \{100 \times 2^i, i = 0, 1, \dots, 8\} \times \{1/5, 1/4, 1/3, 1/2, 2/3, 1\}$. We also include the worst case $\epsilon = 0$ as a baseline.

Scenario III. In the third scenario, we intend to verify that $\epsilon \asymp n^{-1/2}$ is the dividing line between whether or not the excess risk converges. To this end, we consider three kinds of adaptive ϵ , including $\epsilon \asymp n^{-1/2}$, $\epsilon \asymp \log(n)n^{-1/2}$, and $\epsilon \asymp \log^{-1}(n)n^{-1/2}$. The size of $\mathcal{D}$ is set as $\{100 \times 3^i, i = 0, 1, \dots, 7\}$. For all cases, we report the averaged empirical excess risk in 1,000 replications as well as their 95% confidence intervals.

For Scenario I and Scenario II, we report the averaged empirical excess risk and the classification error in 1,000 replications for each setting in Figure 2 and Figure 3, respectively. From the left panel of Figure 2, we can see that the empirical excess risks and the classification errors with a fixed ϵ converge to 0 regardless of the value of ϵ , showing that the randomized response mechanism with a fixed ϵ fails to prevent the third party from learning the optimal classifier based on $\tilde{\mathcal{D}}$. Moreover, as seen from Figure 3, when $\zeta < 1/2$ the estimated excess risks present a decreasing pattern as the sample size increases, whereas that of the case $\zeta = 1$ deteriorates steadily and the curve finally overlaps with that of the worst case $\epsilon = 0$. It is also interesting to observe that the curve of $\zeta = 1/2$ remains unaffected by the sample size. All these phenomenons are in accordance with the results of Theorem 1.

As can be seen in Figure 4, the curve of the case $\epsilon \asymp n^{-1/2}$ remains unchanged as sample size increases as in Scenario II. This is due to the offset of information gain yielded by increasing the sample size and the

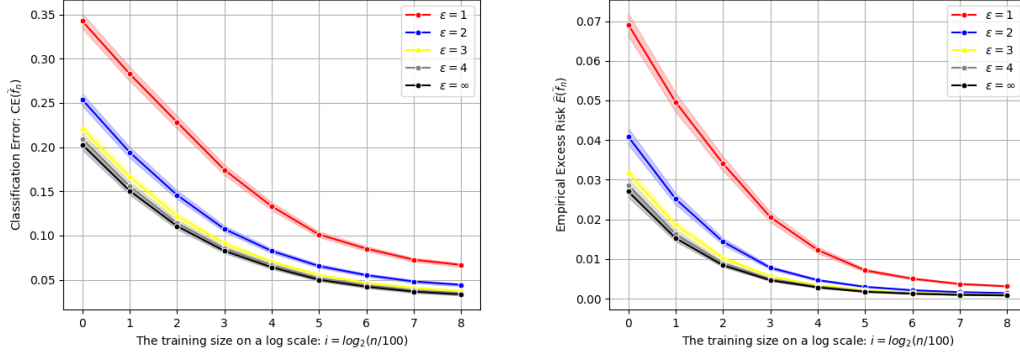


Figure 2: The averaged classification errors (Left) and averaged empirical excess risks (Right) of all settings with $n_{test} = 50,000$ in Scenario I.

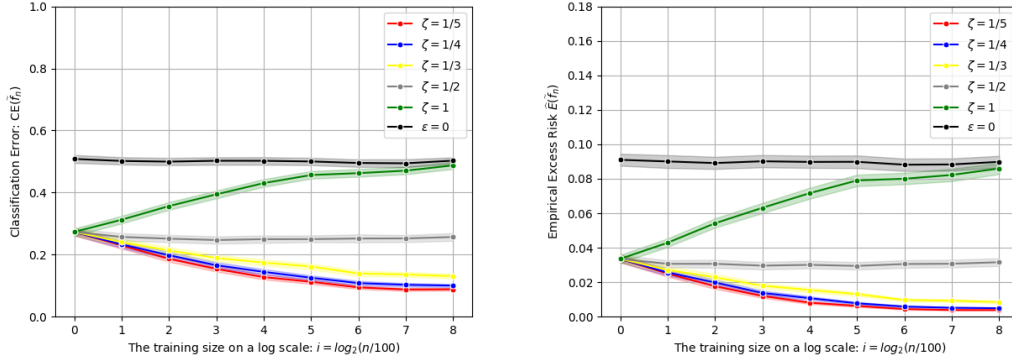


Figure 3: The averaged classification errors (Left) and averaged empirical excess risks (Right) of all settings with $n_{test} = 50,000$ in Scenario II.

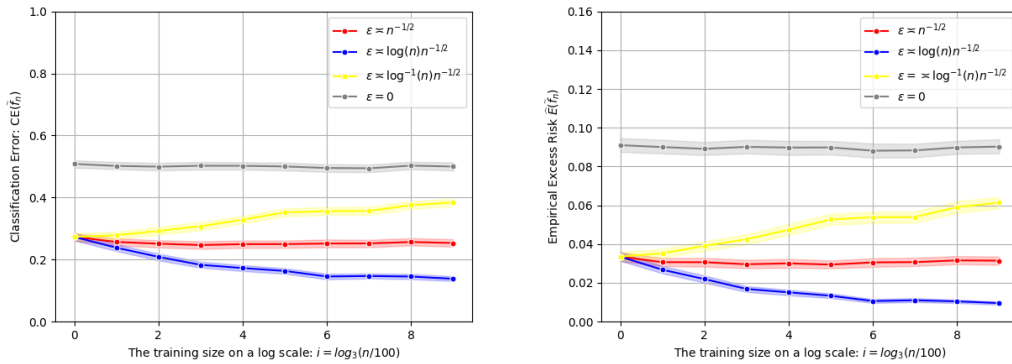


Figure 4: The averaged classification errors (Left) and averaged empirical excess risks (Right) of all settings with $n_{test} = 50,000$ in Scenario III.

information loss in the label-flipping mechanism. As expected, the additional logarithmic term significantly alters the original curve pattern. Specifically, for the case with $\epsilon \asymp \log^{-1}(n)n^{-1/2}$, the performance of $\hat{f}_n$ deteriorates significantly and approaches the worst case with $\epsilon = 0$. On the contrary, the performance of $\hat{f}_n$ in the case $\epsilon \asymp \log(n)n^{-1/2}$ improves significantly with $\hat{E}(\hat{f}_n)$ converging to 0. Therefore, $\epsilon \asymp n^{-1/2}$ appears

to be the dividing line that determines whether the excess risk converges, which completely matches our theoretical results.

5.2 Deep Neural Network Classifier

The simulation in this section aims to verify our theoretical results in Theorem 2, which mainly lies in two aspects. First, we intend to verify the effect of ϵ of Label DP on the optimal structure of deep neural networks for classification problems. Specifically, as stated in Theorem 2, if label noise yielded by the randomized response mechanism increases, a smaller deep neural network should be employed to strike a better balance between the approximation error and the estimation error to achieve a better generalization error. Second, the consistency in estimating the decision boundary is prohibited provided that ϵ is adaptive to the training size n at some specific orders. To these ends, we consider generating training datasets $\tilde{\mathcal{D}} = \{(\mathbf{x}_i, \tilde{y}_i)\}_{i=1}^n$ as follows. First, we set the regression function as $\eta_{nn}(\mathbf{x}_i) = \sum_{j=1}^4 \sin(2\pi x_{ij})/8 + 1/2$ with $x_{ij} \sim \text{Unif}(0, 1)$ for any i, j . Then we generate y_i from $\{1, -1\}$ with probabilities $\eta_{nn}(\mathbf{x}_i)$ and $1 - \eta_{nn}(\mathbf{x}_i)$, respectively. As in the last simulation, we then apply the randomized response mechanism $\mathcal{A}_\theta$ to each y_i to generate $\tilde{y}_i = \mathcal{A}_\theta(y_i)$ with $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$. Then we set the hypothesis space to be the class of L -layer fully connected neural network with equal width and the ReLU activation function, where h denotes the widths in all hidden layers.

The overall training process of the neural network is implemented in Tensorflow (Abadi et al., 2016) with the Adam optimizer and learning rate being 0.001. Additionally, we employ the early-stopping technique to monitor the training error with patience 10 and maintain the parameter with the smallest training error. Let $\tilde{f}_{nn}$ denote the resultant neural network obtained from minimizing (4). We construct the associated plug-in classifier as $\tilde{s}_{nn} = \text{sign}(\tilde{f}_{nn} - 1/2)$ and evaluate its performance by the empirical excess risk and the classification error as in Section 5.1.

Scenario I. In the first scenario, we consider privacy guarantees $\epsilon \in \{1, 2, \infty\}$ and neural network structures with $L = 2$ and $h \in \{8, 12, 16, 20, 24\}$ with . We report the averaged empirical excess risks and the classification errors of all cases in 100 replications as well as their 95% confidence intervals in Figure 5. Clearly, if $\epsilon = 1$, the optimal neural network structure is $h = 8$, whereas those of the cases $\epsilon = 2$ and $\epsilon = \infty$ are $h = 20$ and $h = 24$, respectively. Such results show that a smaller neural network is preferred when stronger privacy protection of labels (smaller ϵ) is considered, which coincides with our theoretical results in Theorem 2 that the optimal neural network structure to achieve the fastest convergence rate of excess risk should diminish as ϵ decreases. Moreover, as the training sample size n increases from 2,000 to 4,000, the optimal structure of the neural network enlarges for the cases $\epsilon = 1$ and $\epsilon = 2$ due to their new tradeoffs between the estimation and approximation errors.

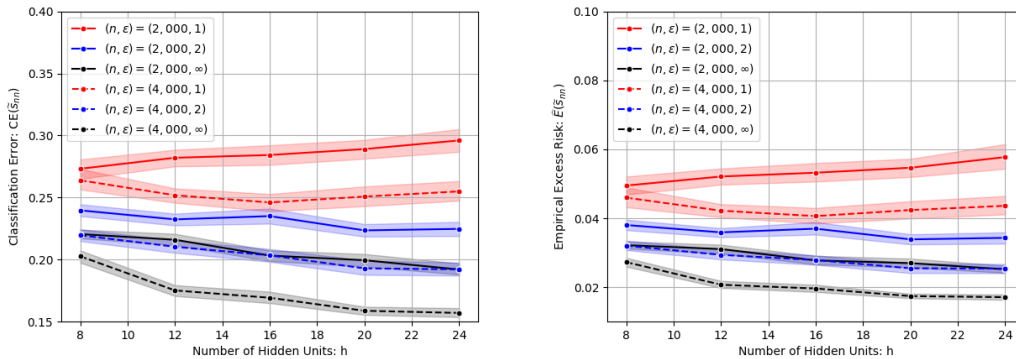


Figure 5: The averaged classification errors (Left) and averaged empirical excess risks (Right) of all cases with $n_{test} = 50,000$.

Scenario II. In the second scenario, we fix the neural network structure as $(L, h) = (2, 16)$ and consider training sample sizes $n \in \{2^i \times 10^3; i = 0, 1, 2, 3, 4\}$. To verify our theoretical results, we compare two privacy

schemes $\epsilon \asymp n^{-1/2}$ and $\epsilon \asymp \log(n)n^{-1/2}$. We also include the case with invariant privacy scheme $\epsilon = 4$ as a baseline. For comparing their difference in trend patterns, the multiplicative constants of two adaptive privacy schemes are chosen such that they have the same starting point as $\epsilon = 4$ when $n = 1,000$. We report the averaged empirical excess risks, the classification errors of all cases in 100 replications, and their 95% confidence intervals in Figure 6. Clearly, the performances of the cases $\epsilon = 4$ and $\epsilon \asymp \log(n)n^{-1/2}$ improve as n increases. However, the performance of case $\epsilon \asymp n^{-1/2}$ presents a different pattern in generalization performance as n increases. Most notably, as n increases from 8,000 to 16,000, its performance deteriorates while the other two cases still observe significant improvements, showing that the additional logarithmic term plays a deterministic role in the convergence of excess risk, which perfectly aligns with our theoretical findings in Theorem 2.

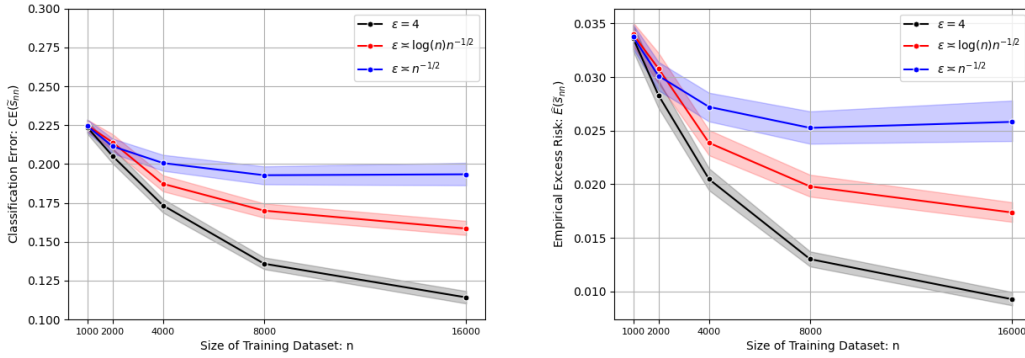


Figure 6: The averaged classification errors (Left) and averaged empirical excess risks (Right) of all cases with $n_{test} = 50,000$.

6 Real Applications - Mnist Dataset

This experiment considers similar settings of the privacy parameter ϵ as Section 5.2 in order to verify our theoretical findings of DNN on the MNIST dataset (LeCun, 1998). The MNIST dataset consists of 60,000 training images and 10,000 testing images, and each sample is a 28×28 grey-scale pixel image of one of the 10 digits. In this experiment, we consider a binary classification problem by only including samples of digits 2 and 3 for training and testing due to their similarity in appearance. The resultant training dataset contains 12,089 training samples and 2,042 testing samples.

For the hyperparameters setting, we consider the neural network with three convolution layers with ReLU activation and two fully-connected layers. For the three convolution layers, we set their kernel sizes and numbers of channels as 4×4 and 4, respectively. Additionally, each convolution layer is followed by a max pooling layer with size 2×2 . The first fully-connected layer has 10 hidden units with ReLU activation, and the last layer outputs the probability of an image being digit 2. As in Section 5.2, the neural network is trained with the Adam optimizer and learning rate being 0.001, and the early-stopping technique to monitor the training error with patience 10 and maintain the parameter with the smallest training error. We evaluate the trained model by the testing error on 2,042 testing samples. We consider training sample size $n \in \{2 \times i \times 10^3; i = 1, 2, 3, 4, 5\}$. We mainly consider two scenarios, including the fixed privacy guarantee with $\epsilon \in \{1, 2, \infty\}$ and the adaptive privacy schemes with $\epsilon \asymp \log(n)\sqrt{n^{-1}}$, $\epsilon \asymp n^{-1/2}$, and $\epsilon \asymp \log^{-1}(n)n^{-1/2}$. The averaged testing error of each case in 50 replications is reported in Figure 7.

