

Excess risk bounds in robust empirical risk minimization

TIMOTHÉE MATHIEU

*Laboratoire de Mathématiques d'Orsay, Univ. Paris-Sud, CNRS, Université Paris-Saclay,
91405 Orsay, France
timothee.mathieu@u-psud.fr*

STANISLAV MINSKER*,

Department of Mathematics, University of Southern California, Los Angeles, CA 90089.

*Corresponding author: minsker@usc.edu

[Received on 26 January 2021]

This paper investigates robust versions of the general empirical risk minimization algorithm, one of the core techniques underlying modern statistical methods. Success of the empirical risk minimization is based on the fact that for a “well-behaved” stochastic process $\{f(X), f \in \mathcal{F}\}$ indexed by a class of functions $f \in \mathcal{F}$, averages $\frac{1}{N} \sum_{j=1}^N f(X_j)$ evaluated over a sample $X_1, \dots, X_N$ of i.i.d. copies of X provide good approximation to the expectations $\mathbb{E}f(X)$, uniformly over large classes $f \in \mathcal{F}$. However, this might no longer be true if the marginal distributions of the process are heavy-tailed or if the sample contains outliers. We propose a version of empirical risk minimization based on the idea of replacing sample averages by robust proxies of the expectations, and obtain high-confidence bounds for the excess risk of resulting estimators. In particular, we show that the excess risk of robust estimators can converge to 0 at fast rates with respect to the sample size N , referring to the rates faster than $N^{-1/2}$. We discuss implications of the main results to the linear and logistic regression problems, and evaluate the numerical performance of proposed methods on simulated and real data.

Keywords: Keywords: robust estimation, excess risk, median-of-means, regression, classification

1. Introduction

This work is devoted to robust algorithms in the framework of statistical learning. A recent Forbes article [53] states that “Machine learning algorithms are very dependent on accurate, clean, and well-labeled training data to learn from so that they can produce accurate results” and “According to a recent report from AI research and advisory firm Cognilytica, over 80% of the time spent in AI projects are spent dealing with and wrangling data.” While some abnormal elements of the sample, or outliers, can be detected and filtered during the preprocessing steps, others are more difficult to detect: for instance, a sophisticated adversary might try to “poison” data to force a desired outcome [42]. Other seemingly abnormal observations could be inherent to the underlying data-generating process. An “ideal” learning method should not discard informative samples, while limiting the effect of individual observation on the output of the learning algorithm at the same time. We are interested in robust methods that are model-free, and require minimal assumptions on the underlying distribution. We study two types of robustness: robustness to heavy tails expressed in terms of the moment requirements, as well as robustness to (a variant of) adversarial contamination. Heavy tails can be used to model variation and randomness naturally occurring in the sample, while adversarial contamination is a convenient way to model outliers of unknown nature.

The statistical framework used throughout the paper is defined as follows. Let $(S, \mathcal{S})$ be a measurable space, and let $X \in S$ be a random variable with distribution P . Suppose that $X_1, \dots, X_N$ are i.i.d. copies of X . Moreover, assume that $\mathcal{F}$ is a class of measurable functions from S to $\mathbb{R}$ and $\ell : \mathbb{R} \rightarrow \mathbb{R}_+$, where $\mathbb{R}_+$ is a set of non-negative integers, is a loss function. Many problems in statistical learning theory can be formulated as risk minimization of the form

$$\mathbb{E} \ell(f(X)) \rightarrow \min_{f \in \mathcal{F}}.$$

We will frequently write $P\ell(f)$ or simply $\mathcal{L}(f)$ in place of the expected loss $\mathbb{E} \ell(f(X))$. Throughout the paper, we will also assume that the minimum above is attained for some (unique) $f_* \in \mathcal{F}$ (however, f_* does not necessarily coincide with the global minimizer of $\mathcal{L}(f)$ over all measurable functions that might not belong to $\mathcal{F}$). For example, in the context of regression, $X = (Z, Y) \in \mathbb{R}^d \times \mathbb{R}$, $f(Z, Y) = Y - g(Z)$ for some g in a class $\mathcal{G}$ (such as the class of linear functions), $\ell(x) = x^2$, and $f_*(z, y) = y - g_*(z)$, where $g_* = \operatorname{argmin}_{g \in \mathcal{G}} \mathbb{E}(Y - g(Z))^2$. As the true distribution P is usually unknown, a proxy of f_* is obtained via *empirical risk minimization* (ERM), namely

$$\tilde{f}_N := \operatorname{argmin}_{f \in \mathcal{F}} \mathcal{L}_N(f), \quad (1.1)$$

where P_N is the empirical distribution based on the sample $X_1, \dots, X_N$ and

$$\mathcal{L}_N(f) := P_N \ell(f) = \frac{1}{N} \sum_{j=1}^N \ell(f(X_j)).$$

Performance of any $f \in \mathcal{F}$ (in particular, $\tilde{f}_N$) is measured via the excess risk $\mathcal{E}(f) := P\ell(f) - P\ell(f_*)$. The excess risk of $\tilde{f}_N$ is a random variable defined as

$$\mathcal{E}(\tilde{f}_N) := P\ell(\tilde{f}_N) - P\ell(f_*) = \mathbb{E}[\ell(\tilde{f}_N(X)) | X_1, \dots, X_N] - \mathbb{E}\ell(f_*(X)).$$

General bounds for the excess risk have been extensively studied; a small subsample of the relevant works includes the papers [57, 58, 31, 4, 7, 55] and references therein. However, until recently sharp estimates were known only in the situation when the functions in the class $\ell(\mathcal{F}) := \{\ell(f), f \in \mathcal{F}\}$ are uniformly bounded, or when the envelope $F_\ell(x) := \sup_{f \in \mathcal{F}} |\ell(f(x))|$ of the class $\ell(\mathcal{F})$ possesses finite exponential moments. Our focus is on the situation when marginal distributions of the process $\{\ell(f(X)), f \in \mathcal{F}\}$ indexed by $\mathcal{F}$ are allowed to be heavy-tailed, meaning that they possess finite moments of low order only (in this paper, “low order” usually means between 2 to 4). In such cases, the tail probabilities of the random variables $\left\{ \frac{1}{\sqrt{N}} \sum_{j=1}^N (\ell(f(X_j)) - \mathbb{E}\ell(f(X))), f \in \mathcal{F} \right\}$ decay polynomially, thus rendering many existing techniques ineffective. Moreover, we consider a challenging framework of *adversarial contamination* where the initial dataset of cardinality N is merged with a set of $\mathcal{O} < N$ outliers which are generated by an adversary who has an opportunity to inspect the data, and the combined dataset of cardinality $N^\circ = N + \mathcal{O}$ is presented to an algorithm; in this paper, we assume that the proportion of contamination $\frac{\mathcal{O}}{N}$ (or its upper bound) is known.

The approach that we propose is based on replacing the sample mean at the core of ERM by a more “robust” estimator of $\mathbb{E} \ell(f(X))$ that exhibits tight concentration under minimal moment assumptions. Well known examples of such estimators include the median-of-means estimator [48, 2, 38] and Catoni’s estimator [14]. Both the median-of-means and Catoni’s estimators gain robustness at the cost of being biased. The ways that the bias of these estimators is controlled is based on different principles however.

Informally speaking, Catoni’s estimator relies on delicate “truncation” of the data, while the median-of-means (MOM) estimator exploits the fact that the median and the mean of a symmetric distribution both coincide with its center of symmetry. In this paper, we will use “hybrid” estimators that take advantage of both symmetry and truncation. This family of estimators has been introduced and studied in [46, 47], and we review the construction below.

1.1 Organization of the paper.

The main ideas behind the proposed estimators are explained in Section 1.3, followed by the high-level overview of the main theoretical results and comparison to existing literature in Section 1.4. The complete statements of the key results are given in Section 2, and in Section 3 we deduce the corollaries of these results for specific examples. Finally, the main ideas and key inequalities necessary for the proofs is explained in Section 4. The remaining technical arguments are contained in the Supplementary material [41]. Finally, in Section C of the Supplement we discuss practical implementation and numerical performance of our methods on synthetic and real data.