Figure 7 presents similar results as in Section 5.2. First, when ϵ is fixed, differentially private classifiers improve in generalization performance as the training sample size increases and attain competitive performance as the non-private classifier ($\epsilon = \infty$) when the training sample size is large enough. This result accords with our theoretical findings in Theorem 1 that fixed privacy guarantee in the label DP slows down the convergence to the optimal classifier with a multiplicative constant. In stark contrast, as shown in the right plot of Figure 7, the convergence to the optimal classifier is prevented if $\epsilon \asymp n^{-1/2}$, as boosting the

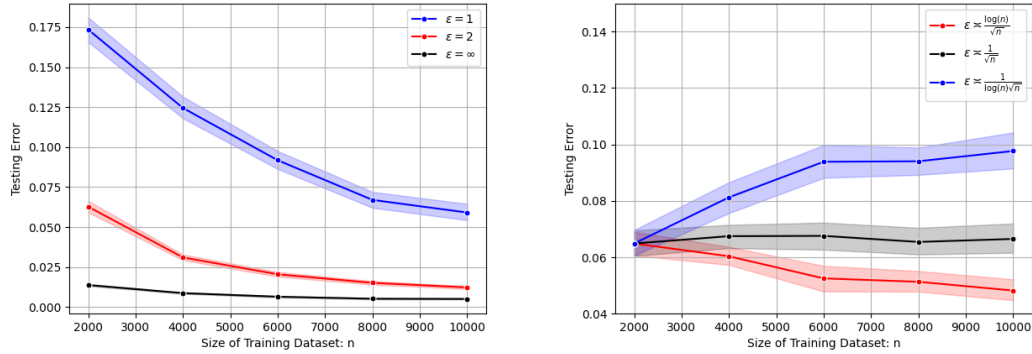


Figure 7: The averaged testing errors with fixed ϵ (Left) and adaptive ϵ (Right) under different training sample sizes in MNIST dataset

training size from 2,000 to 10,000 fails to significantly improve the testing error. This again aligns with our Theorem 1

References

- Martín Abadi, Paul Barham, Jianmin Chen, Zhifeng Chen, Andy Davis, Jeffrey Dean, Matthieu Devin, Sanjay Ghemawat, Geoffrey Irving, Michael Isard, et al. {TensorFlow}: a system for {Large-Scale} machine learning. In *12th USENIX Symposium on Operating Systems Design and Implementation (OSDI 16)*, pp. 265–283, 2016.
- Mário S Alvim, Miguel E Andrés, Konstantinos Chatzikokolakis, Pierpaolo Degano, and Catuscia Palamidessi. Differential privacy: on the trade-off between utility and information leakage. In *International Workshop on Formal Aspects in Security and Trust*, pp. 39–54. Springer, 2012.
- Pathum Chamikara Mahawaga Arachchige, Peter Bertok, Ibrahim Khalil, Dongxi Liu, Seyit Camtepe, and Mohammed Atiquzzaman. Local differential privacy for deep learning. *IEEE Internet of Things Journal*, 7(7):5827–5842, 2019.
- Sanjeev Arora, László Babai, Jacques Stern, and Z Sweedyk. The hardness of approximate optima in lattices, codes, and systems of linear equations. *Journal of Computer and System Sciences*, 54(2):317–331, 1997.
- Jean-Yves Audibert and Alexandre B Tsybakov. Fast learning rates for plug-in classifiers. *The Annals of statistics*, 35(2):608–633, 2007.
- Peter L Bartlett and Shahar Mendelson. Rademacher and gaussian complexities: Risk bounds and structural results. *Journal of Machine Learning Research*, 3(Nov):463–482, 2002.
- Peter L Bartlett and Marten H Wegkamp. Classification with a reject option using a hinge loss. *Journal of Machine Learning Research*, 9(8), 2008.
- Peter L Bartlett, Michael I Jordan, and Jon D McAuliffe. Convexity, classification, and risk bounds. *Journal of the American Statistical Association*, 101(473):138–156, 2006.
- Peter L Bartlett, Nick Harvey, Christopher Liaw, and Abbas Mehrabian. Nearly-tight VC-dimension and pseudodimension bounds for piecewise linear neural networks. *The Journal of Machine Learning Research*, 20(1):2285–2301, 2019.
- Raef Bassily, Mehryar Mohri, and Ananda Theertha Suresh. Open problem: Better differentially private learning algorithms with margin guarantees. In Po-Ling Loh and Maxim Raginsky (eds.), *Proceedings of Thirty Fifth Conference on Learning Theory*, volume 178 of *Proceedings of Machine Learning Research*, 2023.

- ing Research*, pp. 5638–5643. PMLR, 02–05 Jul 2022. URL <https://proceedings.mlr.press/v178/open-problem-bassily22a.html>.
- Graeme Blair, Kosuke Imai, and Yang-Yang Zhou. Design and analysis of the randomized response technique. *Journal of the American Statistical Association*, 110(511):1304–1319, 2015.
- Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K Warmuth. Learnability and the vapnik-chervonenkis dimension. *Journal of the ACM (JACM)*, 36(4):929–965, 1989.
- Olivier Bousquet, Stéphane Boucheron, and Gábor Lugosi. Introduction to statistical learning theory. In *Summer School on Machine Learning*, pp. 169–207. Springer, 2003.
- Robert Istvan Busa-Fekete, Umar Syed, Sergei Vassilvitskii, et al. Population level privacy leakage in binary classification with label noise. In *NeurIPS 2021 Workshop Privacy in Machine Learning*, 2021.
- Timothy I Cannings, Yingying Fan, and Richard J Samworth. Classification with imperfect training labels. *Biometrika*, 107(2):311–330, 04 2020. doi: 10.1093/biomet/asaa011. URL <https://doi.org/10.1093/biomet/asaa011>.
- Olivier Chapelle, Eren Manavoglu, and Romer Rosales. Simple and scalable response prediction for display advertising. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 5(4):1–34, 2014.
- Kamalika Chaudhuri and Daniel Hsu. Sample complexity bounds for differentially private learning. In *Proceedings of the 24th Annual Conference on Learning Theory*, pp. 155–186. JMLR Workshop and Conference Proceedings, 2011.
- Corinna Cortes and Vladimir Vapnik. Support-vector networks. *Machine Learning*, 20(3):273–297, 1995.
- Teddy Cunningham, Konstantin Klemmer, Hongkai Wen, and Hakan Ferhatosmanoglu. Geopointgan: Synthetic spatial data with local label differential privacy. *arXiv preprint arXiv:2205.08886*, 2022.
- Jinshuo Dong, Aaron Roth, and Weijie J Su. Gaussian differential privacy. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 84(1):3–54, 2022.
- Cynthia Dwork. Differential privacy: A survey of results. In *International Conference on Theory and Applications of Models of Computation*, pp. 1–19. Springer, 2008.
- Úlfar Erlingsson, Vasyl Pihur, and Aleksandra Korolova. Rappor: Randomized aggregatable privacy-preserving ordinal response. In *Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security*, pp. 1054–1067, 2014.
- Úlfar Erlingsson, Vitaly Feldman, Ilya Mironov, Ananth Raghunathan, Kunal Talwar, and Abhradeep Thakurta. Amplification by shuffling: From local to central differential privacy via anonymity. In *Proceedings of the Thirtieth Annual ACM-SIAM Symposium on Discrete Algorithms*, pp. 2468–2479. SIAM, 2019.
- Hossein Esfandiari, Vahab Mirrokni, Umar Syed, and Sergei Vassilvitskii. Label differential privacy via clustering. In *International Conference on Artificial Intelligence and Statistics*, pp. 7055–7075. PMLR, 2022.
- Badih Ghazi, Noah Golowich, Ravi Kumar, Pasin Manurangsi, and Chiyuan Zhang. Deep learning with label differential privacy. *Advances in Neural Information Processing Systems*, 34:27131–27145, 2021.
- Badih Ghazi, Pritish Kamath, Ravi Kumar, Ethan Leeman, Pasin Manurangsi, Avinash Varadarajan, and Chiyuan Zhang. Regression with label differential privacy. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=h900wsmL-cT>.
- Antonious Girgis, Deepesh Data, Suhas Diggavi, Peter Kairouz, and Ananda Theertha Suresh. Shuffled model of differential privacy in federated learning. In *International Conference on Artificial Intelligence and Statistics*, pp. 2521–2529. PMLR, 2021.

- Xiao Guo, Xiang Li, Xiangyu Chang, and Shujie Ma. Privacy-preserving community detection for locally distributed multiple networks. *arXiv preprint arXiv:2306.15709*, 2023.
- Jonathan Hehir, Aleksandra Slavkovic, and Xiaoyue Niu. Consistent spectral clustering of network block models under local differential privacy. *Journal of Privacy and Confidentiality*, 12(2), 2022.
- Matthew Joseph, Aaron Roth, Jonathan Ullman, and Bo Waggoner. Local differential privacy for evolving data. *Advances in Neural Information Processing Systems*, 31, 2018.
- Peter Kairouz, Sewoong Oh, and Pramod Viswanath. Extremal mechanisms for local differential privacy. *The Journal of Machine Learning Research*, 17(1):492–542, 2016.
- Vishesh Karwa, Pavel N Krivitsky, and Aleksandra B Slavković. Sharing social network data: differentially private estimation of exponential family random-graph models. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 66(3):481–500, 2017.
- Yongdai Kim, Ilsang Ohn, and Dongha Kim. Fast convergence rates of deep neural networks for classification. *Neural Networks*, 138:179–197, 2021.
- Vladimir Koltchinskii. *Oracle inequalities in empirical risk minimization and sparse recovery problems: École D’Été de Probabilités de Saint-Flour XXXVIII-2008*, volume 2033. Springer Science & Business Media, 2011.
- Guillaume Lécué. Optimal rates of aggregation in classification under low noise assumption. *Bernoulli*, 13(4):1000–1022, 2007.
- Yann LeCun. The mnist database of handwritten digits. <http://yann.lecun.com/exdb/mnist/>, 1998.
- Wee Sun Lee, Peter L Bartlett, and Robert C Williamson. Lower bounds on the VC-dimension of smoothly parametrized function classes. In *Proceedings of the Seventh Annual Conference on Computational Learning Theory*, pp. 362–367, 1994.
- Yunwen Lei, Lixin Ding, and Yingzhou Bi. Local rademacher complexity bounds based on covering numbers. *Neurocomputing*, 218:320–330, 2016.
- Ximing Li, Chendi Wang, and Guang Cheng. Statistical theory of differentially private marginal-based data synthesis algorithms. In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net, 2023a. URL <https://openreview.net/pdf?id=hxUwnEGxW87>.
- Zhechen Li, Ao Liu, Lirong Xia, Yongzhi Cao, and Hanpin Wang. Differentially private condorcet voting. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pp. 5755–5763, 2023b.
- Yi Lin. A note on margin-based loss functions in classification. *Statistics & Probability Letters*, 68(1):73–82, 2004.
- Mani Malek Esmaeili, Ilya Mironov, Karthik Prasad, Igor Shilov, and Florian Tramer. Antipodes of label differential privacy: Pate and alibi. *Advances in Neural Information Processing Systems*, 34:6934–6945, 2021.
- H Brendan McMahan, Gary Holt, David Sculley, Michael Young, Dietmar Ebner, Julian Grady, Lan Nie, Todd Phillips, Eugene Davydov, Daniel Golovin, et al. Ad click prediction: a view from the trenches. In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 1222–1230, 2013.
- Frank McSherry and Ilya Mironov. Differentially private recommender systems: Building privacy into the netflix prize contenders. In *Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 627–636, 2009.

- Ryumei Nakada and Masaaki Imaizumi. Adaptive approximation and generalization of deep neural network with intrinsic dimensionality. *Journal of Machine Learning Research*, 21(174):1–38, 2020.
- Tapan K Nayak and Samson A Adeshiyan. A unified framework for analysis and comparison of randomized response surveys of binary characteristics. *Journal of Statistical Planning and Inference*, 139(8):2757–2766, 2009.
- Igal Sason and Sergio Verdú. f -divergence inequalities. *IEEE Transactions on Information Theory*, 62(11):5973–6006, 2016.
- Johannes Schmidt-Hieber. Nonparametric regression using deep neural networks with relu activation function. *The Annals of Statistics*, 48(4):1875–1897, 2020.
- Shai Shalev-Shwartz and Shai Ben-David. *Understanding machine learning: From theory to algorithms*. Cambridge university press, 2014.
- Xiaotong Shen and Wing Hung Wong. Convergence rate of sieve estimates. *The Annals of Statistics*, pp. 580–615, 1994.
- Xiaotong Shen, George C Tseng, Xuegong Zhang, and Wing Hung Wong. On ψ -learning. *Journal of the American Statistical Association*, 98(463):724–734, 2003.
- Steve Smale and Ding-Xuan Zhou. Estimating the approximation error in learning theory. *Analysis and Applications*, 1(01):17–41, 2003.
- Jun Tang, Aleksandra Korolova, Xiaolong Bai, Xueqiang Wang, and Xiaofeng Wang. Privacy loss in apple’s implementation of differential privacy on macos 10.12. *arXiv preprint arXiv:1709.02753*, 2017.
- Alexander B Tsybakov. Optimal aggregation of classifiers in statistical learning. *The Annals of Statistics*, 32(1):135–166, 2004.
- Brendan van Rooyen and Robert C. Williamson. A theory of learning with corrupted labels. *Journal of Machine Learning Research*, 18(228):1–50, 2018. URL <http://jmlr.org/papers/v18/16-315.html>.
- Vladimir N Vapnik and A Ya Chervonenkis. On the uniform convergence of relative frequencies of events to their probabilities. In *Measures of Complexity*, pp. 11–30. Springer, 2015.
- Chendi Wang, Buxin Su, Jiayuan Ye, Reza Shokri, and Weijie Su. Unified enhancement of privacy bounds for mixture mechanisms via f -differential privacy. In *NeurIPS*, 2023.
- Di Wang and Jinhui Xu. On sparse linear regression in the local differential privacy model. In *International Conference on Machine Learning*, pp. 6628–6637. PMLR, 2019.
- Junhui Wang, Xiaotong Shen, Yiwen Sun, and Annie Qu. Classification with unstructured predictors and an application to sentiment analysis. *Journal of the American Statistical Association*, 111(515):1242–1253, 2016a.
- Tianhao Wang, Jeremiah Blocki, Ninghui Li, and Somesh Jha. Locally differentially private protocols for frequency estimation. In *26th USENIX Security Symposium (USENIX Security 17)*, pp. 729–745, 2017.
- Yue Wang, Xintao Wu, and Donghui Hu. Using randomized response for differential privacy preserving data collection. In *EDBT/ICDT Workshops*, volume 1558, pp. 0090–6778, 2016b.
- Stanley L Warner. Randomized response: A survey technique for eliminating evasive answer bias. *Journal of the American Statistical Association*, 60(309):63–69, 1965.
- Jon Wellner et al. *Weak convergence and empirical processes: with applications to statistics*. Springer Science & Business Media, 2013.
- Yu Xin and Tommi Jaakkola. Controlling privacy in recommender systems. *Advances in Neural Information Processing Systems*, 27, 2014.

- Shirong Xu, Ben Dai, and Junhui Wang. Sentiment analysis with covariate-assisted word embeddings. *Electronic Journal of Statistics*, 15(1):3015–3039, 2021.
- Ziqi Yan, Qiong Wu, Meng Ren, Jiqiang Liu, Shaowu Liu, and Shuo Qiu. Locally private jaccard similarity estimation. *Concurrency and Computation: Practice and Experience*, 31(24):e4889, 2019.
- Dmitry Yarotsky. Error bounds for approximations with deep relu networks. *Neural Networks*, 94:103–114, 2017.
- Bin Yu. Assouad, Fano, and Le Cam. In *Festschrift for Lucien Le Cam: Research Papers in Probability and Statistics*, pp. 423–435. Springer, 1997.
- Sen Yuan, Milan Shen, Ilya Mironov, and Anderson CA Nascimento. Practical, label private deep learning training based on secure multiparty computation and differential privacy. *Cryptology ePrint Archive*, 2021.
- Tong Zhang. Statistical behavior and consistency of classification methods based on convex risk minimization. *The Annals of Statistics*, 32(1):56–85, 2004.
- Ding-Xuan Zhou. The covering number in learning theory. *Journal of Complexity*, 18(3):739–767, 2002.

A Appendix

Lemma 1 (Restatement of Lemma 1). *Suppose the randomized response mechanism $\tilde{Y} = \mathcal{A}_\theta(Y)$ satisfies ϵ -Label DP with $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$, then for any $\mathbf{x} \in \mathcal{X}$ it holds that*

$$\mathcal{L}(\epsilon, \mathbf{x}) \leq D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) \leq \mathcal{U}(\epsilon, \mathbf{x}),$$

where $D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}})$ denotes the Kullback-Leibler divergence (KL) divergence between $\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}$ and $\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}$, $\mathcal{L}(\epsilon, \mathbf{x}) = 2(2\eta(\mathbf{x}) - 1)^2(1 + \exp(\epsilon))^{-2}$, and $\mathcal{U}(\epsilon, \mathbf{x}) = \min\{U_1(\epsilon, \mathbf{x}), U_2(\epsilon, \mathbf{x})\}$ with $U_1(\epsilon, \mathbf{x}) = (2\eta(\mathbf{x}) - 1)^2\eta^{-1}(\mathbf{x})(1 - \eta(\mathbf{x}))^{-1}(1 + \exp(\epsilon))^{-2}$ and $U_2(\epsilon, \mathbf{x}) = (2\eta(\mathbf{x}) - 1)^2 \exp(-\epsilon)$.