1.2 Notation.

For two sequences $\{a_j\}_{j \geq 1} \subset \mathbb{R}$ and $\{b_j\}_{j \geq 1} \subset \mathbb{R}$ for $j \in \mathbb{N}$, the expression $a_j \lesssim b_j$ means that there exists a constant $c > 0$ such that $a_j \leq cb_j$ for all $j \in \mathbb{N}$; $a_j \asymp b_j$ means that $a_j \lesssim b_j$ and $b_j \lesssim a_j$. Absolute constants will be denoted c, c_1, C, C' , etc, and may take different values in different parts of the paper. For a function $h : \mathbb{R}^d \mapsto \mathbb{R}$, we define

$$\operatorname{argmin}_{y \in \mathbb{R}^d} h(y) = \{y \in \mathbb{R}^d : h(y) \leq h(x) \text{ for all } x \in \mathbb{R}^d\},$$

and $\|h\|_\infty := \operatorname{ess\,sup}\{|h(y)| : y \in \mathbb{R}^d\}$. Moreover, $L(h)$ will stand for a Lipschitz constant of h . For $f \in \mathcal{F}$, let $\sigma^2(\ell, f) = \operatorname{Var}(\ell(f(X)))$ and for any subset $\mathcal{F}' \subseteq \mathcal{F}$, denote $\sigma^2(\ell, \mathcal{F}') = \sup_{f \in \mathcal{F}'} \sigma^2(\ell, f)$. Additional notation and auxiliary results are introduced on demand.

1.3 Robust mean estimators.

Let $k \leq N$ be an integer, and assume that $G_1, \dots, G_k$ are disjoint subsets of the index set $\{1, \dots, N\}$ of cardinality $|G_j| = n \geq \lfloor N/k \rfloor$ each. Given $f \in \mathcal{F}$, let

$$\overline{\mathcal{L}}_j(f) := \frac{1}{n} \sum_{i \in G_j} \ell(f(X_i))$$

be the empirical mean evaluated over the subsample indexed by G_j . Given a convex, even function $\rho : \mathbb{R} \mapsto \mathbb{R}_+$ and $\Delta > 0$, set

$$\widehat{\mathcal{L}}^{(k)}(f) := \operatorname{argmin}_{y \in \mathbb{R}} \sum_{j=1}^k \rho \left(\sqrt{n} \frac{\overline{\mathcal{L}}_j(f) - y}{\Delta} \right). \quad (1.2)$$

Clearly, if $\rho(x) = x^2$, $\widehat{\mathcal{L}}^{(k)}(f)$ is equal to the sample mean. If $\rho(x) = |x|$, then $\widehat{\mathcal{L}}^{(k)}(f)$ is the median-of-means estimator [48, 2, 19]. We will be interested in the situation when ρ is smooth and “shaped” like Huber’s loss, in particular, that ρ' is bounded and Lipschitz continuous (exact conditions imposed

on ρ are specified in Assumption 1 below). Note that (1.2) defines a whole family of estimators for different values of k and n . It is instructive to consider two cases: first, when $k = N$ (so that $n = 1$) and the scaling factor $\Delta \asymp \sqrt{\text{Var}(\ell(f(X)))\sqrt{N}}$, $\widehat{\mathcal{L}}^{(k)}(f)$ is akin to Catoni's estimator [14], and when n is large (e.g. $\sqrt{N} \ll n \ll N$ and $\Delta \asymp \sqrt{\text{Var}(\ell(f(X)))}$), $\widehat{\mathcal{L}}^{(k)}(f)$ is the “median-of-means type” estimator. Let us elaborate on these two cases further. Generally speaking, the estimator $\widehat{\mathcal{L}}^{(k)}(f)$ is biased, and, as we already mentioned in the introduction, one way to understand the difference between Catoni's and median-of-means type estimators is via the difference in mechanisms used to control the bias. In the case of Catoni's estimator, this mechanism is based on truncating each observation at the level of order $\sqrt{N}$ ¹ encoded in the choice of $\Delta \asymp \sqrt{\text{Var}(\ell(f(X)))\sqrt{N}}$, while in the case of the median-of-means estimator it relies on the approximate symmetry, implied by the Central Limit Theorem, of the distribution of the empirical averages $\overline{\mathcal{L}}_j(f)$, and in particular the fact that any reasonable estimator of location for this distribution will be close to the mean $\mathcal{L}(f)$ when $|G_j|$ is large. In Section 2.1, we formally introduce the key quantities that allow us to control the bias under various moment assumptions on the underlying classes.

We also construct a permutation-invariant version of the estimator $\widehat{\mathcal{L}}^{(k)}(f)$ that does not depend on the specific choice of the subgroups $G_1, \dots, G_k$. We conjecture that this estimator is more efficient than $\widehat{\mathcal{L}}^{(k)}(f)$; see remark 1.1 below for more details. Next, let

$$\mathcal{A}_N^{(n)} := \{J : J \subseteq \{1, \dots, N\}, |J| = n\}.$$

Let h be a measurable, permutation-invariant function of n variables. Recall that a U-statistic of order n with kernel h based on an i.i.d. sample $X_1, \dots, X_N$ is defined as [25]

$$U_{N,n} = \frac{1}{\binom{N}{n}} \sum_{J \in \mathcal{A}_N^{(n)}} h(\{X_j\}_{j \in J}). \quad (1.3)$$

Given $J \in \mathcal{A}_N^{(n)}$, let $\overline{\mathcal{L}}(f; J) := \frac{1}{n} \sum_{i \in J} f(X_i)$. Consider U-statistics of the form

$$U_{N,n}(z; f) = \sum_{J \in \mathcal{A}_N^{(n)}} \rho \left(\sqrt{n} \frac{\overline{\mathcal{L}}(f; J) - z}{\Delta} \right).$$

Then the permutation-invariant version of $\widehat{\mathcal{L}}^{(k)}(f)$ is defined as

$$\widehat{\mathcal{L}}_U^{(k)}(f) := \underset{z \in \mathbb{R}}{\operatorname{argmin}} U_{N,n}(z; f).$$

Finally, assuming that $\widehat{\mathcal{L}}^{(k)}(f)$ provides good approximation of the expected loss $\mathcal{L}(f)$ of each individual $f \in \mathcal{F}$, it is natural to consider

$$\widehat{f}_N := \underset{f \in \mathcal{F}}{\operatorname{argmin}} \widehat{\mathcal{L}}^{(k)}(f), \quad (1.4)$$

as well as its permutation-invariant analogue

$$\widehat{f}_N^U := \underset{f \in \mathcal{F}}{\operatorname{argmin}} \widehat{\mathcal{L}}_U^{(k)}(f) \quad (1.5)$$

¹Reference to truncation can be made explicit by setting $\rho(x) = \min(x^2/2, |x| - 1/2)$ to be Huber's loss and considering the gradient descent iteration for the optimization problem (1.2).

as an alternative to standard empirical risk minimization (1.1). The main goal of this paper is to obtain general bounds for the excess risk of the estimators $\hat{f}_N$ and $\hat{f}_N^U$ under minimal assumptions on the stochastic process $\{\ell(f(X)), f \in \mathcal{F}\}$. More specifically, we are interested in scenarios when the excess risk converges to 0 at fast, or “optimistic” rates, referring to the rates faster than $N^{-1/2}$. Rate of order $N^{-1/2}$ (“slow rates”) are easier to establish: in particular, results of this type follow from bounds on the uniform deviations $\sup_{f \in \mathcal{F}} |\widehat{\mathcal{L}}^{(k)}(f) - \mathcal{L}(f)|$ that have been investigated in [47]. Proving fast rates is a more technically challenging task: to achieve the goal, we develop Bahadur-type representations [6] of the estimators $\widehat{\mathcal{L}}^{(k)}(f)$ and $\widehat{\mathcal{L}}_U^{(k)}(f)$ that provide linear, in $\ell(f)$, approximations of these nonlinear statistics that are easier to study, and carefully analyze the remainder terms. Introduction of such representations in the framework of median-of-means estimation is one of the main technical novelties of the paper; the tools we develop could prove useful in other related problems, such as study of the asymptotic distributions of the robust estimators $\hat{f}_N$ and $\hat{f}_N^U$.

REMARK 1.1 The main reason we introduce the permutation-invariant estimator $\hat{f}_N^U$ is our conjecture that it has superior, compared to $\hat{f}_N$, performance. We were able to confirm this fact numerically in our experiments; however, complete theoretical confirmation is not yet available, and requires new technical tools beyond those developed in the present work. Specifically, we conjecture that $\hat{f}_N^U$ is more *efficient* than $\hat{f}_N$: when $\mathcal{F}$ is finite dimensional, this means, informally, that the asymptotic distribution of $\sqrt{N}(\hat{f}_N^U - f_*)$ has smaller variance than the asymptotic distribution of $\sqrt{N}(\hat{f}_N - f_*)$. In other words, the conjectured difference in performance is about the constant factors rather than the rates. Such improvements are too subtle to be captured by the non-asymptotic bounds for the excess risk that are being pursued in this work, nevertheless they are clearly noticeable in the simulations.

It should also be acknowledged that exact evaluation of the U-statistics-based estimators $\widehat{\mathcal{L}}_U^{(k)}(f)$ and $\hat{f}_N^U$ is not feasible due to the number of summands $\binom{N}{n}$ being very large even for small values of n . However, exact computation is typically not required, and throughout our detailed simulation studies, gradient descent methods proved to be very efficient for the problem (1.5) in scenarios like least-squares and logistic regression. These points, as well as comparison of the numerical performance of the estimators $\hat{f}_N^U$ and $\hat{f}_N$, are further discussed in Section C of the Supplementary material [41].

1.4 Overview of the main results and comparison to existing bounds.

Our main contribution is the proof of high-confidence bounds for the excess risk of the estimators $\hat{f}_N$ and $\hat{f}_N^U$. First, we show (see Theorem 2.1 and (2.4)) that the excess risk is bounded from above by the quantity of order $N^{-1/2}$ (referred to as “slow rates”) with exponentially high probability if

$$\sigma^2(\ell, \mathcal{F}) = \sup_{f \in \mathcal{F}} \sigma^2(\ell, f) < \infty \text{ and } \mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{\sqrt{N}} \sum_{j=1}^N (\ell(f(X_j)) - \mathbb{E}\ell(f(X))) < \infty.$$

The latter is true if the class $\{\ell(f), f \in \mathcal{F}\}$ is P-Donsker [24], in other words, if the empirical process $f \mapsto \frac{1}{\sqrt{N}} \sum_{j=1}^N (\ell(f(X_j)) - \mathbb{E}\ell(f(X)))$ converges weakly to a Gaussian limit. This result is analogous to its counterpart in the standard empirical risk minimization framework. Moreover, it is known [43, 36] that in general, the $N^{-1/2}$ rate for the excess risk of the empirical risk minimizers can not be improved. Our main contribution is the proof of the fact that that under additional assumption requiring that any $f \in \mathcal{F}$ with small excess risk is itself close to f_* (that minimizes the expected loss), $\hat{f}_N$ and $\hat{f}_N^U$ attain *fast rates*. This fact is well-known in the usual empirical risk minimization framework [10, 31] but is

new for the type of robust estimators considered here. We state the bounds below only for $\widehat{f}_N$ while the results for the U-statistics based $\widehat{f}_N^U$ are similar, up to the change in absolute constants. In order to avoid excessive technical details at this stage, we will first illustrate our general results by stating corollaries for the popular frameworks of logistic regression and regression with quadratic loss, while the most general versions of the theorems and additional examples will be stated afterwards.

Binary classification and logistic regression. Assume that $(Z, Y) \in S \times \{\pm 1\}$ is a random couple where Z is an instance and Y is a binary label, and let $g_*(z) := \mathbb{E}[Y|Z = z]$ be the regression function. It is well-known that the binary classifier $b_*(z) := \text{sign}(g_*(z))$ achieves smallest possible misclassification error defined as $P(Y \neq g(Z))$. Let $\mathcal{F}$ be a given convex class of functions mapping S to $\mathbb{R}$, $\ell : \mathbb{R} \mapsto \mathbb{R}_+$ a convex, nondecreasing, Lipschitz loss function, and let

$$h_* = \underset{\text{all measurable } f}{\operatorname{argmin}} \mathbb{E}\ell(Yf(Z)).$$

The loss ℓ is classification-calibrated if $\text{sign}(h_*(z)) = b_*(z)$ P-almost surely; we refer the reader to [9] for a detailed exposition. In the case of logistic regression considered below, $S = \mathbb{R}^d$,

$$\ell(y, f(z)) = \ell(yf(z)) := \log(1 + e^{-yf(z)})$$

is the classification-calibrated loss and $\mathcal{F} = \{f_\beta(\cdot) = \langle \cdot, v \rangle, v \in \mathbb{R}^d, \|v\|_2 \leq R\}$. Note that results stated below hold without assuming that $h_* \in \mathcal{F}$.

Regression with quadratic loss. Let $(Z, Y) \in S \times \mathbb{R}$ be a random couple satisfying $Y = f_*(Z) + \eta$ where the noise variable η is independent of Z and $f_*(z) = \mathbb{E}[Y|Z = z]$ is the regression function. Linear regression with quadratic loss corresponds to $S = \mathbb{R}^d$,

$$\ell(y, f(z)) = \ell(y - f(z)) := (y - f(z))^2$$

and $\mathcal{F} = \{f_\beta(\cdot) = \langle \cdot, v \rangle, v \in \mathbb{R}^d, \|v\|_2 \leq R\}$. In this case, we will assume that $f_* \in \mathcal{F}$; it is possible to avoid this assumption at the cost of additional technicalities and taking advantage of the deep results of S. Mendelson [45] on the multiplier inequalities.

In the statements below, we will assume that we are given an i.i.d. sample $(Z_1, Y_1), \dots, (Z_N, Y_N)$ having the same distribution as (Z, Y) where the marginal distribution of Z is supported on a compact set. Moreover, suppose that $\mathbb{E}|\eta|^8 < \infty$ in the case of regression with quadratic loss; Section 3 contains other examples covering a wider class of distributions and classes $\mathcal{F}$.

THEOREM 1.1 (Informal) Assume the framework of either logistic regression or linear regression with quadratic loss. Then, for appropriately chosen k and Δ ,

$$\mathcal{E}(\widehat{f}_N) \leq C(R, P, \rho) \left(\frac{d}{N} + \frac{s}{N^{3/4}} + \left(\frac{\mathcal{O}}{N} \right)^{3/4} \right)$$

with probability at least $1 - e^{-s}$ for all $s \lesssim k$.

Moreover, we construct a two-step estimator $\widehat{f}_N''$ based on $\widehat{f}_N$ that is capable of achieving further improved rates.

THEOREM 1.2 (Informal) Assume the framework of either logistic regression or linear regression with quadratic loss. There exists an estimator $\widehat{f}_N''$, defined later in the paper, such that

$$\mathcal{E}(\widehat{f}_N'') \leq C(R, P, \rho) \left(\frac{d}{N} + \frac{s}{N} + \frac{\mathcal{O}}{N} \right)$$

with probability at least $1 - e^{-s}$ for all $1 \leq s \leq s_{\max}$ where $s_{\max} := s_{\max}(N) \rightarrow \infty$ as $N \rightarrow \infty$.

The estimator $\hat{f}_N''$ mentioned in Theorem 1.2 is based on a two-step procedure, where $\hat{f}_N$ serves as an initial approximation that is refined on the second step via risk minimization restricted to a “small neighborhood” of $\hat{f}_N$. All of the bounds in this paper have the form $\mathcal{E}(\hat{f}_N) \leq \bar{\delta} + C(\mathcal{F}, P) \left(\frac{s}{N^\gamma} + \left(\frac{\rho}{N} \right)^\gamma \right)$, where $\frac{1}{2} \leq \gamma \leq 1$ and $\bar{\delta}$ is the quantity (formally defined in (2.5)) that often coincides, up to log-factors, with the optimal rate for the excess risk [3, 40] – for instance, $\bar{\delta} \asymp \frac{d}{N}$ in the examples above. In the standard empirical risk minimization, the excess risk bounds in the linear and logistic regression admit the bounds of order $\frac{d}{N} + \frac{s}{N}$, albeit under more restrictive assumptions and in the corruption-free framework. Therefore, the bound of Theorem 1.1 is suboptimal in these cases due to the “remainder terms” being of order $N^{-3/4}$, and the improvement achieved by the two-step estimator $\hat{f}_N''$, as described in Theorem 1.2, becomes important.