Proof of Lemma 1: We first prove the lower bound of $D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}})$, which is mainly based on the Pinsker’s inequality (Sason & Verdú, 2016). Without loss of generality, we assume that $\eta(\mathbf{x}) > 1/2$. Define

$$S(p, q) = p \log \frac{p}{q} + (1 - p) \log \frac{1 - p}{1 - q} - 2(p - q)^2,$$

where $p, q \in [0, 1]$. Let $S(p, q)$ take the partial derivative with respect to q , we get

$$\frac{\partial S(p, q)}{\partial q} = -\frac{p}{q} + \frac{1 - p}{1 - q} + 4(p - q) = -(p - q)\left(\frac{1}{q(1 - q)} - 4\right) \leq 0, \text{ for } q \leq p,$$

where the last inequality follows from that fact that $q(1 - q) \leq 1/4$ for $q \in [0, 1]$ and the equality holds when $p = q$. Suppose that $\mathcal{A}_\theta$ satisfies ϵ -Label DP, which is equivalent to set $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$. Therefore, it holds that

$$D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}} | \mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) \geq 2(\eta(\mathbf{x}) - \tilde{\eta}(\mathbf{x}))^2 = 2(2\eta(\mathbf{x}) - 1)^2(1 + \exp(\epsilon))^{-2} \triangleq \mathcal{L}(\epsilon, \mathbf{x}).$$

Next, we proceed to prove the upper bound. For any pair of distribution $\mathbb{P}$ and $\mathbb{Q}$, we have

$$D_{KL}(\mathbb{P}|\mathbb{Q}) = D_{KL}(\mathbb{P}|\mathbb{Q}) + 1 - 1 \leq \exp(D_{KL}(\mathbb{P}|\mathbb{Q})) - 1.$$

Recall that Y and $\tilde{Y}$ take values in $\{-1, 1\}$, we have

$$\begin{aligned} D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}|\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) &\leq \exp(D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}|\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}})) - 1 \\ &\leq \eta(\mathbf{x})\frac{\eta(\mathbf{x})}{\tilde{\eta}(\mathbf{x})} + (1-\eta(\mathbf{x}))\frac{1-\eta(\mathbf{x})}{1-\tilde{\eta}(\mathbf{x})} - 1 = \eta(\mathbf{x})\left(\frac{\eta(\mathbf{x})}{\tilde{\eta}(\mathbf{x})} - 1\right) + (1-\eta(\mathbf{x}))\left(\frac{1-\eta(\mathbf{x})}{1-\tilde{\eta}(\mathbf{x})} - 1\right) \\ &= \eta(\mathbf{x})\left(\frac{\eta(\mathbf{x}) - \tilde{\eta}(\mathbf{x})}{\tilde{\eta}(\mathbf{x})}\right) + (1-\eta(\mathbf{x}))\left(\frac{\tilde{\eta}(\mathbf{x}) - \eta(\mathbf{x})}{1-\tilde{\eta}(\mathbf{x})}\right) = \frac{(\eta(\mathbf{x}) - \tilde{\eta}(\mathbf{x}))^2}{\tilde{\eta}(\mathbf{x})(1-\tilde{\eta}(\mathbf{x}))}. \end{aligned}$$

Note that

$$\begin{aligned} \tilde{\eta}(\mathbf{x})(1-\tilde{\eta}(\mathbf{x})) &= (\theta\eta(\mathbf{x}) + (1-\theta)(1-\eta(\mathbf{x})))((1-\theta)\eta(\mathbf{x}) + \theta(1-\eta(\mathbf{x}))) \\ &= \theta(1-\theta)\eta^2(\mathbf{x}) + \eta(\mathbf{x})(1-\eta(\mathbf{x}))(\theta^2 + (1-\theta)^2) + \theta(1-\theta)(1-\eta(\mathbf{x}))^2 \\ &\geq \max\{\eta(\mathbf{x})(1-\eta(\mathbf{x})), \theta(1-\theta)\}. \end{aligned}$$

It then follows that

$$D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}|\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) \leq \frac{(\eta(\mathbf{x}) - \tilde{\eta}(\mathbf{x}))^2}{\max\{\eta(\mathbf{x})(1-\eta(\mathbf{x})), \theta(1-\theta)\}}. \quad (6)$$

Plugging $\theta = \exp(\epsilon)/(1 + \exp(\epsilon))$ into (6) yields that

$$D_{KL}(\mathbb{P}_{Y|\mathbf{X}=\mathbf{x}}|\mathbb{P}_{\tilde{Y}|\mathbf{X}=\mathbf{x}}) \leq \frac{(2\eta(\mathbf{x}) - 1)^2}{\max\{\eta(\mathbf{x})(1-\eta(\mathbf{x}))(1 + \exp(\epsilon))^2, \exp(\epsilon)\}} \triangleq \mathcal{U}(\theta, \mathbf{x}).$$

This completes the proof. $\square$

Lemma 2 (Restatement of Lemma 2). *If $\tilde{Y} = \mathcal{A}_\theta(Y)$ with $\theta > 1/2$, then $f^*(\mathbf{x}) = \tilde{f}^*(\mathbf{x})$ for any $\mathbf{x} \in \mathcal{X}$ and $\tilde{D}(f, \tilde{f}^*) = (2\theta - 1)D(f, f^*)$ for any f .*

Proof of Lemma 2: By the definition of $\tilde{\eta}(\mathbf{x})$, we can obtain

$$2\tilde{\eta}(\mathbf{x}) - 1 = (2\theta - 1)(2\eta(\mathbf{x}) - 1). \quad (7)$$

Clearly, $2\eta(\mathbf{x}) > 1$ indicates that $2\tilde{\eta}(\mathbf{x}) > 1$ provided that $\theta > 1/2$. By the fact that $f^*(\mathbf{x}) = \text{sign}(\eta(\mathbf{x}) - 1/2)$, it holds that $f^*(\mathbf{x}) = \tilde{f}^*(\mathbf{x})$ for any $\mathbf{x} \in \mathcal{X}$. Recall that the excess risk in terms of 0-1 loss can be written as

$$R(f) - R(f^*) = \mathbb{E}[|\text{sign}(f(\mathbf{X})) \neq \text{sign}(f^*(\mathbf{X}))||2\eta(\mathbf{X}) - 1|].$$

Combined with (7), it holds that

$$\begin{aligned} \tilde{R}(f) - \tilde{R}(\tilde{f}^*) &= \mathbb{E}[|\text{sign}(f(\mathbf{X})) \neq \text{sign}(\tilde{f}^*(\mathbf{X}))||2\tilde{\eta}(\mathbf{X}) - 1|] \\ &= (2\theta - 1)\mathbb{E}[|\text{sign}(f(\mathbf{X})) \neq \text{sign}(f^*(\mathbf{X}))||2\eta(\mathbf{X}) - 1|] \\ &= (2\theta - 1)(R(f) - R(f^*)). \end{aligned}$$

This completes the proof. $\square$

Lemma 3 (Restatement of Lemma 3). *Denote that $\kappa_\epsilon = (e^\epsilon - 1)/(e^\epsilon + 1)$. Under Assumption 1, for any $\theta \in (1/2, 1]$, it holds that*

$$(1) \quad \mathbb{P}(|2\tilde{\eta}(\mathbf{X}) - 1| \leq t) \leq c\kappa_\epsilon^{-\gamma}t^\gamma, \text{ for any } t \in [0, 1),$$

$$(2) \quad \tilde{R}_\phi(f) - \tilde{R}_\phi(f^*) \geq \kappa_\epsilon(4c)^{-1/\gamma} \left(\mathbb{E}[|f(\mathbf{X}) - f^*(\mathbf{X})|] \right)^{\frac{\gamma+1}{\gamma}} \text{ for any } f \in \mathcal{F},$$

Proof of Lemma 3: By the assumption that $\|f\|_\infty \leq 1$ and the fact that 0-1 loss is upper bounded by hinge loss, we obtain $R_\phi(f) - R_\phi(f^*) \geq R(f) - R(f^*)$. Since ϕ is set as hinge loss, it holds that

$$\begin{aligned} R_\phi(f) - R_\phi(f^*) &= \mathbb{E}_{\mathbf{X}} [|f(\mathbf{X}) - f^*(\mathbf{X})| |1 - 2\eta(\mathbf{X})|] \\ &\geq t \mathbb{E}_{\mathbf{X}} [|f(\mathbf{X}) - f^*(\mathbf{X})| I(|1 - 2\eta(\mathbf{X})| > t)]. \end{aligned}$$

By the fact that $I(|1 - 2\eta(\mathbf{X})| > t) = 1 - I(|1 - 2\eta(\mathbf{X})| \leq t)$, it then follows that

$$\begin{aligned} R_\phi(f) - R_\phi(f^*) &\geq t \left(\mathbb{E}_{\mathbf{X}} [|f(\mathbf{X}) - f^*(\mathbf{X})|] - 2\mathbb{P}(|1 - 2\eta(\mathbf{X})| \leq t) \right) \\ &\geq t \left(\mathbb{E}_{\mathbf{X}} [|f(\mathbf{X}) - f^*(\mathbf{X})|] - 2ct^\gamma \right) \\ &= t \mathbb{E}_{\mathbf{X}} [|f(\mathbf{X}) - f^*(\mathbf{X})|] - 2ct^{\gamma+1} \end{aligned}$$

where the last inequality follows from Assumption 1. Choosing t such that $t \mathbb{E} [|f(\mathbf{X}) - f^*(\mathbf{X})|] = 4ct^{\gamma+1}$, we get

$$R_\phi(f) - R_\phi(f^*) \geq (4c)^{-1/\gamma} \left(\mathbb{E} [|f(\mathbf{X}) - f^*(\mathbf{X})|] \right)^{\frac{\gamma+1}{\gamma}}.$$

Next, we proceed to establish the relation between $\tilde{R}_\phi(f) - \tilde{R}_\phi(f^*)$ and $\mathbb{E} [|f(\mathbf{X}) - f^*(\mathbf{X})|]$. For each $\mathbf{x} \in \mathcal{X}$, the conditional risk is given as

$$\begin{aligned} \mathbb{E}_{\tilde{Y}} \left(\phi(f(\mathbf{X})\tilde{Y}) | \mathbf{X} = \mathbf{x} \right) &= \tilde{\eta}(\mathbf{x})\phi(f(\mathbf{x})) + (1 - \tilde{\eta}(\mathbf{x}))\phi(-f(\mathbf{x})) \\ &= [\theta\eta(\mathbf{x}) + (1 - \theta)(1 - \eta(\mathbf{x}))]\phi(f(\mathbf{x})) + [\theta(1 - \eta(\mathbf{x})) + (1 - \theta)\eta(\mathbf{x})]\phi(-f(\mathbf{x})) \\ &= \left(\theta\eta(\mathbf{x})\phi(f(\mathbf{x})) + \theta(1 - \eta(\mathbf{x}))\phi(-f(\mathbf{x})) \right) + \left((1 - \theta)(1 - \eta(\mathbf{x}))\phi(f(\mathbf{x})) + (1 - \theta)\eta(\mathbf{x})\phi(-f(\mathbf{x})) \right). \end{aligned}$$

Taking the expectation with respect to $\mathbf{X}$ yields that

$$\tilde{R}_\phi(f) = \mathbb{E}_{(\mathbf{X}, \tilde{Y})} \left(\phi(f(\mathbf{X})\tilde{Y}) \right) = \theta R_\phi(f) + (1 - \theta) R_\phi(-f) \quad (8)$$

Notice that ϕ is hinge loss, hence

$$R_\phi(f) - R_\phi(f^*) = R_\phi(-f^*) - R_\phi(-f),$$

for any f with $\|f\|_\infty \leq 1$. It follows that

$$\begin{aligned} \tilde{R}_\phi(f) - \tilde{R}_\phi(f^*) &= \theta(R_\phi(f) - R_\phi(f^*)) + (1 - \theta)(R_\phi(-f) - R_\phi(-f^*)) \\ &= (2\theta - 1)(R_\phi(f) - R_\phi(f^*)) \\ &\geq (2\theta - 1)(4c)^{-1/\gamma} \left(\mathbb{E} [|f(\mathbf{X}) - f^*(\mathbf{X})|] \right)^{\frac{\gamma+1}{\gamma}}. \end{aligned}$$

This completes the proof. $\square$

Lemma 4 (Restatement of Lemma 4). *Denote that $\mathcal{M} = \max\{EFNE(f) - EFNE(f^*), EFPE(f) - EFPE(f^*)\}$. Under Assumption 1, for any margin classifier f , it holds that*

$$R(f) - R(f^*) \leq \mathcal{M} \leq \frac{1}{2 \min\{\mathbb{P}(Y = 1), \mathbb{P}(Y = -1)\}} \left(2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*) \right),$$

where c and γ are as defined in Assumption 1. Particularly, if Assumption 1 holds with $\gamma = \infty$, for any margin classifier f , (3) becomes

$$R(f) - R(f^*) \leq \mathcal{M} \leq \frac{3}{2 \min\{\mathbb{P}(Y = 1), \mathbb{P}(Y = -1)\}} (R(f) - R(f^*)).$$

Proof of Lemma 4: We first prove the left-hand side of (3). By the relation among $R(f)$, $\text{EFNE}(f)$, and $\text{EFPE}(f)$, one has

$$\begin{aligned} R(f) - R(f^*) &= (\text{EFNE}(f) - \text{EFNE}(f^*))\mathbb{P}(Y = 1) + (\text{EFPE}(f) - \text{EFPE}(f^*))\mathbb{P}(Y = -1) \\ &\leq \max\{\text{EFNE}(f) - \text{EFNE}(f^*), \text{EFPE}(f) - \text{EFPE}(f^*)\}(\mathbb{P}(Y = 1) + \mathbb{P}(Y = -1)) \\ &= \max\{\text{EFNE}(f) - \text{EFNE}(f^*), \text{EFPE}(f) - \text{EFPE}(f^*)\}. \end{aligned}$$

Next, we prove the right-hand side of (3). We first define

$$S_1(f) = \{\mathbf{x} \in \mathcal{X} : \text{sign}(f(\mathbf{x})) = 1\} \text{ and } S_{-1}(f) = \{\mathbf{x} \in \mathcal{X} : \text{sign}(f(\mathbf{x})) = -1\}.$$

Clearly, it can be verified that $S_1(f^*) = \{\mathbf{x} : \eta(\mathbf{x}) > 1/2\}$ and $S_{-1}(f^*) = \{\mathbf{x} : \eta(\mathbf{x}) < 1/2\}$. With these, we further get

$$\begin{aligned} \text{EFNE}(f) - \text{EFNE}(f^*) &= \frac{1}{\mathbb{P}(Y = 1)} \left(\int_{S_{-1}(f)} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{S_{-1}(f^*)} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right) \\ &= \frac{1}{\mathbb{P}(Y = 1)} \left(\int_{S_{-1}(f) \setminus \Delta S_{-1}} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{S_{-1}(f^*) \setminus \Delta S_{-1}} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right), \\ \text{EFPE}(f) - \text{EFPE}(f^*) &= \frac{1}{\mathbb{P}(Y = -1)} \left(\int_{S_1(f)} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{S_1(f^*)} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right) \\ &= \frac{1}{\mathbb{P}(Y = -1)} \left(\int_{S_1(f) \setminus \Delta S_1} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{S_1(f^*) \setminus \Delta S_1} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right), \end{aligned}$$

where $\Delta S_{-1} = S_{-1}(f^*) \cap S_{-1}(f)$ and $\Delta S_1 = S_1(f^*) \cap S_1(f)$. We can easily verify that $S_{-1}(f) \setminus \Delta S_{-1} = S_{-1}(f^*) \setminus \Delta S_1$ and $S_1(f) \setminus \Delta S_1 = S_{-1}(f^*) \setminus \Delta S_{-1}$. Therefore, it follows that

$$\begin{aligned} R(f) - R(f^*) &= \left(\int_{\Delta_1} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{\Delta_2} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right) + \left(\int_{\Delta_2} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{\Delta_1} (1 - \eta(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right) \\ &\geq 2 \left(\int_{\Delta_1} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{\Delta_2} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \right) - \int_{\Delta_1 \cup \Delta_2} \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x}, \end{aligned}$$

where $\Delta_1 = S_{-1}(f) \setminus \Delta S_{-1}$ and $\Delta_2 = S_1(f) \setminus \Delta S_1$.