Next, we provide a brief overview of the literature on the topic and compare our results to the state of the art. Robustness of statistical learning algorithms has been studied extensively in recent years. Existing research has mainly focused on addressing robustness to heavy tails as well as adversarial contamination. One line of work investigated robust versions of the gradient descent method for the optimization problem (1.1) based on variants of the multivariate median-of-means technique [51, 16, 59, 1], as well as Catoni’s estimator [28]. The line works initiated in the theoretical computer science community [33, 20, 22, also see the survey paper [23]] tackled the problem of optimal mean estimation in the adversarial contamination framework by establishing deep connections between the mean and covariance estimation problems that culminated in the family of powerful filtering algorithms; these algorithms can also be used as subroutines in robust gradient descent-type methods [21, 17]. While these algorithms admit strong theoretical guarantees, they require robustly estimating the gradient vector at every step (with the exception of [21] that offers a more efficient approach) hence are computationally demanding; moreover, results are weaker for losses that are not strongly convex (for instance, the hinge loss). The line of research that is closest in spirit to the approach of this paper includes the works that employ robust risk estimators based on Catoni’s idea [5, 13, 27] and the median-of-means technique, such as “tournaments” and the “min-max median-of-means” [40, 39, 34, 35, 18], also see [17, 29] for the computationally efficient algorithms related to the tournament-type procedures. As it was mentioned in the introduction, the core of our methods can be viewed as a “hybrid” between Catoni’s and the median-of-means estimators. We provide a more detailed comparison to the results of the aforementioned papers:

1. We show that risk minimization based on a version of Catoni’s estimator is capable of achieving fast rates, thus improving the results and weakening the assumptions stated in [13] that only allowed the slow rates to be established;
2. We develop new tools and techniques to analyze proposed estimators. In particular, we do not rely on the “small ball” method [32, 44] and the standard “majority vote-based” analysis [34, 40] of the median-of-means estimators. Instead, we provide accurate bounds for the bias and investigate the remainder terms for the Bahadur-type linear approximations of the estimators defined in (1.2). In particular, we demonstrate that the order of typical deviations of the estimator $\hat{\mathcal{P}}^{(k)}(f)$ around $\mathcal{L}(f)$ are significantly smaller than the deviations of the subsample averages $\bar{\mathcal{L}}_j(f)$, which is not easy to do using the majority vote-based proof techniques; consequently, this fact allows us to “decouple” the confidence parameter s that controls the deviation probabilities from parameters k and ρ responsible for the number of subsamples and the degree of contamination respectively. Unlike the tournaments-based estimators, in some regimes our algorithms admit a

“universal” choice of k that is independent of the parameter $\bar{\delta}$ controlling the optimal rate. In the previous works, parameter k was often overloaded as it controlled the deviation probabilities while depending on $\bar{\delta}$ (or a closely related quantity) at the same time. Finally, our techniques allow us to establish bounds that are uniform over a certain range of confidence parameter s while the previously existing deviation results were only available for $s \asymp k$.

3. We are able to simultaneously treat the case of Lipschitz as well as non-Lipschitz (e.g., quadratic) loss functions ℓ . At the same time, in some situations (e.g. linear regression with quadratic loss), the required assumptions are stronger compared to the best results in the literature tailored specifically to the task, e.g. [34, 40] that treat the case of regression with quadratic loss.
4. Existing approaches based on the median-of-means estimators are either computationally intractable [40], or outputs of practically efficient algorithms do not admit strong theoretical guarantees [34, 35, 18]. We design numerical algorithms specifically for the estimators $\hat{f}_N$ and $\tilde{f}_N^U$ defined via (1.4) and (1.5), and show that they enjoy good performance in numerical experiments as well as strong theoretical guarantees.

2. Theoretical guarantees for the excess risk.

In this section, we give complete statements of the main results and explain the high-level ideas behind their proofs.

2.1 Preliminaries.

We start by introducing the main quantities that appear in our results, and state the key assumptions. Recall that $\sigma^2(\ell, \mathcal{F}')$ stands for $\sup_{f \in \mathcal{F}'} \sigma^2(\ell, f)$, where $\mathcal{F}' \subseteq \mathcal{F}$. The loss functions ρ that will be of interest to us satisfy the following assumption.

Assumption 1 Suppose that the function $\rho : \mathbb{R} \mapsto \mathbb{R}$ is convex, even, 5 times continuously differentiable and such that

- (i) $\rho'(z) = z$ for $|z| \leq 1$ and $\rho'(z) = \text{const}$ for $z \geq 2$,
- (ii) $z - \rho'(z)$ is nondecreasing.

An example of a function ρ satisfying required assumptions is given by “smoothed” Huber’s loss defined as follows. Let

$$H(y) = \frac{y^2}{2} I\{|y| \leq 3/2\} + \frac{3}{2} \left(|y| - \frac{3}{4} \right) I\{|y| > 3/2\}$$

be the usual Huber’s loss. Moreover, let ϕ be the “bump function” $\phi(x) = C \exp\left(-\frac{4}{1-4x^2}\right) I\{|x| \leq \frac{1}{2}\}$ where C is chosen so that $\int_{\mathbb{R}} \phi(x) dx = 1$. Then ρ given by the convolution $\rho(x) = (H * \phi)(x)$ satisfies Assumption 1.

REMARK 2.1 (a) The requirements that ρ is 5 times continuously differentiable is of the technical nature and is likely not necessary. It appears due to the fact that we need to control higher order terms in the Bahadur-Kiefer type representations of the estimator $\widehat{\mathcal{L}}^{(k)}(f)$, as well as rely on the Lindeberg replacement-type arguments in our proofs.

- (b) The derivative ρ' has a natural interpretation of being a smooth version of the truncation function. Moreover, observe that $\rho'(2) - 2 \leq \rho'(1) - 1 = 0$ by (ii), hence $\|\rho'\|_{\infty} \leq 2$. It is also easy to see

that for any $x > y$, $\rho'(x) - \rho'(y) = y - \rho'(y) - (x - \rho'(x)) + x - y \leq x - y$, hence ρ' is Lipschitz continuous with Lipschitz constant $L(\rho') = 1$.

In Section 1.3, we have briefly discussed the bias of robust mean estimators and various ways that it can be controlled. Now we will introduce the key quantities necessary to make the bounds precise. Everywhere below, $\Phi(\cdot)$ stands for the cumulative distribution function of the standard normal random variable and $W(f)$ denotes a random variable with distribution $N(0, \sigma^2(f))$. For $f \in \mathcal{F}$ such that $\sigma(f) > 0$, $n \in \mathbb{N}$ and $t > 0$, define

$$\mathcal{M}_f(t, n) := \left| \Pr \left(\frac{\sum_{j=1}^n (f(X_j) - Pf)}{\sigma(f)\sqrt{n}} \leq t \right) - \Phi(t) \right|,$$

where $Pf := \mathbb{E}f(X)$. In other words, $\mathcal{M}_f(t, n)$ controls the rate of convergence in the central limit theorem. It follows from the results of L. Chen and Q.-M. Shao (Theorem 2.2 in [15]) that

$$\begin{aligned} \mathcal{M}_f(t, n) \leq g_f(t, n) := C \left(\frac{\mathbb{E}(f(X) - \mathbb{E}f(X))^2 I \left\{ \frac{|f(X) - \mathbb{E}f(X)|}{\sigma(f)\sqrt{n}} > 1 + \left| \frac{t}{\sigma(f)} \right| \right\}}{\sigma^2(f) \left(1 + \left| \frac{t}{\sigma(f)} \right| \right)^2} \right. \\ \left. + \frac{1}{\sqrt{n}} \frac{\mathbb{E}|f(X) - \mathbb{E}f(X)|^3 I \left\{ \frac{|f(X) - \mathbb{E}f(X)|}{\sigma(f)\sqrt{n}} \leq 1 + \left| \frac{t}{\sigma(f)} \right| \right\}}{\sigma^3(f) \left(1 + \left| \frac{t}{\sigma(f)} \right| \right)^3} \right) \end{aligned}$$

given that the absolute constant C is large enough. Note that, crucially, the control of the rate in terms of $g_f(t, n)$ is non-uniform, since $g_f(t, n)$ is a decreasing function of t . Moreover, let

$$G_f(n, \Delta) := \int_0^\infty g_f \left(\Delta \left(\frac{1}{2} + t \right), n \right) dt.$$

The quantity $\frac{G_f(n, \Delta)}{\sqrt{n}}$ plays the key role in controlling the bias of the estimator $\widehat{\mathcal{L}}^{(k)}(f)$: it decreases both as Δ get large and as the subsample size n increases, referring to different bias-controlling mechanisms of Catoni's and the median-of-means type estimators, see discussion after (1.2). The following statement provides simple upper bounds for $g_f(t, n)$ and $G_f(n, \Delta)$ that depend on the tail properties of $f(X)$; its proof can be found in [47, Section 4.4].