Then

$$\begin{aligned} &\int_{\Delta_1} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{\Delta_2} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \leq \int_{\Delta_1 \cup \Delta_2} \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} + R(f) - R(f^*) \\ &\leq \int_{\Delta_1 \cup \Delta_2} I(|2\eta(\mathbf{x}) - 1| \leq t) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} + \int_{\Delta_1 \cup \Delta_2} I(|2\eta(\mathbf{x}) - 1| > t) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} + R(f) - R(f^*) \\ &\leq ct^\gamma + \int_{\Delta_1 \cup \Delta_2} I(|2\eta(\mathbf{x}) - 1| > t) \frac{|2\eta(\mathbf{x}) - 1|}{t} \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} + R(f) - R(f^*) \\ &\leq ct^\gamma + t^{-1}(R(f) - R(f^*)) + R(f) - R(f^*) \end{aligned}$$

where the last inequality follows from the low-noise assumption. Choosing $t = (R(f) - R(f^*))^{\frac{1}{1+\gamma}} / c^{\frac{1}{1+\gamma}}$, we get

$$\int_{\Delta_1} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} - \int_{\Delta_2} \eta(\mathbf{x}) \mathbb{P}_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \leq 2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*)$$

Finally, we have

$$\text{EFNE}(f) - \text{EFNE}(f^*) \leq \frac{1}{2\mathbb{P}(Y = 1)} \left(2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*) \right). \quad (9)$$

Using similar steps, we also have

$$\text{EFPE}(f) - \text{EFPE}(f^*) \leq \frac{1}{2\mathbb{P}(Y = -1)} \left(2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*) \right). \quad (10)$$

Combining (9) and (10) yields that

$$\begin{aligned} & \max\{\text{EFNE}(f) - \text{EFNE}(f^*), \text{EFPE}(f) - \text{EFPE}(f^*)\} \\ & \leq \frac{1}{2\min\{\mathbb{P}(Y = 1), \mathbb{P}(Y = -1)\}} \left(2c^{-\frac{1}{1+\gamma}} (R(f) - R(f^*))^{\frac{\gamma}{1+\gamma}} + R(f) - R(f^*) \right). \end{aligned}$$

The desired results immediately follows by letting γ go to infinity, which completes the proof. $\square$

Theorem 1 (Restatement of Theorem 1). *Under Assumptions 1 and 2, for any minimizer $\tilde{f}_n$ of (2), there exist some positive constants A_1 and A_2 such that*

$$A_2 \left\{ \left(\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2} \right)^{\frac{\gamma+1}{\gamma+2}} + \tau_n \right\} \leq \sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{f}_n) - R(f^*)] \leq A_1 \left\{ \left(\frac{\mathcal{V}_1(\Theta) \log(n)}{n\kappa_\epsilon^2} \right)^{\frac{\gamma+1}{\gamma+2}} + s_n \right\},$$

where $\mathcal{P}_\gamma$ be a class of distributions of $(\mathbf{X}, Y)$ satisfying Assumption 1, $s_n = \sup_{\pi \in \mathcal{P}_\gamma} \inf_{f \in \mathcal{F}} H(f, f^*)$, and $\tau_n = \sup_{\pi \in \mathcal{P}_\gamma} \inf_{f \in \mathcal{F}} D(f, f^*)$.

Proof of Theorem 1: Proof of the upper bound. The proof of the upper bound mainly utilizes a uniform concentration inequality. We first denote that $\tilde{H}(f, f^*) = \tilde{R}_\phi(f) - \tilde{R}_\phi(f^*)$. Next, we proceed to prove that $\mathbb{P}(\tilde{H}(\tilde{f}_n, f^*) \geq \delta_n)$ converges to zero as n goes to infinity for some convergent sequence $\delta_n > 0$. For any $\delta_n > 0$, we let $\mathcal{F}_{\delta_n} = \{f \in \mathcal{F} : \tilde{H}(f, f^*) \geq \delta_n\}$. If $\tilde{f}_n \in \mathcal{F}_{\delta_n}$, one has

$$\sup_{f \in \mathcal{F}_{\delta_n}} \frac{1}{n} \sum_{i=1}^n \phi(f_{\mathcal{F}}^*(\mathbf{x}_i) \tilde{y}_i) + \lambda_n J(f_{\mathcal{F}}^*) - \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i) \tilde{y}_i) - \lambda_n J(f) \geq 0.$$

This follows from the optimality of $\tilde{f}_n$ for minimizing (2). Therefore,

$$\mathbb{P}(\tilde{H}(\tilde{f}_n, f^*) \geq \delta_n) \leq \mathbb{P} \left(\sup_{f \in \mathcal{F}_{\delta_n}} \frac{1}{n} \sum_{i=1}^n \phi(f_{\mathcal{F}}^*(\mathbf{x}_i) \tilde{y}_i) + \lambda_n J(f_{\mathcal{F}}^*) - \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i) \tilde{y}_i) - \lambda_n J(f) \geq 0 \right) \equiv I.$$

Next, it suffices to prove the convergence of I with n . To this end, we proceed to provide an upper bound for I . Define that $\mathcal{H}_{ij} = \{f \in \mathcal{F} : 2^{i-1}\delta_n \leq \tilde{H}(f, f^*) \leq 2^i\delta_n, 2^{j-1}J_0 \leq J(f) \leq 2^jJ_0\}$ for $i \geq 1$ and $j \geq 1$ and $\mathcal{H}_{i0} = \{f \in \mathcal{F} : 2^{i-1}\delta_n \leq \tilde{H}(f, f^*) \leq 2^i\delta_n, J(f) \leq J_0\}$. It can be easily verified that $\mathcal{F}_{\delta_n}$ admits the decomposition as $\mathcal{F}_{\delta_n} = \cup_{i=1}^{\infty} \cup_{j=0}^{\infty} \mathcal{H}_{ij}$. With this, I can be upper bounded as

$$\begin{aligned} I &= \mathbb{P} \left(\sup_{f \in \mathcal{F}_{\delta_n}} \frac{1}{n} \sum_{i=1}^n \phi(f_{\mathcal{F}}^*(\mathbf{x}_i) \tilde{y}_i) + \lambda_n J(f_{\mathcal{F}}^*) - \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i) \tilde{y}_i) - \lambda_n J(f) \geq 0 \right) \\ &\leq \sum_{i=1}^{\infty} \sum_{j=0}^{\infty} \mathbb{P} \left(\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n \phi(f_{\mathcal{F}}^*(\mathbf{x}_i) \tilde{y}_i) + \lambda_n J(f_{\mathcal{F}}^*) - \frac{1}{n} \sum_{i=1}^n \phi(f(\mathbf{x}_i) \tilde{y}_i) - \lambda_n J(f) \geq 0 \right) \\ &\leq \sum_{i=1}^{\infty} \sum_{j=0}^{\infty} \mathbb{P} \left(\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n l_\phi(f_{\mathcal{F}}^*, z_i) - \frac{1}{n} \sum_{i=1}^n l_\phi(f, z_i) \geq \lambda_n \left(\inf_{f \in \mathcal{H}_{ij}} J(f) - J_0 \right) + \inf_{f \in \mathcal{H}_{ij}} \mathbb{E}(\phi(f(\mathbf{X}_i) \tilde{Y}_i)) - \mathbb{E}(\phi(f_{\mathcal{F}}^*(\mathbf{X}_i) \tilde{Y}_i)) \right), \end{aligned}$$

where $z_i = (\mathbf{x}_i, \tilde{y}_i)$ and $l_\phi(f, z_i) = \phi(f(\mathbf{x}_i) \tilde{y}_i) - \mathbb{E}(\phi(f(\mathbf{X}_i) \tilde{Y}_i))$. Here it is important to note that

$$\begin{aligned} & \inf_{f \in \mathcal{H}_{ij}} \mathbb{E}(\phi(f(\mathbf{X}_i) \tilde{Y}_i)) - \mathbb{E}(\phi(f_{\mathcal{F}}^*(\mathbf{X}_i) \tilde{Y}_i)) = \inf_{f \in \mathcal{H}_{ij}} \tilde{R}_\phi(f) - \tilde{R}_\phi(f^*) - \tilde{R}_\phi(f_{\mathcal{F}}^*) + \tilde{R}_\phi(f^*) \\ &= \inf_{f \in \mathcal{H}_{ij}} \tilde{H}(f, f^*) + \tilde{H}(f_{\mathcal{F}}^*, f^*) = \inf_{f \in \mathcal{H}_{ij}} \tilde{H}(f, f^*) + (2\theta - 1)H(f_{\mathcal{F}}^*, f^*). \end{aligned}$$

Let $V(i, j) = \lambda_n (\inf_{f \in \mathcal{H}_{ij}} J(f) - J_0) + \inf_{f \in \mathcal{H}_{ij}} \mathbb{E}(l_\phi(f, Z)) - \mathbb{E}(l_\phi(f_{\mathcal{F}}^*, Z))$. By the definition of $\mathcal{H}_{ij}$, we get

$$V(i, j) \geq M(i, j) = \lambda_n (2^{j-1} - 1) J_0 + (2^{i-1} - 1/4) \delta_n, \text{ for } i, j \geq 1, \quad (11)$$

where the inequality follows by assuming that $(2\theta - 1)s_n \leq 1/4\delta_n$. Next, we suppose that $\lambda_n J_0 \leq 1/4\delta_n$, we further have

$$M(i, 0) \geq (2^{i-1} - 1/2) \delta_n \geq 2^{i-2} \delta_n, \text{ for } i \geq 1. \quad (12)$$

Plugging (11) and (12) into I , it follows that

$$I \leq \sum_{i=1}^{\infty} \sum_{j=0}^{\infty} \mathbb{P} \left(\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n l_\phi(f_{\mathcal{F}}^*, z_i) - \frac{1}{n} \sum_{i=1}^n l_\phi(f, z_i) \geq M(i, j) \right) = \sum_{i=1}^{\infty} \sum_{j=0}^{\infty} P_{ij}.$$

Therefore, bounding I reduces to bounding P_{ij} separately. Let $Q_n = \frac{1}{n} \sum_{i=1}^n (l_\phi(f_{\mathcal{F}}^*, z_i) - l_\phi(f, z_i))$, then P_{ij} can be written as

$$P_{ij} = \mathbb{P} \left(\sup_{f \in \mathcal{H}_{ij}} Q_n - \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} Q_n \right] \geq M(i, j) - \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} Q_n \right] \right).$$

Next, we proceed to bound P_{ij} by the Talagrand's inequality (see Theorem 2.6 in Koltchinskii, 2011). To this end, we first establish the relation between $\mathbb{E} [\sup_{f \in \mathcal{H}_{ij}} Q_n]$ and $M(i, j)$. Let $q(f, z_i) = \phi(f_{\mathcal{F}}^*(\mathbf{x}_i) y_i) - \phi(f(\mathbf{x}_i) y_i)$ and $\mathbf{z}' = (z'_1, \dots, z'_n)$ be a ghost sample.

$$\begin{aligned} \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} Q_n \right] &= \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n q(f, z_i) - \mathbb{E}(q(f, Z)) \right] \\ &= \mathbb{E}_{\mathbf{z}} \left[\sup_{f \in \mathcal{H}_{ij}} \mathbb{E}_{\mathbf{z}'} \left(\frac{1}{n} \sum_{i=1}^n q(f, z_i) - \frac{1}{n} \sum_{i=1}^n q(f, z'_i) | \mathbf{z} \right) \right] \\ &\leq \mathbb{E}_{\mathbf{z}, \mathbf{z}'} \left[\sup_{f \in \mathcal{H}_{ij}} \left(\frac{1}{n} \sum_{i=1}^n q(f, z_i) - \frac{1}{n} \sum_{i=1}^n q(f, z'_i) \right) \right] \\ &= \mathbb{E}_{\mathbf{z}, \mathbf{z}', \sigma} \left[\sup_{f \in \mathcal{H}_{ij}} \left(\frac{1}{n} \sum_{i=1}^n \sigma_i (q(f, z_i) - q(f, z'_i)) \right) \right] \\ &\leq 2 \mathbb{E}_{\mathbf{z}} \mathbb{E}_{\sigma} \left[\sup_{f \in \mathcal{H}_{ij}} \left(\frac{1}{n} \sum_{i=1}^n \sigma_i q(f, z_i) \right) \right] = 2 \mathcal{R}_n(\mathcal{H}_{ij}). \end{aligned}$$

By Theorem 3.11 in Koltchinskii (2011), it follows that there exists some constant C_2 such that

$$\mathcal{R}_n(\mathcal{H}_{ij}) \leq \frac{C_2}{\sqrt{n}} \mathbb{E} \int_0^{\sigma_{ij}} \sqrt{\log \mathcal{N}(u, \mathcal{H}_{ij}, L_2(P_n))} du \leq \frac{C_2}{\sqrt{n}} \mathbb{E} \int_0^{\sigma_{ij}} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1} \mathcal{V}_2(\Theta))} du,$$

where $\sigma_{ij} = \sqrt{\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n q^2(f, z_i)}$, $\|f\|_{L_2(P_n)} = \sqrt{\frac{1}{n} \sum_{i=1}^n q^2(f, z_i)}$, and $\mathcal{N}(\mathcal{H}_{ij}, L_2(P_n), u)$ is the minimal number of $L_2(P_n)$ -balls of radius u to cover $\mathcal{H}_{ij}$.

Notice that $\int_0^{\sigma_{ij}} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1} \mathcal{V}_2(\Theta))} du$ is concave function with respect to σ_{ij} , therefore

$$\mathbb{E} \left\{ \int_0^{\sigma_{ij}} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1} \mathcal{V}_2(\Theta))} du \right\} \leq \int_0^{\mathbb{E} \sigma_{ij}} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1} \mathcal{V}_2(\Theta))} du \leq \int_0^{\sqrt{\mathbb{E} \sigma_{ij}^2}} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1} \mathcal{V}_2(\Theta))} du.$$

By the definition of σ_{ij} , we have $\mathbb{E} [\sigma_{ij}^2] = \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} \frac{1}{n} \sum_{i=1}^n q^2(f, z_i) \right]$. Using symmetrization and contraction inequalities, we get

$$\mathbb{E} [\sigma_{ij}^2] \leq \sup_{f \in \mathcal{H}_{ij}} \mathbb{E} [q^2(f, Z)] + 8 \mathcal{R}_n(f, \mathcal{H}_{ij}).$$

Next, we proceed to bound $\mathbb{E}[q^2(f, Z)]$.

$$\begin{aligned}
\mathbb{E}[q^2(f, Z)] &= \mathbb{E} \left[\phi(f_{\mathcal{F}}^*(\mathbf{X})\tilde{Y}) - \phi(f(\mathbf{X})\tilde{Y}) \right]^2 \\
&\leq 2\mathbb{E} \left[\phi(f_{\mathcal{F}}^*(\mathbf{X})\tilde{Y}) - \phi(f_{\phi}^*(\mathbf{X})\tilde{Y}) \right]^2 + 2\mathbb{E} \left[\phi(f_{\mathcal{F}}^*(\mathbf{X})\tilde{Y}) - \phi(f(\mathbf{X})\tilde{Y}) \right]^2 \\
&\leq 2\mathbb{E} \left[\phi(f_{\mathcal{F}}^*(\mathbf{X})\tilde{Y}) - \phi(f^*(\mathbf{X})\tilde{Y}) \right] + 2\mathbb{E} \left[|\phi(f(\mathbf{X})\tilde{Y}) - \phi(f_{\mathcal{F}}^*(\mathbf{X})\tilde{Y})| \right] \\
&\leq 2s_n + 2\mathbb{E} [|f(\mathbf{X}) - f_{\mathcal{F}}^*(\mathbf{X})|] \leq 2^{-1}\delta_n + 2\mathbb{E} [|f(\mathbf{X}) - f_{\mathcal{F}}^*(\mathbf{X})|], \tag{13}
\end{aligned}$$

where the second inequality follows from the assumption that $\|f\|_{\infty} \leq 1$. Combining this with Lemma 3 yields that

$$\sup_{f \in \mathcal{H}_{ij}} \mathbb{E}[q^2(f, Z)] \leq 2^{-1}\delta_n + 2C_1((2\theta - 1)^{-1}(\tilde{R}_{\phi}(f) - \tilde{R}_{\phi}(f^*)))^{\frac{\gamma}{\gamma+1}} = 4C_1(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}(2^i\delta_n)^{\frac{\gamma}{\gamma+1}},$$

where $C_1 = (4c)^{\frac{1}{\gamma+1}}$. Consequently, we get

$$\mathcal{R}_n(\mathcal{H}_{ij}) \leq \frac{C_2}{\sqrt{n}} \int_0^{U_{ij}(f)} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1}\mathcal{V}_2(\Theta))} du, \tag{14}$$

where $U_{ij}(f) = \min \left\{ \sqrt{4C_1(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}(2^i\delta_n)^{\frac{\gamma}{\gamma+1}} + 8\mathcal{R}_n(\mathcal{H}_{ij})}, 1 \right\}$ due to the fact that $|q(f, Z)| \leq 1$. Then, the right-hand side of (14) can be upper bounded as

$$\begin{aligned}
&\frac{C_2}{\sqrt{n}} \int_0^{U_{ij}(f)} \sqrt{\mathcal{V}_1(\Theta) \log(u^{-1}\mathcal{V}_2(\Theta))} du = \frac{C_2\mathcal{V}_2(\Theta)\sqrt{\mathcal{V}_1(\Theta)}}{\sqrt{n}} \int_0^{U_{ij}(f)/\mathcal{V}_2(\Theta)} \sqrt{\log\left(\frac{1}{u}\right)} du \\
&= \frac{C_2\mathcal{V}_2(\Theta)\sqrt{\mathcal{V}_1(\Theta)}}{\sqrt{n}} \int_{\frac{\mathcal{V}_2(\Theta)}{U_{ij}(f)}}^{+\infty} \frac{1}{u^2} \sqrt{\log(u)} du \leq \frac{2C_2\sqrt{\mathcal{V}_1(\Theta)}U_{ij}(f)}{\sqrt{n}} \sqrt{\log\left(\frac{\mathcal{V}_2(\Theta)}{U_{ij}(f)}\right)}. \tag{15}
\end{aligned}$$

Combining (15) with (14) yields that

$$\left(\mathcal{R}_n(\mathcal{H}_{ij})\right)^2 \leq \frac{4C_2^2\mathcal{V}_1(\Theta)\left(4C_1(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}(2^i\delta_n)^{\frac{\gamma}{\gamma+1}} + 8\mathcal{R}_n(\mathcal{H}_{ij})\right)}{n} \log\left(\frac{\mathcal{V}_2(\Theta)}{U_{ij}(f)}\right). \tag{16}$$

Solving (16) gives $\mathcal{R}_n(f, \mathcal{H}_{ij}) \lesssim \max\{\mathcal{V}_1(\Theta)n^{-1}, n^{-1/2}\mathcal{V}_1(\Theta)^{1/2}(2\theta - 1)^{-\frac{\gamma}{2(\gamma+1)}}\delta_n^{\frac{\gamma}{2(\gamma+1)}}\}$. Therefore, we get $\mathcal{R}_n(f, \mathcal{H}_{ij}) \leq 1/4\delta_n$ provided that

$$(\mathcal{V}_1(\Theta)/n)^{\frac{\gamma+1}{\gamma+2}}(2\theta - 1)^{-\frac{\gamma}{\gamma+2}} \log(n/\mathcal{V}_1(\Theta)) \leq C\delta_n, \tag{17}$$

for some large constants C . Hence, as n goes to infinity, it follows that

$$\mathcal{R}_n(f, \mathcal{H}_{ij}) \leq 1/4\delta_n \leq 1/4M(i, j).$$

With this, we get

$$\mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} Q_n \right] \leq \mathcal{R}_n(f, \mathcal{H}_{ij}) \leq 1/4M(i, j).$$

Therefore, P_{ij} can be further bounded as

$$P_{ij} \leq \mathbb{P} \left(\sup_{f \in \mathcal{H}_{ij}} Q_n - \mathbb{E} \left[\sup_{f \in \mathcal{H}_{ij}} Q_n \right] \geq 1/2M(i, j) \right). \tag{18}$$

Applying Talagrand's inequality to the right-hand side, it follows that there exists some positive constants C_4 such that

$$P_{ij} \leq C_4 \exp \left(-\frac{nM(i,j)}{2C_4} \log \left(1 + \frac{M(i,j)}{2\mathbb{E}[\sigma_{ij}^2]} \right) \right).$$

As proved above, we can verify that there exists some constant C_5

$$\mathbb{E}[\sigma_{ij}^2] \leq \sup_{f \in \mathcal{H}_{ij}} \mathbb{E}[q^2(f, Z)] + 8\mathcal{R}_n(f, \mathcal{H}_{ij}) \leq C_5 \left((2\theta - 1)^{-\frac{\gamma}{\gamma+1}} (2^i \delta_n)^{\frac{\gamma}{\gamma+1}} + M(i, j) \right).$$

Notice that $\frac{M(i,j)}{2\mathbb{E}[\sigma_{ij}^2]}$ converges to 0 as n increases, therefore there exists some constant $0 < C_7 < 1$ such that $\log(1+x) \geq C_7 x$ for $x \in [0, 1/(2C_5)]$. It then follows that

$$P_{ij} \leq C_4 \exp \left(-\frac{C_7 n M^2(i, j)}{4C_4 C_5 ((2\theta - 1)^{-\frac{\gamma}{\gamma+1}} (M(i, j))^{\frac{\gamma}{\gamma+1}} + M(i, j))} \right).$$

Since $M(i, j) \ll (2\theta - 1)^{-\frac{\gamma}{\gamma+1}} (M(i, j))^{\frac{\gamma}{\gamma+1}}$ when $\theta - 1/2 = o(1)$ and $M(i, j) = o(1)$, we have

$$P_{ij} \leq C_4 \exp \left(-C_8 n (2\theta - 1)^{\frac{\gamma}{\gamma+1}} M^{\frac{\gamma+2}{\gamma+1}}(i, j) \right),$$

where $C_8 = C_7/(4C_4 C_5)$. Therefore, we have

$$\sum_{i=1}^n \sum_{j=0}^n P_{ij} \leq \sum_{i=1}^n \sum_{j=1}^n C_4 \exp \left(-\frac{C_8 n M^{\frac{\gamma+2}{\gamma+1}}(i, j)}{(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right) + \sum_{i=1}^n C_4 \exp \left(-\frac{C_8 n M^{\frac{\gamma+2}{\gamma+1}}(i, 0)}{(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right) \equiv I_1 + I_2.$$

Note that for any $i, j \geq 1$,

$$\begin{aligned} M^{\frac{\gamma+2}{\gamma+1}}(i, j) &\geq ((2^{i-1} - 1/4)\delta_n + \lambda_n(2^{j-1} - 1)J_0)^{\frac{\gamma+2}{\gamma+1}} \\ &\geq (2^{i-1} - 1/4)\delta_n^{\frac{\gamma+2}{\gamma+1}} + (2^{j-1} - 1)(\lambda_n J_0)^{\frac{\gamma+2}{\gamma+1}} \\ &\geq i/2\delta_n^{\frac{\gamma+2}{\gamma+1}} + (j-1)(\lambda_n J_0)^{\frac{\gamma+2}{\gamma+1}}. \end{aligned}$$

Hence I_1 is upper bounded as

$$\begin{aligned} I_1 &\leq \sum_{i=1}^n \sum_{j=1}^n C_4 \exp \left(-C_8 n \frac{i/2\delta_n^{\frac{\gamma+2}{\gamma+1}} + (j-1)(\lambda_n J_0)^{\frac{\gamma+2}{\gamma+1}}}{(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right) \\ &= C_4 \frac{\exp \left(\frac{-C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right)}{1 - \exp \left(\frac{-C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right)} \cdot \frac{1}{1 - \exp \left(\frac{-C_8 n (\lambda_n J_0)^{\frac{\gamma+2}{\gamma+1}}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right)} \leq 4C_4 \exp \left(\frac{-C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right), \end{aligned}$$

where the last inequality holds when $\max \left\{ \exp \left(\frac{-C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right), \exp \left(\frac{-C_8 n \delta_n^{2-1/\gamma}}{2(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right) \right\} \leq 1/2$, which holds true when n goes to infinity. Similarly, for I_2 , we get

$$I_2 \leq \sum_{i=1}^n C_4 \exp \left(-\frac{i/4C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right) \leq 2C_4 \exp \left(-\frac{C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{4(2\theta - 1)^{-\frac{\gamma}{\gamma+1}}} \right).$$

Combining I_1 and I_2 , it follows that

$$I \leq 6C_4 \exp \left(-\frac{C_8 n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{4(2\theta-1)^{-\frac{\gamma}{\gamma+1}}} \right) = T_1 \exp \left(-T_2 \frac{n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{(2\theta-1)^{-\frac{\gamma}{\gamma+1}}} \right),$$

where $T_1 = 6C_4$ and $T_2 = C_8/4$. Therefore, we conclude that

$$\mathbb{P}(\tilde{H}(\tilde{f}_n, f^*) \geq \delta_n) \leq T_1 \exp \left(-T_2 \frac{n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{(2\theta-1)^{-\frac{\gamma}{\gamma+1}}} \right).$$

With a choice of δ_n such that $\delta_n \geq C_9((2\theta-1)^{-\frac{\gamma}{\gamma+1}}|\Theta|n^{-1}\log(n/|\Theta|))^{\frac{\gamma+1}{\gamma+2}}$ for some positive constants $C_9 > 0$, we have $\tilde{H}(\tilde{f}_n, f^*) = o_p(1)$, which implies that $\mathbb{E}(\tilde{H}(\tilde{f}_n, f^*)) = O_p(\delta_n)$. Notice that T_1 and T_2 are both independent of π , therefore we further have

$$\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{P}(\tilde{H}(\tilde{f}_n, f^*) \geq \delta_n) \leq T_1 \exp \left(-T_2 \frac{n \delta_n^{\frac{\gamma+2}{\gamma+1}}}{(2\theta-1)^{-\frac{\gamma}{\gamma+1}}} \right). \quad (19)$$

By the relation between excess risk and excess ϕ -risk $\mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{D}(\tilde{f}_n, f^*)) \leq \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{H}(\tilde{f}_n, f^*))$ (Bartlett et al., 2006), we further have

$$\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{D}(\tilde{f}_n, f^*)) = O(\delta_n).$$

Combining this with the Lemma 2 that $D(\tilde{f}_n, f^*) = (2\theta-1)^{-1}\tilde{D}(\tilde{f}_n, f^*)$, we get

$$\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}(D(\tilde{f}_n, f^*)) = O(\delta_n(2\theta-1)^{-1}).$$

The fastest rate for δ_n can be obtained by choosing the fastest rate such that the right-hand side of (19) converges to zero with n and (17) holds, which yields that

$$\delta \asymp (\mathcal{V}_1(\Theta)/n)^{\frac{\gamma+1}{\gamma+2}}(2\theta-1)^{-\frac{\gamma}{\gamma+2}}\log(n).$$

This completes the proof of the upper bound.

Proof of lower bound. We first define the minimax excess risk as

$$W_n = \inf_f \sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}[R(f) - R(f^*)],$$

which admits the decomposition as

$$W_n = \inf_f \sup_{\pi \in \mathcal{P}_\gamma} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}}[R(f) - \inf_{f \in \mathcal{F}} R(f)] + \inf_{f \in \mathcal{F}} R(f) - R(f^*) \right\},$$

where the second term on the right-hand side denotes the approximation error under the 0-1 risk. In what follows, we proceed to consider a sub-family of $\mathcal{P}_\gamma$ such that $\inf_{f \in \mathcal{F}} R(f) = R(f^*)$. For example, let $\mathcal{P}' \subset \mathcal{P}_\gamma$ be such a sub-family, then we have

$$\begin{aligned} W_n &\geq \inf_f \sup_{\pi \in \mathcal{P}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}}[R(f) - \inf_{f \in \mathcal{F}} R(f)] \right\} + \sup_{\pi \in \mathcal{P}_\gamma} \left(\inf_{f \in \mathcal{F}} R(f) - R(f^*) \right) \\ &= \inf_f \sup_{\pi \in \mathcal{P}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}}[R(f) - R(f^*)] \right\} + \sup_{\pi \in \mathcal{P}_\gamma} \left(\inf_{f \in \mathcal{F}} R(f) - R(f^*) \right) \\ &= (2\theta-1)^{-1} \inf_f \sup_{\tilde{\pi} \in \tilde{\mathcal{P}}'(\theta)} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(f) - \tilde{R}(f^*)] \right\} + \sup_{\pi \in \mathcal{P}_\gamma} \left(\inf_{f \in \mathcal{F}} R(f) - R(f^*) \right), \end{aligned}$$

where the last inequality follows from Lemma 2 and $\tilde{\mathcal{P}}'(\theta)$ is a set of probability measures $\tilde{\pi}$ on $(\mathbf{X}, \tilde{Y})$ such that any probability distribution $\tilde{\pi} \in \tilde{\mathcal{P}}'(\theta)$ is associated with an $\pi \in \mathcal{P}'$ with $\tilde{\pi}$ and π having the same marginal distribution on $\mathbf{X}$ and $\tilde{Y} = \mathcal{A}_\theta(Y)$.

To provide a lower bound for W_n , it suffices to bound $\inf_f \sup_{\tilde{\pi} \in \tilde{\mathcal{P}}'(\theta)} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(f) - \tilde{R}(f^*)] \right\}$. For ease of notation, we denote $\widetilde{W}_n = \inf_f \sup_{\tilde{\pi} \in \tilde{\mathcal{P}}'(\theta)} \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(f) - \tilde{R}(f^*)]$ and $\tilde{\mathcal{P}}_\gamma(\theta) = \{\tilde{\pi}(\mathbf{X}, \tilde{Y}) : \tilde{Y} = \mathcal{A}_\theta(Y), \pi \in \mathcal{P}_\gamma\}$. Then we proceed to construct $\tilde{\mathcal{P}}'(\theta) \subset \tilde{\mathcal{P}}_\gamma(\theta)$. First, following from Lemma 3, we have for any $\tilde{\pi} \in \tilde{\mathcal{P}}_\gamma(\theta)$

$$\mathbb{P}\left(|2\tilde{\eta}(\mathbf{X}) - 1| \leq t\right) \leq c(2\theta - 1)^{-\gamma} t^\gamma,$$

where the probability is measured under $\tilde{\pi}$.

By Assumption 2, we know $VC(\mathcal{G}_{\mathcal{F}}) \asymp \mathcal{V}_1(\Theta)$. For any N such that $N \asymp \mathcal{V}_1(\Theta)$, there exist N distinct points $\mathbf{x}_1, \dots, \mathbf{x}_N$ such that $\{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ is shattered by $\mathcal{G}_{\mathcal{F}}$. Therefore, we consider distribution supported on $\{\mathbf{x}_1, \dots, \mathbf{x}_N\}$. Let $w \in (0, 1)$ be a number satisfying $(N - 1)w \leq 1$. Let Q be the probability measure on $\mathcal{X}$ such that

$$Q(\mathbf{x}_i) = \begin{cases} w, & i = 1, \dots, N - 1, \\ 1 - (N - 1)w, & i = N. \end{cases}$$

In what follows, we consider the hypercube $\mathcal{C} = \{-1, 1\}^{N-1}$. For any $\boldsymbol{\sigma} = (\sigma_1, \dots, \sigma_{N-1}) \in \mathcal{C}$, we define the regression function as

$$\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x}_i) = \begin{cases} \frac{1 + \sigma_i h}{2}, & i = 1, \dots, N - 1, \\ 1, & i = N. \end{cases}$$

Next, we let $\tilde{\pi}_{\boldsymbol{\sigma}}$ denote the associated probability measure on $\mathcal{X} \times \{-1, 1\}$ with Q being marginal distribution on $\mathcal{X}$ and $\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x})$ being the regression function. With this, we obtain

$$\mathbb{P}(|2\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{X}) - 1| \leq t) = (N - 1)wI(h \leq t).$$

By assuming $(N - 1)w \leq c(2\theta - 1)^{-\gamma} h^\gamma$, we have $\mathbb{P}(|2\eta_{\boldsymbol{\sigma}}(\mathbf{X}) - 1| \leq t) = c(2\theta - 1)^{-\gamma} t^\gamma$ for any $t \in [0, 1]$, which indicates that $\tilde{\pi}_{\boldsymbol{\sigma}} \in \tilde{\mathcal{P}}_\gamma(\theta)$. By setting $\tilde{\mathcal{P}}'(\theta) = \{\tilde{\pi}_{\boldsymbol{\sigma}} : \boldsymbol{\sigma} \in \mathcal{C}\}$, we have

$$\widetilde{W}_n = \inf_f \sup_{\tilde{\pi}_{\boldsymbol{\sigma}} \in \tilde{\mathcal{P}}'(\theta)} \mathbb{E}[\tilde{R}(f) - \tilde{R}(f_{\boldsymbol{\sigma}}^*)],$$

where $f_{\boldsymbol{\sigma}}^*(\mathbf{x}_i) = \sigma_i$ for $i = 1, \dots, N - 1$.

Under $\tilde{\pi}_{\boldsymbol{\sigma}}$,

$$\begin{aligned} \tilde{R}(f) - \tilde{R}(f_{\boldsymbol{\sigma}}^*) &= \mathbb{E}_{\tilde{\pi}_{\boldsymbol{\sigma}}} [|\text{sign}(f(\mathbf{X})) - f_{\boldsymbol{\sigma}}^*(\mathbf{X})| |1 - 2\tilde{\eta}(\mathbf{X})|] \\ &= \frac{1}{2} \mathbb{E}_{\tilde{\pi}_{\boldsymbol{\sigma}}} [|\text{sign}(f(\mathbf{X})) - f_{\boldsymbol{\sigma}}^*(\mathbf{X})| |1 - 2\tilde{\eta}(\mathbf{X})|] \\ &\geq \frac{t}{2} \mathbb{E}_{\tilde{\pi}_{\boldsymbol{\sigma}}} \left[|\text{sign}(f(\mathbf{X})) - f_{\boldsymbol{\sigma}}^*(\mathbf{X})| \left(1 - \mathbb{P}(|1 - 2\tilde{\eta}(\mathbf{X})| \leq t)\right) \right] \\ &\geq (2\theta - 1)(4^{1-\gamma} c)^{-1/\gamma} \left(\mathbb{E}_{\tilde{\pi}_{\boldsymbol{\sigma}}} [|\text{sign}(f(\mathbf{X})) - f_{\boldsymbol{\sigma}}^*(\mathbf{X})|] \right)^{\frac{\gamma+1}{\gamma}} \\ &\geq (2\theta - 1)(4^{1-\gamma} c)^{-1/\gamma} \left(w \sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - f_{\boldsymbol{\sigma}}^*(\mathbf{x}_i)| \right)^{\frac{\gamma+1}{\gamma}}. \end{aligned}$$

Under the distribution $\tilde{\pi}_{\boldsymbol{\sigma}}$, we have $f^*(\mathbf{x}_i) = \sigma_i$ for $i = 1, \dots, N - 1$, which then implies that

$$\tilde{R}(f) - \tilde{R}(f_{\boldsymbol{\sigma}}^*) \geq (2\theta - 1)(4^{1-\gamma} c)^{-1/\gamma} w^{\frac{\gamma+1}{\gamma}} \left(\sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - \sigma_i| \right)^{\frac{\gamma+1}{\gamma}}.$$

Taking the expectation of both sides with respect to $\tilde{\mathcal{D}}$ yields that

$$\begin{aligned}\mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{R}(f) - \tilde{R}(f_{\sigma}^*)) &\geq (2\theta - 1)(4^{1-\gamma}c)^{-1/\gamma} w^{\frac{\gamma+1}{\gamma}} \mathbb{E}_{\tilde{\mathcal{D}}} \left[\left(\sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - \sigma_i| \right)^{\frac{\gamma+1}{\gamma}} \right] \\ &\geq (2\theta - 1)(4^{1-\gamma}c)^{-1/\gamma} w^{\frac{\gamma+1}{\gamma}} \mathbb{E}_{\tilde{\mathcal{D}}}^{\frac{\gamma+1}{\gamma}} \left[\left(\sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - \sigma_i| \right) \right],\end{aligned}$$

where the second inequality follows from Jensen's inequality.