LEMMA 2.1 Let $X_1, \dots, X_n$ be i.i.d. copies of X , and assume that $\text{Var}(f(X)) < \infty$. Then $g_f(t, n) \rightarrow 0$ as $|t| \rightarrow \infty$ and $g_f(t, n) \rightarrow 0$ as $n \rightarrow \infty$, with the convergence being monotone. Moreover, if $\mathbb{E}|f(X) - \mathbb{E}f(X)|^{2+\delta} < \infty$ for some $\delta \in [0, 1]$, then for all $t > 0$

$$\begin{aligned} g_f(t, n) &\leq C' \frac{\mathbb{E}|f(X) - \mathbb{E}f(X)|^{2+\delta}}{n^{\delta/2} (\sigma(f) + |t|)^{2+\delta}} \leq C' \frac{\mathbb{E}|f(X) - \mathbb{E}f(X)|^{2+\delta}}{n^{\delta/2} |t|^{2+\delta}}, \\ G_f(n, \Delta) &\leq C'' \frac{\mathbb{E}|f(X) - \mathbb{E}f(X)|^{2+\delta}}{\Delta^{2+\delta} n^{\delta/2}}, \end{aligned} \tag{2.1}$$

where $C', C'' > 0$ are absolute constants.

We can rewrite the bound for $\sup_{f \in \mathcal{F}} G_f(n, \Delta)$ as $\sup_{f \in \mathcal{F}} G_f(n, \Delta) \leq C'' \frac{\sup_{f \in \mathcal{F}} \mathbb{E}(|f(X) - \mathbb{E}f(X)|/\sigma(\ell, \mathcal{F}))^{2+\delta}}{(\Delta/\sigma(\ell, \mathcal{F}))^{2+\delta} n^{\delta/2}}$,

where the numerator $\sup_{f \in \mathcal{F}} \mathbb{E}(|f(X) - \mathbb{E}f(X)|/\sigma(\ell, \mathcal{F}))^{2+\delta}$ is the quantity akin the kurtosis while the

ratio $M_\Delta := \frac{\Delta}{\sigma(\ell, \mathcal{F})}$ appearing in the denominator can be interpreted as a truncation level expressed in the “units” of $\sigma(\ell, \mathcal{F})$. This “truncation level,” along with the subgroup size n , are the two main quantities controlling the bias of the estimators $\widehat{\mathcal{L}}^{(k)}(f)$, $f \in \mathcal{F}$.

2.2 Slow rates for the excess risk.

Let

$$\begin{aligned}\widehat{\delta}_N &:= \mathcal{E}(\widehat{f}_N) = \mathcal{L}(\widehat{f}_N) - \mathcal{L}(f_*), \\ \widehat{\delta}_N^U &:= \mathcal{E}(\widehat{f}_N^U) = \mathcal{L}(\widehat{f}_N^U) - \mathcal{L}(f_*)\end{aligned}$$

be the excess risk of $\widehat{f}_N$ and its permutation-invariant analogue $\widehat{f}_N^U$ which are the main objects of our interest. The following bound for the excess risk is well known in the empirical risk minimization literature [31], and it easily leads to control of the excess risk in terms of the uniform deviations of robust mean estimators.

$$\begin{aligned}\mathcal{E}(\widehat{f}_N) &= \mathcal{L}(\widehat{f}_N) - \mathcal{L}(f_*) \\ &= \mathcal{L}(\widehat{f}_N) + \widehat{\mathcal{L}}^{(k)}(\widehat{f}_N) - \widehat{\mathcal{L}}^{(k)}(\widehat{f}_N) + \widehat{\mathcal{L}}^{(k)}(f_*) - \widehat{\mathcal{L}}^{(k)}(f_*) - \mathcal{L}(f_*) \\ &= \left(\mathcal{L}(\widehat{f}_N) - \widehat{\mathcal{L}}^{(k)}(\widehat{f}_N) \right) - \left(\mathcal{L}(f_*) - \widehat{\mathcal{L}}^{(k)}(f_*) \right) + \underbrace{\widehat{\mathcal{L}}^{(k)}(\widehat{f}_N) - \widehat{\mathcal{L}}^{(k)}(f_*)}_{\leq 0} \\ &\leq 2 \sup_{f \in \mathcal{F}} \left| \widehat{\mathcal{L}}^{(k)}(f) - \mathcal{L}(f) \right|. \quad (2.2)\end{aligned}$$

The first result, Theorem 2.1 below, together with the inequality (2.2) immediately implies the “slow rate bound” (meaning rate not faster than $N^{-1/2}$) for the excess risk. This result has been previously established in [47]. Define

$$\widetilde{\Delta} := \max(\Delta, \sigma(\ell, \mathcal{F})).$$

THEOREM 2.1 There exist absolute constants $c, C > 0$ such that for all $s > 0$, n and k satisfying

$$\frac{1}{\Delta} \left(\frac{1}{\sqrt{k}} \mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{\sqrt{N}} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) + \sigma(\ell, \mathcal{F}) \sqrt{\frac{s}{k}} \right) + \sup_{f \in \mathcal{F}} G_f(n, \Delta) + \frac{s}{k} + \frac{\mathcal{O}}{k} \leq c, \quad (2.3)$$

the following inequality holds with probability at least $1 - 2e^{-s}$:

$$\begin{aligned}\sup_{f \in \mathcal{F}} \left| \widehat{\mathcal{L}}^{(k)}(f) - \mathcal{L}(f) \right| &\leq C \left[\frac{\widetilde{\Delta}}{\Delta} \left(\mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{N} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) + \sigma(\ell, \mathcal{F}) \sqrt{\frac{s}{N}} \right) \right. \\ &\quad \left. + \widetilde{\Delta} \left(\sqrt{n} \frac{s}{N} + \frac{\sup_{f \in \mathcal{F}} G_f(n, \Delta)}{\sqrt{n}} + \frac{\mathcal{O}}{k\sqrt{n}} \right) \right].\end{aligned}$$

Moreover, the same bounds hold for the permutation-invariant estimators $\widehat{\mathcal{L}}_U^{(k)}(f)$, up to the change in absolute constants.

An immediate corollary is the bound for the excess risk

$$\begin{aligned} \mathcal{E}(\widehat{f}_N) \leq C \left[\frac{\widetilde{\Delta}}{\Delta} \left(\mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{N} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) + \sigma(\ell, \mathcal{F}) \sqrt{\frac{s}{N}} \right) \right. \\ \left. + \widetilde{\Delta} \sqrt{n} \left(\frac{s}{N} + \frac{\sup_{f \in \mathcal{F}} G_f(n, \Delta)}{n} + \frac{\mathcal{O}}{N} \right) \right] \quad (2.4) \end{aligned}$$

that holds under the assumptions of Theorem 2.1 with probability at least $1 - 2e^{-s}$. When the class $\{\ell(f), f \in \mathcal{F}\}$ is P-Donsker [24], $\limsup_{N \rightarrow \infty} \left| \mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{\sqrt{N}} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) \right|$ is bounded, hence condition (2.3) holds for N large enough whenever s is not too big and Δ and k are not too small, namely, $s \leq c'k$ and $\Delta \sqrt{k} \geq c''\sigma(\mathcal{F})$. The bound of Theorem 2.1 also suggests that the natural “unit” to measure the magnitude of the parameter Δ is $\sigma(\ell, \mathcal{F})$.