For ease of notation, we let $X(w, \gamma, \theta) = (2\theta - 1)(4^{1-\gamma}c)^{-1/\gamma} w^{\frac{\gamma+1}{\gamma}}$.

$$\begin{aligned}\sup_{\tilde{\pi}_{\sigma} \in \tilde{\mathcal{P}}'(\theta)} \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{R}(f) - \tilde{R}(f_{\sigma}^*)) &\geq \sup_{\tilde{\pi}_{\sigma} \in \tilde{\mathcal{P}}'(\theta)} X(w, \gamma, \theta) \mathbb{E}_{\tilde{\mathcal{D}}}^{\frac{\gamma+1}{\gamma}} \left[\left(\sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - \sigma_i| \right) \right], \\ &\geq \frac{1}{2^N} \sum_{\sigma \in \mathcal{C}} X(w, \gamma, \theta) \mathbb{E}_{\tilde{\mathcal{D}}}^{\frac{\gamma+1}{\gamma}} \left[\left(\sum_{i=1}^{N-1} |\text{sign}(f(\mathbf{x}_i)) - \sigma_i| \right) \right] \\ &= \frac{1}{2^{N-1}} \sum_{\sigma \in \mathcal{C}} X(w, \gamma, \theta) \mathbb{E}_{\tilde{\mathcal{D}}}^{\frac{\gamma+1}{\gamma}} \left[\sum_{i=1}^{N-1} I(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) \right] \\ &= X(w, \gamma, \theta) \left(\frac{1}{2^{N-1}} \sum_{\sigma \in \mathcal{C}} \sum_{i=1}^{N-1} \mathbb{E}_{\tilde{\mathcal{D}}} \left[I(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) \right] \right)^{\frac{\gamma+1}{\gamma}} \\ &= X(w, \gamma, \theta) \left(\sum_{i=1}^{N-1} \frac{1}{2^{N-1}} \sum_{\sigma \in \mathcal{C}} \mathbb{P}_{\tilde{\pi}_{\sigma}}(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) \right)^{\frac{\gamma+1}{\gamma}}.\end{aligned}$$

For each i , we observe that

$$\begin{aligned}&\frac{1}{2^{N-1}} \sum_{\sigma \in \mathcal{C}} \mathbb{P}_{\tilde{\pi}_{\sigma}}(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) \\ &= \frac{1}{2^{N-1}} \sum_{\sigma: \sigma_i=1} \mathbb{P}_{\tilde{\pi}_{\sigma}}(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) + \frac{1}{2^{N-1}} \sum_{\sigma: \sigma_i=-1} \mathbb{P}_{\tilde{\pi}_{\sigma}}(\text{sign}(f(\mathbf{x}_i)) \neq \sigma_i) \\ &= 2\mathbb{P}_{+i}(\text{sign}(f(\mathbf{x}_i)) \neq 1) + 2\mathbb{P}_{-i}(\text{sign}(f(\mathbf{x}_i)) \neq -1) \geq 2 - 2\text{TV}(\mathbb{P}_{+i}^{\otimes n}, \mathbb{P}_{-i}^{\otimes n}).\end{aligned}$$

where $\mathbb{P}_{+i} = \frac{1}{2^{N-1}} \sum_{\sigma: \sigma_i=1} \mathbb{P}_{\tilde{\pi}_{\sigma}}$, $\mathbb{P}_{-i} = \frac{1}{2^{N-1}} \sum_{\sigma: \sigma_i=-1} \mathbb{P}_{\tilde{\pi}_{\sigma}}$, and $\text{TV}(\mathbb{P}_{+i}^{\otimes n}, \mathbb{P}_{-i}^{\otimes n})$ denotes the total variation between $\mathbb{P}_{+i}^{\otimes n}$ and $\mathbb{P}_{-i}^{\otimes n}$. Notice that for any probability measures P and Q , we have

$$\begin{aligned}\text{TV}(P, Q) &= \frac{1}{2} \int |p(\mathbf{x}) - q(\mathbf{x})| d\mathbf{x} = \frac{1}{2} \int |\sqrt{p(\mathbf{x})} - \sqrt{q(\mathbf{x})}| |\sqrt{p(\mathbf{x})} + \sqrt{q(\mathbf{x})}| d\mathbf{x} \\ &\leq \frac{1}{2} \left(\int (\sqrt{p(\mathbf{x})} - \sqrt{q(\mathbf{x})})^2 d\mathbf{x} \right)^{1/2} \left(\int (\sqrt{p(\mathbf{x})} + \sqrt{q(\mathbf{x})})^2 d\mathbf{x} \right)^{1/2} \\ &\leq \sqrt{H^2(P, Q)},\end{aligned}$$

where $H^2(P, Q)$ denotes the Hellinger distance between P and Q . Let σ and σ' be two indexes such that $\sigma_l = \sigma'_l$ for $l \neq i$ and $\sigma_i = -\sigma'_i = 1$, then we have

$$\text{TV}(\mathbb{P}_{+i}^{\otimes n}, \mathbb{P}_{-i}^{\otimes n}) \leq \sum_{\sigma, \sigma'} \frac{1}{2^{N-2}} \text{TV}(\mathbb{P}_{\pi_{\sigma}}^{\otimes n}, \mathbb{P}_{\pi_{\sigma'}}^{\otimes n}) \leq \max_{\sigma, \sigma'} \text{TV}(\mathbb{P}_{\pi_{\sigma}}^{\otimes n}, \mathbb{P}_{\pi_{\sigma'}}^{\otimes n}) \leq \max_{\sigma, \sigma'} H(\mathbb{P}_{\pi_{\sigma}}^{\otimes n}, \mathbb{P}_{\pi_{\sigma'}}^{\otimes n}).$$

By the definition of $\tilde{\pi}_{\sigma}$ and $\tilde{\pi}_{\sigma'}$, we get

$$\begin{aligned}H^2(\mathbb{P}_{\pi_{\sigma}}^{\otimes n}, \mathbb{P}_{\pi_{\sigma'}}^{\otimes n}) &= w \sum_{i=1}^{N-1} \left(\sqrt{\eta_{\sigma}(\mathbf{x}_i)} - \sqrt{\eta_{\sigma'}(\mathbf{x}_i)} \right)^2 + \left(\sqrt{1 - \eta_{\sigma}(\mathbf{x}_i)} - \sqrt{1 - \eta_{\sigma'}(\mathbf{x}_i)} \right)^2 \\ &= 2w(1 - \sqrt{1 - h^2}) \leq 2wh^2,\end{aligned}$$

for any $h \in [0, 1]$. By the fact that $H^2(P^{\otimes n}, Q^{\otimes n}) = 2 - 2(1 - 2^{-1}H^2(P, Q))^n$, we have

$$\begin{aligned} H^2(\mathbb{P}_{\pi_{\sigma}}^{\otimes n}, \mathbb{P}_{\pi_{\sigma'}}^{\otimes n}) &= 2 - 2[1 - w(1 - \sqrt{1 - h^2})]^n \leq 2 - 2[1 - nw(1 - \sqrt{1 - h^2})] \\ &= 2nw(1 - \sqrt{1 - h^2}) \leq 1/16, \end{aligned}$$

where the last inequality holds by choosing w and h such that $wh^2 \leq n^{-1}/32$, from which it follows that $\text{TV}(\mathbb{P}_{+i}^n, \mathbb{P}_{-i}^n) \leq 1/4$. Notice that w and h are chosen to satisfy $(N-1)w \leq c(2\theta-1)^{-\gamma}h^\gamma$ and $wh^2 \leq n^{-1}/32$. For simplicity, we obtain the following solution by considering the case equalities hold.

$$h = \frac{(N-1)^{\frac{1}{\gamma+2}}(2\theta-1)^{\frac{\gamma}{\gamma+2}}}{(32cn)^{\frac{1}{\gamma+2}}} \quad \text{and} \quad w = \frac{c^{\frac{2}{\gamma+2}}(N-1)^{-\frac{2}{\gamma+2}}(2\theta-1)^{-\frac{2\gamma}{\gamma+2}}}{(32n)^{\frac{\gamma}{\gamma+2}}}.$$

Therefore, we can conclude that

$$\begin{aligned} \sup_{\tilde{\pi}_{\sigma} \in \tilde{\mathcal{P}}'(\theta)} \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{R}(f) - \tilde{R}(f^*)) &\geq (2\theta-1)(4^{1-\gamma})^{-1/\gamma} c^{-\frac{1}{\gamma+2}} \frac{(N-1)^{-\frac{2\gamma+2}{\gamma(\gamma+2)}}(2\theta-1)^{-\frac{2\gamma+2}{\gamma+2}}}{(32n)^{\frac{\gamma+1}{\gamma+2}}} \left(\frac{3(N-1)}{2}\right)^{\frac{\gamma+1}{\gamma}} \\ &= (2\theta-1)(4^{1-\gamma})^{-1/\gamma} c^{-\frac{1}{\gamma+2}} \frac{(N-1)^{\frac{\gamma+1}{\gamma+2}}(2\theta-1)^{-\frac{2\gamma+2}{\gamma+2}}}{(32n)^{\frac{\gamma+1}{\gamma+2}}} \left(\frac{3}{2}\right)^{\frac{\gamma+1}{\gamma}} \\ &= (2\theta-1)(4^{1-\gamma})^{-1/\gamma} c^{-\frac{1}{\gamma+2}} \left(\frac{N-1}{32n(2\theta-1)^2}\right)^{\frac{\gamma+1}{\gamma+2}} \left(\frac{3}{2}\right)^{\frac{\gamma+1}{\gamma}}. \end{aligned}$$

Choosing $N \asymp \mathcal{V}_1(\Theta)$ yields that

$$\inf_f \sup_{\tilde{\pi}_{\sigma} \in \tilde{\mathcal{P}}'(\theta)} \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{R}(f) - \tilde{R}(f^*)) \geq (2\theta-1)A_2 \left(\frac{\mathcal{V}_1(\Theta)}{n(2\theta-1)^2}\right)^{\frac{\gamma+1}{\gamma+2}},$$

where $A_2 \asymp (4^{1-\gamma})^{-1/\gamma} c^{-\frac{1}{\gamma+2}} (2^{-1})^{\frac{6\gamma+6}{\gamma+2}} (3/2)^{\frac{\gamma+1}{\gamma}}$.

Following from Lemma 2, it holds that

$$\inf_f \sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}(\tilde{R}(f) - R(f^*)) \geq A_2 \left(\frac{\mathcal{V}_1(\Theta)}{n(2\theta-1)^2}\right)^{\frac{\gamma+1}{\gamma+2}} + \sup_{\pi \in \mathcal{P}_\gamma} \left(\inf_{f \in \mathcal{F}} R(f) - R(f^*)\right).$$

This completes the proof of the lower bound. $\square$

Corollary 1. Suppose that $\mathcal{F}$ is chosen such that $s_n = 0$ and ϵ is adaptive to n such that $\epsilon = o(1)$. There exists some constants $A_3 > 0$ such that $\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(\tilde{f}_n) - R(f^*)] \geq A_3$ provided that $\epsilon \lesssim n^{-1/2}$.

Proof of Corollary 1: By Theorem 1, we have

$$\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(\tilde{f}_n) - R(f^*)] \geq A_2 \left\{ \left(\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2}\right)^{\frac{\gamma+1}{\gamma+2}} + \tau_n \right\}.$$

Since $s_n = 0$ implies $\tau_n = 0$, we further have

$$\sup_{\pi \in \mathcal{P}_\gamma} \mathbb{E}_{\tilde{\mathcal{D}}}[\tilde{R}(\tilde{f}_n) - R(f^*)] \geq A_2 \left\{ \left(\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2}\right)^{\frac{\gamma+1}{\gamma+2}} \right\}.$$

Without loss of generality, we assume that $\epsilon < 1$ since $\epsilon = o(1)$. Then, by the definition of κ_ϵ , $\kappa_\epsilon \asymp \epsilon$ when $\epsilon \rightarrow 0$. Therefore, we have

$$\frac{\mathcal{V}_1(\Theta)}{n\kappa_\epsilon^2} = \frac{\mathcal{V}_1(\Theta)(\exp(\epsilon) + 1)^2}{n(\exp(\epsilon) - 1)^2} \geq \frac{\mathcal{V}_1(\Theta)}{ne^2\epsilon^2}, \quad (20)$$

where the last inequality follows from the fact that $e^x - 1 > ex$ for any $x \in (0, 1]$. Further, if $\epsilon = o(1/\sqrt{n})$, we have $ne^2 = o(1)$. Therefore, it follows that $\mathcal{V}_1(\Theta)n^{-1}\kappa_\epsilon^{-2} \geq C\mathcal{V}_1(\Theta)e^{-2}$ for some constants C . The desired result immediately follows by setting $A_3 = (C\mathcal{V}_1(\Theta)e^{-2})^{\frac{\gamma+1}{\gamma+2}}$ and this completes the proof. $\square$

Theorem 2 (Restatement of Theorem 2). *Let $\mathcal{P}_{\gamma,\beta}$ be a class of probability measures on $\mathcal{X} \times \{-1, 1\}$ satisfying Assumption 1 and $\eta(\mathbf{X}) \in \mathcal{H}(\beta, [0, 1]^p, M)$. For any minimizer $\tilde{f}_{nn}$ in (4) with $L_n \asymp \log(\kappa_\epsilon n / \log(n))$, $N_n \asymp (\kappa_\epsilon n / \log(n))^{\frac{2p}{2\beta+p}}$, $B_n = 1$, and $P_n \asymp N_n \log(\kappa_\epsilon n / \log(n))$, we have*

$$\left(\frac{1}{n\kappa_\epsilon^2}\right)^{\frac{\beta(\gamma+1)}{\beta(\gamma+2)+p}} \lesssim \sup_{\pi \in \mathcal{P}_{\gamma,\beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] \lesssim \left(\frac{\log n}{n\kappa_\epsilon^2}\right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}}.$$

Particularly, $\sup_{\pi \in \mathcal{P}_{\gamma,\beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] = o(1)$ given that $\epsilon \gtrsim n^{-1/2+\zeta}$ for any $\zeta > 0$.