To put these results in perspective, let us consider two examples. First, assume that $n = 1, k = N$ and set $\Delta = \Delta(s) := \sigma(\mathcal{F}) \sqrt{\frac{N}{s}}$ for $s \leq c'N$. Using Lemma 2.1 with $\delta = 0$ to estimate $G_f(n, \Delta)$, we deduce that

$$\mathcal{E}(\widehat{f}_N) \leq C \left[\mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{N} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) + \sigma(\ell, \mathcal{F}) \left(\sqrt{\frac{s}{N}} + \frac{\mathcal{O}}{\sqrt{N}} \right) \right]$$

with probability at least $1 - 2e^{-s}$. This inequality improves upon excess risk bounds obtained for Catoni-type estimators in [13], as it does not require functions in $\mathcal{F}$ to be uniformly bounded.

The second case we consider is when $N \gg n \geq 2$. For the choice of $\Delta \asymp \sigma(\ell, \mathcal{F})$, the estimator $\widehat{\mathcal{P}}^{(k)}(f)$ most closely resembles the median-of-means estimator, as we have explained in Section 1.3. In this case, Theorem 2.1 yields the excess risk bound of the form

$$\mathcal{E}(\widehat{f}_N) \leq C \left[\mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{N} \sum_{j=1}^N (\ell(f(X_j)) - P\ell(f)) + \sigma(\ell, \mathcal{F}) \left(\sqrt{\frac{s}{N}} + \sqrt{\frac{k}{N}} \sup_{f \in \mathcal{F}} G_f(n, \sigma(\mathcal{F})) + \frac{\mathcal{O}}{k} \sqrt{\frac{k}{N}} \right) \right]$$

that holds with probability $\geq 1 - 2e^{-s}$ for all $s \leq c'k$. As $\sup_{f \in \mathcal{F}} G_f(n, \Delta)$ is small for large n and $\frac{\mathcal{O}}{k} \sqrt{\frac{k}{N}} \leq \sqrt{\frac{\mathcal{O}}{N}}$ whenever $\mathcal{O} \leq k$, this bound improves upon Theorem 2 in [35] that provides bounds for the excess risk for robust classifiers based on the the median-of-means estimators.

2.3 Towards fast rates for the excess risk.

It is well known that in regression and binary classification problems, the excess risk often converges to 0 at a rate faster than $N^{-1/2}$, and could be as fast as N^{-1} . Such rates are often referred to as “fast” or “optimistic” rates. In particular, this is the case when there exists a “link” between the excess risk and the variance of the loss class, namely, if for some convex nondecreasing and nonnegative function ϕ such that $\phi(0) = 0$,

$$\mathcal{E}(f) = P\ell(f) - P\ell(f_*) \geq \phi \left(\sqrt{\text{Var}(\ell(f(X)) - \ell(f_*(X)))} \right).$$

It is thus natural to ask if fast rates can be attained by estimators produced by the robust algorithms proposed above. Results presented in this section give an affirmative answer to this question. Let us

introduce the main quantities that commonly appear in the excess risk bounds [31, 40]. For $\delta > 0$, let

$$\begin{aligned}\mathcal{F}(\delta) &:= \{\ell(f) : f \in \mathcal{F}, \mathcal{E}(f) \leq \delta\}, \\ v(\delta) &:= \sup_{\ell(f) \in \mathcal{F}(\delta)} \sqrt{\text{Var}(\ell(f(X)) - \ell(f_*(X)))}, \\ \omega(\delta) &:= \mathbb{E} \sup_{\ell(f) \in \mathcal{F}(\delta)} \left| \frac{1}{\sqrt{N}} \sum_{j=1}^N \left((\ell(f) - \ell(f_*))(X_j) - P(\ell(f) - \ell(f_*)) \right) \right|.\end{aligned}$$

Moreover, define

$$\mathfrak{B}(\ell, \mathcal{F}) := \frac{\sup_{f \in \mathcal{F}} \mathbb{E}^{1/4}(\ell(f(X)) - \mathbb{E}\ell(f(X)))^4}{\sigma(\ell, \mathcal{F})}.$$

The following condition, known as *Bernstein's condition* following [10], plays the crucial role in the analysis of excess risk bounds.

Assumption 2 There exist constants $D > 0$, $\delta_B > 0$ such that

$$\text{Var}(\ell(f(X)) - \ell(f_*(X))) \leq D^2 \mathcal{E}(f)$$

whenever $\mathcal{E}(f) \leq \delta_B$.

Informally speaking, Assumption 2 postulates that any $f \in \mathcal{F}$ (more precisely, the loss $\ell(f)$ induced by it) with small excess risk is itself close to f_* . If this is true, it turns out that one can avoid global bounds on the expected supremum of the empirical process used to obtain “slow” rates, and instead rely on the modulus of continuity $\omega(\delta)$ of the empirical process locally in the neighborhood of $\ell(f_*)$ in order to get better upper bounds on the excess risk. The basics of this approach in the classical empirical risk minimization frameworks are clearly explained in [31, Chapter 1.2], and we rely on similar ideas below.

Assumption 2 is known to hold in many concrete cases of prediction and classification tasks, and we provide examples and references in Section 3 below. More general versions of the Bernstein's condition are often considered in the literature: for instance, it can be replaced by assumption requiring that $\text{Var}(\ell(f(X)) - \ell(f_*(X))) \leq D^2 (\mathcal{E}(f))^\tau$ for some $\tau \in (0, 1]$, as was done in [10]; clearly, our assumption corresponds to $\tau = 1$. Results of this paper admit straightforward extensions to the slightly less restrictive scenario when $\tau < 1$; we omit the details to reduce the level of technical burden on the statements of our results.

Following [31, Chapter 4], we will say that the function $\psi : \mathbb{R}_+ \mapsto \mathbb{R}_+$ is of concave type if it is nondecreasing and $x \mapsto \frac{\psi(x)}{x}$ is decreasing. Moreover, if for some $\gamma \in (0, 1)$ $x \mapsto \frac{\psi(x)}{x^\gamma}$ is decreasing, we will say that ψ is of strictly concave type with exponent γ . We will assume that $\omega(\delta)$ admits an upper bound $\tilde{\omega}(\delta)$ of strictly concave type (with some exponent γ), and that $v(\delta)$ admits an upper bound $\tilde{v}(\delta)$ of concave type. For instance, when Assumption 2 holds, $v(\delta) \leq D\sqrt{\delta}$ for $\delta \leq \delta_B$, implying that $\tilde{v}(\delta) = D\sqrt{\delta}$ is an upper bound for $v(\delta)$ of strictly concave type with $\gamma = \frac{1}{2}$.² Moreover, the function $\omega(\delta)$ often admits an upper bound of the form $\tilde{\omega}(\delta) = R_1 + \sqrt{\delta}R_2$ where R_1 and R_2 do not depend on δ ; such an upper bound is also of concave type. Next, set

$$\bar{\delta} := \min \left\{ \delta > 0 : C_1(\rho) \frac{1}{\sqrt{N}} \frac{\tilde{\Delta}}{\Delta} \frac{\tilde{\omega}(\delta)}{\delta} \leq \frac{1}{7} \right\}, \quad (2.5)$$

²This is only true in some neighborhood of 0, but is sufficient for our purposes.

where $C_1(\rho)$ is a sufficiently large positive constant that depends only on ρ . The quantity $\bar{\delta}$ often coincides with the optimal rates for the excess risk in the classical empirical risk minimization framework: for example, it is of order $\frac{d}{N}$ up to logarithmic factors in linear regression with quadratic loss and in logistic regression when Bernstein's condition is satisfied; in general, the order of $\bar{\delta}$ ranges between the pessimistic $N^{-1/2}$ in “hard” problems and “optimistic” N^{-1} where the rates between correspond to weaker versions of Assumptions 2, for instance, see [9]. The theorems below provide estimates for the excess risk of robust risk minimizers under various conditions on the tails of the random variables $\{f(X), f \in \mathcal{F}\}$. All these bounds have the same structure that includes the term $\bar{\delta}$ as well as the “remainder terms” that account for the bias of the robust risk estimators $\widehat{\mathcal{L}}^{(k)}(f)$ as well as the outlier contamination proportion $\frac{\ell}{N}$; naturally, stricter moment conditions result in better remainder terms.