Proof of Theorem 2: Proof of the upper bound. Let $\tilde{Z} = (\tilde{Y} + 1)/2$. We intend to establish the connection between the excess risk of $\tilde{f}_{nn}$ and $\|\tilde{f}_{nn} - \eta\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2$. Here the proof is mainly based on Lemma 5.2 in Audibert & Tsybakov (2007).

$$\begin{aligned} \tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*) &= \mathbb{E} \left[|2\tilde{\eta}(\mathbf{X}) - 1| \cdot I(\tilde{s}_{nn}(\mathbf{X}) \neq f^*(\mathbf{X})) \right] \\ &\leq 2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - 1/2| \cdot I(\tilde{s}_{nn}(\mathbf{X}) \neq f^*(\mathbf{X})) \cdot I(|\tilde{\eta}(\mathbf{X}) - 1/2| \leq t) \right] \\ &\quad + 2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - 1/2| \cdot I(\tilde{s}_{nn}(\mathbf{X}) \neq f^*(\mathbf{X})) \cdot I(|\tilde{\eta}(\mathbf{X}) - 1/2| > t) \right] \\ &\leq 2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})| \cdot I(|\tilde{\eta}(\mathbf{X}) - 1/2| \leq t) \right] \\ &\quad + 2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})| \cdot I(|\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})| > t) \right], \end{aligned} \quad (21)$$

where the last inequality follows from the fact that $|\tilde{\eta}(\mathbf{X}) - 1/2| \leq |\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})|$, when $\tilde{s}_{nn}(\mathbf{X}) \neq f^*(\mathbf{X})$. Next, by Cauchy–Schwarz inequality, (21) can be further bounded as

$$\begin{aligned} &2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})| \cdot I(|\tilde{\eta}(\mathbf{X}) - 1/2| \leq t) \right] + 2\mathbb{E} \left[|\tilde{\eta}(\mathbf{X}) - \tilde{f}_{nn}(\mathbf{X})| \cdot I(|\tilde{\eta}(\mathbf{X}) - 1/2| > t) \right] \\ &\leq 2\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})} \sqrt{\mathbb{P}(|\tilde{\eta}(\mathbf{X}) - 1/2| \leq t)} + 2\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2/t \\ &\leq 2^{1-\gamma/2} \|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})} c\kappa_\epsilon^{-\gamma/2} t^{\gamma/2} + 2\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2/t. \end{aligned}$$

Choosing $t = \|\tilde{f}_{nn} - \tilde{\eta}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^{\frac{2}{\gamma+2}} \kappa_\epsilon^{\frac{\gamma}{\gamma+2}}$, we get

$$\tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*) \lesssim \kappa_\epsilon^{-\frac{\gamma}{\gamma+2}} \|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^{\frac{2\gamma+2}{\gamma+2}}.$$

Subsequently, by Lemma 2, it follows that

$$R(\tilde{s}_{nn}) - R(f^*) = \kappa_\epsilon^{-1} \left(\tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*) \right) \lesssim \left(\kappa_\epsilon^{-2} \|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \right)^{\frac{\gamma+1}{\gamma+2}}. \quad (22)$$

Next, we proceed to establish the convergence rate of $\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2$. Let $\mathcal{F}_{n,\delta_n}^{NN} = \{f \in \mathcal{F}_n^{NN} : \|\tilde{\eta} - f\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n\}$. For any $\delta_n > 0$,

$$\mathbb{P} \left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n \right) \leq \mathbb{P} \left(\inf_{f \in \mathcal{F}_{n,\delta_n}^{NN}} n^{-1} \sum_{i=1}^n (f_{nn}^*(\mathbf{x}_i) - \tilde{z}_i)^2 - n^{-1} \sum_{i=1}^n (f(\mathbf{x}_i) - \tilde{z}_i)^2 \geq \delta_n \right),$$

where $f_{nn}^* = \arg \min_{f \in \mathcal{F}_n^{NN}} \|f - \tilde{\eta}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2$.

For ease of notation, we denote that $U_n(f) = n^{-1} \sum_{i=1}^n (f(\mathbf{x}_i) - \tilde{z}_i)^2$ and $U(f) = \mathbb{E}(f(\mathbf{X}) - \tilde{Z})^2$. Then, we have

$$\mathbb{P} \left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n \right) \leq \mathbb{P} \left(\inf_{f \in \mathcal{F}_{n,\delta_n}^{NN}} U_n(f_{nn}^*) - U_n(f) \geq 0 \right).$$

Notice that $\mathcal{F}_{n,\delta_n}^{NN}$ admits the decomposition as $\mathcal{F}_{n,\delta_n}^{NN} = \cup_{i=1}^n \mathcal{H}_i$ with $\mathcal{H}_i = \{f \in \mathcal{F}_n^{NN} : 2^{i-1}\delta_n \leq \|\tilde{\eta} - f\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \leq 2^i\delta_n\}$. Therefore, we further have

$$\mathbb{P}\left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n\right) \leq \sum_{i=1}^n \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} U_n(f_{nn}^*) - U_n(f) \geq 0\right).$$

Clearly, it suffices to bound $\mathbb{P}\left(\inf_{f \in \mathcal{H}_i} U_n(f_{nn}^*) - U_n(f) \geq 0\right)$ for upper bounding $\mathbb{P}\left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n\right)$. For $i \geq 1$,

$$\begin{aligned} & \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} U_n(f_{nn}^*) - U_n(f) \geq 0\right) \\ & \leq \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} [U_n(f_{nn}^*) - U(f_{nn}^*)] - [U_n(f) - U(f)] \geq \inf_{f \in \mathcal{H}_i} U(f) - U(f_{nn}^*)\right) \\ & = \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} [U_n(f_{nn}^*) - U(f_{nn}^*)] - [U_n(f) - U(f)] \geq \inf_{f \in \mathcal{H}_i} U(f) - U(\tilde{\eta}) + U(\tilde{\eta}) - U(f_{nn}^*)\right) \\ & \leq \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} [U_n(f_{nn}^*) - U(f_{nn}^*)] - [U_n(f) - U(f)] \geq 2^{i-1}\delta_n - \|f_{nn}^* - \tilde{\eta}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2\right). \end{aligned}$$

Assuming $\|f_{nn}^* - \tilde{\eta}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \leq \delta_n/2$ yields that

$$\mathbb{P}\left(\inf_{f \in \mathcal{H}_i} U_n(f_{nn}^*) - U_n(f) \geq 0\right) \leq \mathbb{P}\left(\inf_{f \in \mathcal{H}_i} [U_n(f_{nn}^*) - U(f_{nn}^*)] - [U_n(f) - U(f)] \geq 2^{i-2}\delta_n\right).$$

Denote that $M_i = 2^{i-2}\delta_n$. We turn to establish the relation between the variance of $(f_{nn}^*(\mathbf{X}) - \tilde{Z})^2 - (f(\mathbf{X}) - \tilde{Z})^2$ and M_i .

$$\begin{aligned} & \sup_{f \in \mathcal{H}_i} \text{Var}\left[(f_{nn}^*(\mathbf{X}) - \tilde{Z})^2 - (f(\mathbf{X}) - \tilde{Z})^2\right] \\ & = \sup_{f \in \mathcal{H}_i} \text{Var}\left[(f_{nn}^*(\mathbf{X}) - f(\mathbf{X}))(f_{nn}^*(\mathbf{X}) + f(\mathbf{X}) - 2\tilde{Z})\right] \\ & \leq \sup_{f \in \mathcal{H}_i} \mathbb{E}\left[(f_{nn}^*(\mathbf{X}) - f(\mathbf{X}))^2 (f_{nn}^*(\mathbf{X}) + f(\mathbf{X}) - 2\tilde{Z})^2\right] \leq \sup_{f \in \mathcal{H}_i} 4V_n^2 \|f_{nn}^* - f\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \\ & \leq \sup_{f \in \mathcal{H}_i} 8V_n^2 \left(\|f_{nn}^* - \tilde{\eta}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 + \|\tilde{\eta} - f\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2\right) \leq 64V_n^2 M_i \equiv V_i. \end{aligned}$$

In the following, we proceed to verify conditions (4.5)-(4.7) in Shen & Wong (1994). First, the relation between M_i and V_i directly implies (4.6) with $T = 32V_n^2$ and $\epsilon = 1/2$. Second, by Lemma 5 of Schmidt-Hieber (2020),

$$\begin{aligned} & \log \mathcal{N}\left(\epsilon, \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n), \|\cdot\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}\right) \leq \\ & 2L_n(P_n + 1) \log\left(\epsilon^{-1}(L_n + 1)(N_n + 1) \max\{B_n, 1\}\right). \end{aligned}$$

It then follows that

$$\begin{aligned} & \int_{\frac{\epsilon}{32}M_i}^{V_i^{1/2}} \sqrt{\log \mathcal{N}\left(\epsilon, \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n), \|\cdot\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}\right)} d\epsilon / M_i \\ & \leq \int_{\frac{\epsilon}{32}M_i}^{V_i^{1/2}} \sqrt{2L_n(P_n + 1) \log\left(\epsilon^{-1}(L_n + 1)(N_n + 1) \max\{B_n, 1\}\right)} d\epsilon / M_i. \end{aligned} \quad (23)$$

Notice that the right-hand side of (23) is non-increasing in i and M_i , it then follows that

$$\begin{aligned} & \int_{\frac{\epsilon}{32}M_i}^{V_i^{1/2}} \sqrt{2L_n(P_n + 1) \log\left(\epsilon^{-1}(L_n + 1)(N_n + 1) \max\{B_n, 1\}\right)} d\epsilon / M_i \\ & \leq \int_{\frac{\epsilon}{32}M_1}^{V_1^{1/2}} \sqrt{2L_n(P_n + 1) \log\left(\epsilon^{-1}(L_n + 1)(N_n + 1) \max\{B_n, 1\}\right)} d\epsilon / M_1. \end{aligned}$$

With this, condition (4.7) can be satisfied by imposing

$$\int_{\frac{\epsilon}{32}M_1}^{V_1^{1/2}} \sqrt{2L_n(P_n + 1) \log \left(\epsilon^{-1}(L_n + 1)(N_n + 1) \max\{B_n, 1\} \right)} d\epsilon/M_1 \lesssim n^{1/2}, \quad (24)$$

which directly implies the condition (4.5) by appropriate choices of L_n and S_n . By Theorem 3 in Shen & Wong (1994), we get

$$\mathbb{P} \left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n \right) \lesssim \sum_{i=1}^{\infty} \exp \left(- \frac{nM_i^2}{512V_n^2M_i + 2M_i/3} \right) \lesssim \sum_{i=1}^{\infty} \exp \left(- n2^{i-2}\delta_n \right).$$

By the fact $2^{i-2} \geq (i-1/2)$ for $i \geq 1$, there exists some constants C such that

$$\mathbb{P} \left(\|\tilde{\eta} - \tilde{f}_{nn}\|_{L^2(\mathbb{P}_{\mathbf{X}})}^2 \geq \delta_n \right) \lesssim \exp(-Cn\delta_n),$$

provided that $n\delta = o(1)$.

Next, we proceed to consider the approximation error of neural network to meet the assumption that $\|f_{nn}^* - \eta\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}^2 \leq \delta_n$. Notice that $\eta(\mathbf{X}) \in \mathcal{H}(\beta, [0, 1]^p, M)$. By Theorem 5 of Schmidt-Hieber (2020), there exists a class of neural networks $\mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)$ such that for any $0 < \psi_n < 1$

$$\inf_{f \in \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)} \|f - \eta\|_{L^\infty(\mathbb{P}_{\mathbf{X}})} \leq \kappa_\epsilon^{-1} \psi_n,$$

where $P_n \asymp (\kappa_\epsilon^{-1} \psi_n)^{-p/\beta} \log(\kappa_\epsilon/\psi_n)$, $N_n \asymp (\kappa_\epsilon^{-1} \psi_n)^{-p/\beta}$, $B_n = 1$, $V_n \geq M + 1$ and $L_n \asymp \log(\kappa_\epsilon/\psi_n)$. Let η_{nn}^* denote the optimal function in $\mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, \infty)$ to approximate η . Suppose that η_{nn}^* is a L -layer neural network and formulated as

$$\eta_{nn}^*(\mathbf{x}) = \mathbf{A}_{L+1} \mathbf{g}_L(\mathbf{x}) + \mathbf{b}_{L+1},$$

where $\mathbf{g}_L(\mathbf{x}) = \mathbf{h}_L \circ \mathbf{h}_{L-1} \circ \dots \circ \mathbf{h}_1(\mathbf{x})$. We construct a new neural network that $\tilde{\eta}_{nn}$ as

$$\begin{aligned} \tilde{\eta}_{nn}(\mathbf{x}) &= \theta \eta_{nn}^*(\mathbf{x}) + (1 - \theta)(1 - \eta_{nn}^*(\mathbf{x})) \\ &= \theta \mathbf{A}_{L+1} \mathbf{g}_L(\mathbf{x}) + \theta \mathbf{b}_{L+1} + (1 - \theta)(-\mathbf{A}_{L+1} \mathbf{g}_L(\mathbf{x}) + (1 - \mathbf{b}_{L+1})) \\ &= (2\theta - 1) \mathbf{A}_{L+1} \mathbf{g}_L(\mathbf{x}) + (2\theta - 1) \mathbf{b}_{L+1} + (1 - \theta) \mathbf{1}_{L+1}. \end{aligned}$$

It can be easily verified that $\tilde{\eta}_{nn} \in \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)$. This along with the fact that $\tilde{\eta}(\mathbf{x}) = \theta \eta(\mathbf{x}) + (1 - \theta)(1 - \eta(\mathbf{x}))$ with $B_n \geq 1$ results in

$$\begin{aligned} \|\tilde{\eta}_{nn} - \tilde{\eta}\|_{L^\infty(\mathbb{P}_{\mathbf{X}})} &= \|\theta \eta_{nn}^* + (1 - \theta)(1 - \eta_{nn}^*) - \theta \eta - (1 - \theta)(1 - \eta)\|_{L^\infty(\mathbb{P}_{\mathbf{X}})} \\ &= (2\theta - 1) \|\eta_{nn}^* - \eta\|_{L^\infty(\mathbb{P}_{\mathbf{X}})} = \kappa_\epsilon \|\eta_{nn}^* - \eta\|_{L^\infty(\mathbb{P}_{\mathbf{X}})} \leq \psi_n. \end{aligned}$$

Therefore,

$$\|f_{nn}^* - \tilde{\eta}\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}^2 = \inf_{f \in \mathcal{F}_n^{NN}(L_n, N_n, P_n, B_n, V_n)} \|f - \tilde{\eta}\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}^2 \leq \|\tilde{\eta}_{nn} - \tilde{\eta}\|_{L^\infty(\mathbb{P}_{\mathbf{X}})}^2 \leq \psi_n^2.$$

Plugging $P_n \asymp (\kappa_\epsilon^{-1} \psi_n)^{-p/\beta} \log(\kappa_\epsilon/\psi_n)$, $N_n \asymp (\kappa_\epsilon^{-1} \psi_n)^{-p/\beta}$, and $L_n \asymp \log(\kappa_\epsilon/\psi_n)$ into (24) yields that $\kappa_\epsilon^{p/(2\beta)} \psi_n^{-p/(2\beta)} \log(\psi_n^{-1}) \lesssim (n\delta_n)^{1/2}$. Combining this with the assumption that $\psi_n^2 \lesssim \delta_n$, it follows that $\delta_n \asymp \psi_n^2 \asymp \kappa_\epsilon^{2p/(2\beta+p)} (\log n/n)^{2\beta/(2\beta+p)}$. Plugging this into (22) yields that

$$\mathbb{E} \left[R(\tilde{s}_{nn}) - R(f^*) \right] \lesssim \left(\frac{\log n}{n\kappa_\epsilon^2} \right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}}. \quad (25)$$

Notice that the proof of (25) is independent of the distribution of $\mathbf{X}$, therefore (25) holds for any distribution in $\mathcal{P}_{\gamma, \beta}$, which implies that

$$\sup_{\pi \in \mathcal{P}_{\gamma, \beta}} \mathbb{E} \left[R(\tilde{s}_{nn}) - R(f^*) \right] \lesssim \left(\frac{\log n}{n\kappa_\epsilon^2} \right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}}.$$

Proof of the lower bound. The proof is based on the well-known Assouad's lemma. The overall proof for the lower bound is mainly based on the proofs Theorem 3.5 and 4.1 in Audibert & Tsybakov (2007). We first introduce the partition $\{\mathcal{X}_i\}_{i=0}^m$ of the cube $[0, 1]^p$ using the grid $G_q \subseteq [0, 1]^p$ defined by

$$G_q = \left\{ \left(\frac{2k_1+1}{2q}, \dots, \frac{2k_p+1}{2q} \right) : k_i \in \{0, \dots, q-1\}, i \in \{1, \dots, p\} \right\},$$

where $q \geq 1$ is an integer. For any $\mathbf{x} \in \mathbb{R}^p$, let $n_q(\mathbf{x}) \in G_q$ be the unique point which is the closest point to $\mathbf{x} \in \mathbb{R}^p$ among all points in G_q . Without loss of generality, we assume the uniqueness of $n_q(\mathbf{x})$ by choosing the one closest to 0. We define a partition $\{\mathcal{X}'_i\}_{i=1}^{q^p}$ of $\mathbb{R}^p$ as follows. For $\mathbf{x}, \mathbf{y} \in \mathbb{R}^p$, $\mathbf{x}$ and $\mathbf{y}$ are in the same cell $\mathcal{X}'_i$ if and only if $n_q(\mathbf{x}) = n_q(\mathbf{y})$. Then, for $m \leq q^p$, we define $\mathcal{X}_i$ as $\mathcal{X}_i = \mathcal{X}'_i$ for $1 \leq i \leq m$ and $\mathcal{X}_0 = \mathbb{R}^p / \cup_{i=1}^m \mathcal{X}_i$.