THEOREM 2.2 Assume that conditions of Theorem 2.1 hold. Additionally, suppose that $M_\Delta := \frac{\Delta}{\sigma(\ell, \mathcal{F})} \geq 1$. Then

$$\widehat{\delta}_N \leq \bar{\delta} + C(\rho) \left(D^2 \left(\frac{1}{M_\Delta^2 n} + \frac{s + \mathcal{O}}{N} \right) + \sigma(\ell, \mathcal{F}) \sqrt{n} M_\Delta \left(\frac{1}{M_\Delta^4 n} + \frac{s + \mathcal{O}}{N} \right) \right).$$

with probability at least $1 - 10e^{-s}$, where the constant $C(\rho)$ depends on ρ only and D is a constant appearing in Assumption 2.

Under stronger moment assumptions, the excess risk bound can be strengthened and take the following form.

THEOREM 2.3 Assume that conditions of Theorem 2.1 hold. Additionally, suppose that

$$\sup_{f \in \mathcal{F}} \mathbb{E}^{1/4} (\ell(f(X)) - \mathbb{E}\ell(f(X)))^4 < \infty$$

and that $M_\Delta := \frac{\Delta}{\sigma(\ell, \mathcal{F})} \geq 1$. Then

$$\widehat{\delta}_N \leq \bar{\delta} + C(\rho) (D^2 + \sigma(\ell, \mathcal{F}) \sqrt{n} M_\Delta) \left(\frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_\Delta^4 n^2} + \frac{s + \mathcal{O}}{N} \right).$$

with probability at least $1 - 10e^{-s}$, where the constant $C(\rho)$ depends on ρ only and D is a constant appearing in Assumption 2.

The main ideas behind the proofs of Theorems 2.2 and 2.3 are explained in the beginning of Section 4.

REMARK 2.2

1. The bounds of Theorems 2.2 and 2.3 hold for the excess risk $\widehat{\delta}_N^U$ of the permutation-invariant estimator $\widehat{f}_N^U$, up to a change in absolute constants.
2. It is evident that whenever $\mathcal{O} = 0$, the best possible rates implied by Theorem 2.2 are of order $N^{-2/3}$ (indeed, this is the case whenever $M_\Delta \sqrt{n} \asymp N^{1/3}$ and $\bar{\delta} \lesssim N^{-2/3}$), while the best possible rates attained by Theorem 2.3 are of order $N^{-3/4}$ (when $M_\Delta \sqrt{n} \asymp N^{1/4}$ and $\bar{\delta} \lesssim N^{-3/4}$); in particular, in this case the choice of M_Δ and n is independent of $\bar{\delta}$. In general, if $\mathcal{O} = \varepsilon N$ for $\varepsilon > 0$, the best rates implied by Theorems 2.2 and 2.3 are $\bar{\delta} + C(\mathcal{F}, \rho, P) \varepsilon^{2/3}$ and $\bar{\delta} + C(\mathcal{F}, \rho, P) \varepsilon^{3/4}$ respectively.
3. Assumption requiring that $M_\Delta \geq 1$ is introduced for convenience: without it, extra powers of the ratio $\frac{\max(\Delta, \sigma(\ell, \mathcal{F}))}{\Delta}$ appear in the bounds.

Our next goal is to describe an estimator that is capable of achieving excess risk rates up to N^{-1} . The approach that we follow is similar in spirit to the “minmax” estimators studied in [5, 38, 34], among

others, as well as the “median-of-means tournaments” introduced in [40]; all these methods focus on estimating the differences $\mathcal{L}(f_1) - \mathcal{L}(f_2)$ for all $f_1, f_2 \in \mathcal{F}$. Recall that $f_* = \operatorname{argmin}_{f \in \mathcal{F}} P\ell(f)$, and observe that for any fixed $f' \in \mathcal{F}$, f_* can be equivalently defined via

$$f_* = \operatorname{argmin}_{f \in \mathcal{F}} P(\ell(f) - \ell(f')).$$

A version of the robust empirical risk minimizer (1.4) corresponding to this problem can be defined as

$$\widehat{\mathcal{L}}^{(k)}(f - f') := \operatorname{argmin}_{y \in \mathbb{R}} \frac{1}{\sqrt{N}} \sum_{j=1}^k \rho \left(\sqrt{n} \frac{(\overline{\mathcal{L}}_j(f) - \overline{\mathcal{L}}_j(f')) - y}{\Delta} \right)$$

for appropriately chosen $\Delta > 0$, and

$$\widehat{f}_N := \operatorname{argmin}_{f \in \mathcal{F}} \widehat{\mathcal{L}}^{(k)}(f - f').$$

Moreover, if $f' \in \mathcal{F}$ is a priori known to be “close” to f_* , then it suffices to search for the minimizer in a neighborhood $\mathcal{F}'$ of f' that contains f_* instead of all $f \in \mathcal{F}$:

$$\widehat{f}_N' := \operatorname{argmin}_{f \in \mathcal{F}'} \widehat{\mathcal{L}}^{(k)}(f - f').$$

The advantage gained by this procedure is expressed by the fact that $\sup_{f \in \mathcal{F}'} \operatorname{Var}(\ell(f(X)) - \ell(f'(X)))$ can be much smaller than $\sigma(\ell, \mathcal{F})$.

We will now formalize this argument and provide performance guarantees; we use the framework of Theorem 2.3 which leads to the bounds that are easier to state and interpret. However, similar reasoning applies to the setting of Theorem 2.2 as well. The presented algorithms also admit straightforward permutation-invariant modifications that we omit. Let

$$\widehat{\mathcal{E}}_N(f) := \widehat{\mathcal{L}}^{(k)}(f) - \widehat{\mathcal{L}}^{(k)}(\widehat{f}_N)$$

be the “empirical excess risk” of f . Indeed, this is a meaningful notion as $\widehat{f}_N$ is the minimizer of $\widehat{\mathcal{L}}^{(k)}(f)$ over $f \in \mathcal{F}$. Assume that the initial sample of size N is split into two disjoint parts S_1 and S_2 of cardinalities that differ by at most 1: $(X_1, Y_1), \dots, (X_N, Y_N) = S_1 \cup S_2$. The algorithm proceeds in the following way:

1. Let $\widehat{f}_{|S_1|}$ be the estimator (1.4) evaluated over subsample S_1 of cardinality $|S_1| \geq \lfloor N/2 \rfloor$, with the scale parameter Δ_1 and the partition parameter k_1 corresponding the group size $n_1 = \lfloor |S_1|/k_1 \rfloor$;
2. Let $\delta' = \overline{\delta} + C(\rho) (D^2 + \sigma(\ell, \mathcal{F}) \sqrt{n} M_{\Delta_1}) \left(\frac{\mathfrak{B}_{\Delta_1}^6(\ell, \mathcal{F})}{M_{\Delta_1}^4 n_1^2} + \frac{s + \mathcal{O}}{N} \right)$ be a known upper bound on the excess risk in Theorem 2.3 (while this condition is restrictive, it is similar to the requirements of existing approaches [13, 40]; discussion of adaptation issues is beyond the scope of this paper and will be addressed elsewhere). Set

$$\widehat{\mathcal{F}}(\delta') := \left\{ f \in \mathcal{F} : \widehat{\mathcal{E}}_N(f) \leq \delta' \right\}.$$

3. Define $\hat{f}_N'' := \operatorname{argmin}_{f \in \widehat{\mathcal{F}}(\delta')} \widehat{\mathcal{L}}^{(k)}(f - \hat{f}_{|S_1|})$ where

$$\widehat{\mathcal{L}}^{(k)}(f - \hat{f}_{|S_1|}) = \operatorname{argmin}_{y \in \mathbb{R}} \sum_{j=1}^{k_2} \rho \left(\sqrt{n} \frac{(\overline{\mathcal{L}}_j(f) - \overline{\mathcal{L}}_j(\hat{f}_{|S_1|})) - y}{\Delta_2} \right)$$

is based on the subsample S_2 of cardinality $|S_2| \geq \lfloor N/2 \rfloor$, a scale parameter Δ_2 and the partition parameter k_2 corresponding the group size $n_2 = \lfloor |S_2|/k_2 \rfloor$.

It will be demonstrated in the course of the proofs that on an event of high probability, $\widehat{\mathcal{F}}(\delta') \subseteq \mathcal{F}(c\delta')$ for an absolute constant $c \leq 7$. Hence, on this event $\sup_{f \in \widehat{\mathcal{F}}(\delta')} \operatorname{Var}(\ell(f(X)) - \ell(f_*(X))) \leq v^2(c\delta') \leq cD^2\delta'$ by the definition of $v(\delta)$ and Assumption 2, thus $\Delta_2 = DM_{\Delta_2}\sqrt{c\delta'}$ with $M_{\Delta_2} \geq 1$ often leads to an estimator with improved performance.