Let $u : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be a non-increasing infinitely differentiable function such that $u = 1$ on $[0, 1/4]$ and $u = 0$ on $[1/2, \infty)$. Let $\phi := C_\phi u(\|\mathbf{x}\|)$, where $C_\phi \leq 1$ is taken small enough such that $\phi \in \mathcal{H}(\beta, [0, 1]^p, M)$. For a given $\boldsymbol{\sigma} = (\sigma_1, \dots, \sigma_m) \in \{\pm 1\}^m$, we construct a distribution $\pi_{\boldsymbol{\sigma}}$ on $\mathbb{R}^p \times \{-1, 1\}$ as follows. Let μ be the Lebesgue measure. For $0 < \omega < \frac{1}{m}$ and $A_0 \subseteq \mathcal{X}_0$ with $\mu(A_0) > 0$, we construct the marginal distribution $\mathbb{P}_{\mathbf{X}}$ on $\mathbb{R}^p$ that has the density function

$$\mathbb{P}_{\mathbf{X}}(\mathbf{x}) = \begin{cases} \frac{\omega}{\mu(\mathcal{B}(\mathbf{0}, q^{-1}/4))}, & \mathbf{x} \in \mathcal{B}(z, q^{-1}/4) \text{ for some } z \in G_q, \\ (1 - m\omega)/\mu(A_0), & \mathbf{x} \in A_0, \\ 0, & \text{otherwise.} \end{cases}$$

The conditional distribution on $\{-1, 1\}$ is defined by

$$\eta_{\boldsymbol{\sigma}}(\mathbf{x}) = \mathbb{P}(Y = 1 | \mathbf{X} = \mathbf{x}) = \begin{cases} \frac{1 + \sigma_j \phi(\mathbf{x})}{2}, & \text{for } \mathbf{x} \in \mathcal{X}_j, j = 1, \dots, m, \\ 1/2, & \mathbf{x} \in \mathcal{X}_0, \end{cases}$$

where $\phi(\mathbf{x}) = q^{-\beta} \phi(q(\mathbf{x} - n_q(\mathbf{x})))$. Correspondingly, for any $\boldsymbol{\sigma}$, $\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x})$ is given as

$$\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x}) = \theta \tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x}) + (1 - \theta)(1 - \tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x})) = \begin{cases} \frac{1 + (2\theta - 1)\sigma_j \phi(\mathbf{x})}{2}, & \text{for } \mathbf{x} \in \mathcal{X}_j, j = 1, \dots, m, \\ 1/2, & \mathbf{x} \in \mathcal{X}_0, \end{cases}$$

Notice that $D^s \varphi = q^{|s|-\beta} D^s(q(\mathbf{x} - n_q(\mathbf{x})))$ for any $s \in \mathbb{N}$ with $|s| \leq \beta$. Thus, $\eta_{\boldsymbol{\sigma}}(\mathbf{x})$ belongs to $\mathcal{H}(\beta, [0, 1]^p, M)$.

$$\begin{aligned} \mathbb{P}(|\eta_{\boldsymbol{\sigma}}(\mathbf{x}) - 1/2| \leq t) &= m \mathbb{P}(\phi(q(\mathbf{x} - n_q(\mathbf{x}))) \leq 2tq^\beta) \\ &= m \int_{\mathcal{B}(\mathbf{x}_0, (4q)^{-1})} I(\phi(q(\mathbf{x} - \mathbf{x}_0)) \leq 2tq^\beta) \frac{w}{\mu(\mathcal{B}(\mathbf{0}, (4q)^{-1}))} d\mathbf{x} \\ &= m \int_{\mathcal{B}(\mathbf{0}, 1/4)} I(\phi(\mathbf{x}) \leq 2tq^\beta) \frac{w}{\mu(\mathcal{B}(\mathbf{0}, 1/4))} d\mathbf{x} = mw I(t \geq C_\phi/(2q^\beta)), \end{aligned}$$

where $\mathbf{x}_0 = (\frac{1}{2K}, \dots, \frac{1}{2K})$. Clearly, the low-noise assumption of $\eta_{\boldsymbol{\sigma}}$ can be satisfied by setting $mw \leq C_\phi/(2q^\beta)^\gamma$. Let $\mathcal{P}_{\gamma, \beta}$ denote the set of joint distributions of $(\mathbf{X}, Y)$ satisfying the low-noise assumption and $\eta(\mathbf{x}) \in \mathcal{H}(\beta, [0, 1]^p, M)$. For any $\boldsymbol{\sigma}$, we have $\pi_{\boldsymbol{\sigma}} \in \mathcal{P}_{\gamma, \beta}$, implying $\mathcal{P}' = \{\pi_{\boldsymbol{\sigma}} : \boldsymbol{\sigma} \in \{-1, 1\}^m\} \subset \mathcal{P}_{\gamma, \beta}$. Therefore,

$$\sup_{\pi \in \mathcal{P}_{\gamma, \beta}} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] \geq \sup_{\pi_{\boldsymbol{\sigma}} \in \mathcal{P}'} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)].$$

Let $\tilde{\mathcal{P}}'$ be a set of probability measures on $(\mathbf{X}, \tilde{Y})$ satisfying that for each $\pi_{\boldsymbol{\sigma}} \in \mathcal{P}'$ there exists an $\tilde{\pi}_{\boldsymbol{\sigma}} \in \tilde{\mathcal{P}}'$ such that $\pi_{\boldsymbol{\sigma}}$ and $\tilde{\pi}_{\boldsymbol{\sigma}}$ have the same marginal distribution of $\mathbf{X}$ and $\tilde{\eta}_{\boldsymbol{\sigma}}(\mathbf{x}) = \theta \eta_{\boldsymbol{\sigma}}(\mathbf{x}) + (1 - \theta)(1 - \eta_{\boldsymbol{\sigma}}(\mathbf{x}))$. It follows that

$$\sup_{\pi_{\boldsymbol{\sigma}} \in \mathcal{P}'} \mathbb{E}_{\tilde{\mathcal{D}}} [R(\tilde{s}_{nn}) - R(f^*)] \geq (2\theta - 1)^{-1} \sup_{\tilde{\pi}_{\boldsymbol{\sigma}} \in \tilde{\mathcal{P}}'} \mathbb{E}_{\tilde{\mathcal{D}}} [\tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*)].$$

Next, we proceed to bound $\sup_{\pi_{\sigma} \in \tilde{\mathcal{P}}'} \mathbb{E}_{\tilde{\mathcal{D}}} [\tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*)]$. Notice that f^* varies with the value of σ , therefore we use f_{σ}^* to characterize its dependence on σ .

$$\begin{aligned} \sup_{\pi_{\sigma} \in \tilde{\mathcal{P}}'} \mathbb{E}_{\tilde{\mathcal{D}}} [\tilde{R}(\tilde{s}_{nn}) - \tilde{R}(f^*)] &= \sup_{\pi_{\sigma} \in \tilde{\mathcal{P}}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[|2\tilde{\eta}_{\sigma}(\mathbf{X}) - 1| \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\} \\ &= (2\theta - 1) \sup_{\pi_{\sigma} \in \tilde{\mathcal{P}}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\}, \end{aligned}$$

where $\mathbb{1}_{\{x\}} = 1$ if the statement x is true.

Let Π be the distribution of a Rademacher random variable σ , that is, $\Pi(\sigma = 1) = \Pi(\sigma = -1) = 1/2$. Then

$$\sup_{\pi_{\sigma} \in \tilde{\mathcal{P}}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq f_{\sigma}^*(\mathbf{X})\}} \right] \right] \right\} \geq \mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq f_{\sigma}^*(\mathbf{X})\}} \right] \right] \right\}.$$

Note that for $\mathbf{x} \in \mathcal{X}_0$, $\|\mathbf{x} - n_q(\mathbf{x})\| \geq (2q)^{-1}$ and $\varphi(\mathbf{x}) = 0$. Thus, we have

$$\mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq f_{\sigma}^*(\mathbf{X})\}} \right] \right] \right\} = \sum_{j=1}^m \mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\}.$$

Let $\sigma_{j,r} = (\sigma_1, \dots, \sigma_{j-1}, r, \sigma_{j+1}, \dots, \sigma_m)$ for $r \in \{0, \pm 1\}$. Here $\sigma_{j,r}$ denotes a vector deduced from σ by fixing its j -th element to r . We have

$$\begin{aligned} &\mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\} \\ &= \mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\pi_{\sigma_{j,0}}^n} \left[\frac{d\tilde{\pi}_{\sigma}^n}{d\tilde{\pi}_{\sigma_{j,0}}^n} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\} \\ &= \mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\pi_{\sigma_{j,0}}^n} \mathbb{E}_{\sigma_j \sim \Pi} \left[\frac{d\tilde{\pi}_{\sigma}^n}{d\tilde{\pi}_{\sigma_{j,0}}^n} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\}. \end{aligned}$$

where $\tilde{\pi}_{\sigma_{j,0}}$ has the same marginal distribution as $\mathbb{P}_{\mathbf{X}}$ and $\tilde{\pi}_{\sigma_{j,0}}$ and $\tilde{\pi}_{\sigma}$ differ in the conditional distribution over the points in $\mathcal{X}_j$. Specifically,

$$\tilde{\eta}_{\sigma_{j,0}}(\mathbf{x}) = \begin{cases} 1/2, & \text{if } \mathbf{x} \in \mathcal{X}_j, \\ \tilde{\eta}_{\sigma}(\mathbf{x}), & \text{otherwise.} \end{cases}$$

It then follows that

$$\begin{aligned} &\mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\pi_{\sigma_{j,0}}^n} \mathbb{E}_{\sigma_j \sim \Pi} \left[\frac{d\tilde{\pi}_{\sigma}^n}{d\tilde{\pi}_{\sigma_{j,0}}^n} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\} \\ &= \mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\pi_{\sigma_{j,0}}^n} \mathbb{E}_{\sigma_j \sim \Pi} \left[\frac{d\tilde{\pi}_{\sigma}^n}{d\tilde{\pi}_{\sigma_{j,0}}^n} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\mathbf{X} \in \mathcal{X}_j\}} \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j\}} \right] \right] \right\} \\ &= \mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\mathbf{X} \in \mathcal{X}_j\}} \mathbb{E}_{\pi_{\sigma_{j,0}}^n} \mathbb{E}_{\sigma_j \sim \Pi} \left[\frac{d\tilde{\pi}_{\sigma}^n}{d\tilde{\pi}_{\sigma_{j,0}}^n} \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j\}} \right] \right] \right\} \\ &\geq \mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\mathbf{X} \in \mathcal{X}_j\}} \left(\frac{1}{2} \mathbb{P}_{\pi_{\sigma_{j,1}}^n}(\tilde{s}_{nn}(\mathbf{X}) \neq 1) + \frac{1}{2} \mathbb{P}_{\pi_{\sigma_{j,-1}}^n}(\tilde{s}_{nn}(\mathbf{X}) \neq -1) \right) \right] \right\} \\ &\geq \frac{1}{2} \mathbb{E}_{\Pi^{m-1}} \left\{ \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\mathbf{X} \in \mathcal{X}_j\}} \right] \left(1 - \text{TV}(\tilde{\pi}_{\sigma_{j,1}}^n, \tilde{\pi}_{\sigma_{j,-1}}^n) \right) \right\}, \end{aligned}$$

where TV is the total variation distance between two distributions. Since $\mathbb{P}_{\mathbf{X}}(\mathcal{X}_j) = \omega$, we have

$$\mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{\{\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j\}} \right] \right] \right\} \geq \frac{\omega}{2} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} (\varphi(\mathbf{X}) | \mathbf{X} \in \mathcal{X}_j) \left(1 - \text{TV}(\tilde{\pi}_{\sigma_{j,1}}^n, \tilde{\pi}_{\sigma_{j,-1}}^n) \right),$$

Notice that the above inequality holds for any index $j \in [m]$. In conclusion, we obtain

$$\sum_{j=1}^m \mathbb{E}_{\Pi^m} \left\{ \mathbb{E}_{\tilde{\pi}_{\sigma}^n} \left[\mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left[\varphi(\mathbf{X}) \mathbb{1}_{[\tilde{s}_{nn}(\mathbf{X}) \neq \sigma_j; \mathbf{X} \in \mathcal{X}_j]} \right] \right] \right\} \geq \frac{m\omega}{2} \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} (\varphi(\mathbf{X}) | \mathbf{X} \in \mathcal{X}_j) \left(1 - \text{TV}(\tilde{\pi}_{\sigma_{1,1}}^n, \tilde{\pi}_{\sigma_{1,-1}}^n) \right).$$

Now we bound $\text{TV}(\tilde{\pi}_{\sigma_{1,1}}^n, \tilde{\pi}_{\sigma_{1,-1}}^n)$. First, it holds

$$\text{TV}(\tilde{\pi}_{\sigma_{1,1}}^n, \tilde{\pi}_{\sigma_{1,-1}}^n) = \sum_{l=1}^n \binom{n}{l} \omega^l (1-\omega)^{n-l} \mathcal{V}_l,$$

where $\mathcal{V}_l = \text{TV}(\tilde{\pi}_{-1}^l, \tilde{\pi}_1^l)$ with $\tilde{\pi}_r = \tilde{\pi}_{\sigma_{1,r}}(\cdot | \mathbf{X} \in \mathcal{X}_1)$. Note that

$$\mathcal{V}_l \leq H(\tilde{\pi}_{-1}^l, \tilde{\pi}_1^l) = \sqrt{2 \left(1 - \left[1 - \frac{H^2(\tilde{\pi}_{-1}, \tilde{\pi}_1)}{2} \right]^l \right)}$$

where H is the Hellinger distance and

$$1 - \frac{H^2(\tilde{\pi}_{-1}, \tilde{\pi}_1)}{2} = \mathbb{E}_{\mathbf{X} \sim \mathbb{P}_{\mathbf{X}}} \left(\sqrt{1 - (2\theta - 1)^2 \varphi(\mathbf{X})} | \mathbf{X} \in \mathcal{X}_1 \right) =: \sqrt{1 - b^2}$$

Since $1 - (1 - x^2)^{l/2} \leq \frac{lx^2}{2}$ for $l \geq 2$ and $x > 0$, we have $\mathcal{V}_l \leq b\sqrt{l}$ and

$$\text{TV}(\tilde{\pi}_{\sigma_{1,1}}^n, \tilde{\pi}_{\sigma_{1,-1}}^n) \leq b \sum_{l=1}^n \mathbb{P} \left(\sum_{i=1}^n \epsilon_{i,\omega} = l \right) \sqrt{l} \leq b\sqrt{n\omega},$$

where ϵ_i are i.i.d. random variables such that $\mathbb{P}(\epsilon_i = 1) = \omega = 1 - \mathbb{P}(\epsilon_i = -1)$. In conclusion, we have

$$\sup_{\tilde{\pi}_{\sigma} \in \tilde{\mathcal{P}}'} \left\{ \mathbb{E}_{\tilde{\mathcal{D}}} \left[\tilde{R}(\tilde{s}_{nn}) \right] - \tilde{R}^* \right\} \geq m\omega b' (1 - b\sqrt{n\omega}),$$

where

$$\begin{aligned} b &= \left[1 - \left(\mathbb{E} \left[\sqrt{1 - (2\theta - 1)^2 \varphi^2(X)} | \mathbf{X} \in \mathcal{X}_j \right] \right)^2 \right]^{1/2} \asymp (2\theta - 1)q^{-\beta}, \\ b' &= \mathbb{E} \left((2\theta - 1)\varphi(\mathbf{X}) | \mathbf{X} \in \mathcal{X}_j \right) \asymp (2\theta - 1)q^{-\beta}. \end{aligned}$$

As a result, we have

$$\sup_{\pi \in \mathcal{P}_{\gamma,\beta}} \mathbb{E}_{\tilde{\mathcal{D}}} \left[R(\tilde{s}_{nn}) - R^* \right] \gtrsim m\omega q^{-\beta} (1 - (2\theta - 1)q^{-\beta} \sqrt{n\omega}).$$

Take $\omega = \frac{q^{2\beta}}{4n(2\theta-1)^2}$ and $m = q^p$, we obtain

$$\sup_{\pi \in \mathcal{P}_{\gamma,\beta}} \mathbb{E}_{\tilde{\mathcal{D}}} \left[R(\tilde{s}_{nn}) - R^* \right] \gtrsim \frac{q^{\beta+p}}{n\kappa_{\epsilon}^2} \asymp \left(\frac{1}{n\kappa_{\epsilon}^2} \right)^{\frac{(\gamma+1)\beta}{(\gamma+2)\beta+p}}$$

by taking $q \asymp (n(2\theta-1)^2)^{\frac{1}{(\gamma+2)\beta+p}}$. Particularly, when $\epsilon \asymp n^{-1/2+\zeta}$, we have

$$\sup_{\pi \in \mathcal{P}_{\gamma,\beta}} \mathbb{E}_{\tilde{\mathcal{D}}} \left[R(\tilde{s}_{nn}) - R^* \right] \lesssim \left(\frac{\log n}{n\kappa_{\epsilon}^2} \right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}} \asymp \left(\frac{\log(n)}{n^{2\zeta}} \right)^{\frac{2\beta(\gamma+1)}{2\beta(\gamma+2)+p(\gamma+2)}} \rightarrow 0 \text{ as } n \rightarrow \infty.$$

This completes the proof. $\square$