THEOREM 2.4 Suppose that

$$\sup_{f \in \mathcal{F}} \mathbb{E}^{1/4}(\ell(f(X)) - \mathbb{E}\ell(f(X)))^4 < \infty$$

and that Δ_1, Δ_2 satisfy $M_{\Delta_1} := \frac{\Delta_1}{\sigma(\ell, \mathcal{F})} \geq 1$ and $M_{\Delta_2} := \frac{\Delta_2}{D\sqrt{7\delta'}} \geq 1$. Moreover, assume that for a sufficiently small absolute constant $c' > 0$, $\sup_{f \in \mathcal{F}} \max(G_f(n_1, \Delta_1), G_f(n_2, \Delta_2)) \leq c'$ and $\frac{s+\mathcal{O}}{\min(k_1, k_2)} \leq c'$. Finally, we require that

$$\begin{aligned} \sqrt{k_1}M_{\Delta_1} &\geq \frac{c'}{\sigma(\ell, \mathcal{F})} \mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{\sqrt{|S_1|}} \sum_{j=1}^{|S_1|} (\ell(f(X_j)) - P\ell(f)) \quad \text{and} \\ \sqrt{k_2}M_{\Delta_2} &\geq c' \frac{\sqrt{N\delta'}}{D}. \end{aligned} \tag{2.6}$$

Then

$$\mathcal{E}(\hat{f}_N'') \leq \bar{\delta} + C(\rho) \left(D^2 + D\sqrt{\delta'}\sqrt{n}M_{\Delta_2} \right) \left(\frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_{\Delta_2}^4 n^2} + \frac{s+\mathcal{O}}{N} \right)$$

with probability at least $1 - 20e^{-s}$, where $C(\rho)$ depends on ρ only and D is the constant appearing in Assumption 2.

The statement of Theorem 2.4 is technical, so let us try to distill the main ideas. The key difference between Theorem 2.3 and Theorem 2.4 is that the “remainder term”

$$\sigma(\ell, \mathcal{F})\sqrt{n}M_{\Delta} \left(\frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_{\Delta}^4 n^2} + \frac{s+\mathcal{O}}{N} \right)$$

is replaced by a potentially much smaller quantity $\sqrt{\delta'}\sqrt{n}M_{\Delta} \left(\frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_{\Delta}^4 n^2} + \frac{s+\mathcal{O}}{N} \right)$. In particular, if $\delta' \ll (nM_{\Delta}^2)^{-1}$, this term often becomes negligible. To be more specific, assume that $\bar{\delta} = \frac{C(\mathcal{F})}{\sqrt{N}} \cdot h(N)$ where $h(N) \rightarrow 0$ as $N \rightarrow \infty$ (meaning that fast rates are achievable) and that $\mathcal{O} = \varepsilon N$ for $\varepsilon \geq \frac{1}{N}$. Moreover, suppose that $\mathfrak{B}(\ell, \mathcal{F})$ is bounded above by a constant. If Δ_1 is chosen such that $\Delta_1 \asymp \sigma(\ell, \mathcal{F})$, then

$\delta' = C \left(\bar{\delta} + \sigma(\ell, \mathcal{F}) \left(\left(\frac{k}{N} \right)^{3/2} + \frac{s+\mathcal{O}}{\sqrt{kN}} \right) \right)$. Hence, if $\max(h(N)\sqrt{N}, N\epsilon^{2/3}) \ll k_j \leq CN\sqrt{\epsilon}$ for $j = 1, 2$ and $\Delta_2 \asymp \sqrt{\delta'}$, then

$$\delta' \cdot nM_{\Delta_2}^2 = O(1),$$

and the excess risk of $\hat{f}_N''$ admits the bound

$$\mathcal{E}(\hat{f}_N'') \leq \bar{\delta} + C(\rho, D) \left(\epsilon + \frac{s}{N} \right)$$

that holds with probability at least $1 - Ce^{-s}$. A possible choice satisfying all the required conditions is $k_j \asymp N\sqrt{\epsilon}$, $j = 1, 2$ (indeed, in this case it is straightforward to check that conditions (2.6) hold for sufficiently large N as $k_j \gtrsim \sqrt{N}$, $j = 1, 2$). Analysis of the case when $\mathcal{O} = 0$ follows similar steps, with several simplifications.

3. Examples.

In this section, we consider two common prediction problems, regression and binary classification, and discuss the implications of our main results for these problems in detail.

3.1 Binary classification with convex surrogate loss.

The key elements of the binary classification framework were outlined in Section 1.4. Here, we recall few popular examples of classification-calibrated losses and present conditions that are sufficient for the Assumption 2 to hold.

Logistic loss $\ell(yf(z)) = \log(1 + e^{-yf(z)})$. Consider two scenarios:

1. Uniformly bounded classes, meaning that for all $f \in \mathcal{F}$, $\sup_{z \in S} |f(z)| \leq B$. In this case, Assumption 2 holds with $D = 2e^B$ for all $f \in \mathcal{F}$. See [8] and Proposition 6.1 in [3].
2. Linear separators and Gaussian design: in this case, we assume that $S = \mathbb{R}^d$, $Z \sim N(0, I)$ is Gaussian, and $\mathcal{F} = \{\langle \cdot, v \rangle : \|v\|_2 \leq R\}$ is a class of linear functions. In this case, according to the Proposition 6.2 in [3], Bernstein's assumption is satisfied with $D = cR^{3/2}$ for some absolute constant $c > 0$.

Hinge loss $\ell(yf(z)) = \max(0, 1 - yf(z))$. In this case, sufficient condition for Assumption 2 to hold is the following: there exists $\tau > 0$ such that $|g_*(Z)| \geq \tau$ almost surely, where $g_*(z) = \mathbb{E}[Y|Z = z]$. It follows from Theorem 7 in [8] (see also [55]) that Assumption 2 holds with $D = \frac{1}{\sqrt{2\tau}}$ in this case.

Bound for $\bar{\delta}$. Let Π stand for the marginal distribution of Z and recall that

$$\omega(\delta) := \mathbb{E} \sup_{\ell(f) \in \mathcal{F}(\delta)} \left| \frac{1}{\sqrt{N}} \sum_{j=1}^N \left((\ell(Y_j f(Z_j)) - \ell(Y_j f_*(Z_j))) - \mathbb{E}(\ell(Y f(Z)) - \ell(Y f_*(Z))) \right) \right|.$$

Since ℓ is Lipschitz continuous by assumption (with Lipschitz constant denoted $L(\ell)$), consequent application of symmetrization and Talagrand's contraction inequalities [37, 56] yields that

$$\omega(\delta) \leq 4L(\ell) \mathbb{E} \sup_{\|f - f_*\|_{L_2(\Pi)} \leq D\sqrt{\delta}} \left| \frac{1}{\sqrt{N}} \sum_{j=1}^N \epsilon_j (f - f_*)(Z_j) \right|$$

where $\varepsilon_1, \dots, \varepsilon_N$ are i.i.d. random signs independent from Y_j 's and Z_j 's. The latter quantity is the modulus of continuity of a Rademacher process, and various upper bounds for it are well known. For instance, if $\mathcal{F}$ is a subset of a linear space of dimension d , then, according to Proposition 3.2 in [31], $\mathbb{E} \sup_{\|f-f_*\|_{L_2(\Pi)} \leq D\sqrt{\delta}} \left| \frac{1}{\sqrt{N}} \sum_{j=1}^N \varepsilon_j (f-f_*)(Z_j) \right| \leq D\sqrt{\delta}\sqrt{d}$, whence $\tilde{\omega}(\delta) := 4DL(\ell)\sqrt{\delta d}$ is an upper bound for $\omega(\delta)$ and is of concave type, implying that

$$\bar{\delta} \leq C(\rho, \ell) D^2 \frac{d}{N}.$$

More generally, assume that the class $\mathcal{F}$ has a measurable envelope $F(z) := \sup_{f \in \mathcal{F}} |f(z)|$ that satisfies $\|F(Z)\|_{\psi_2} < \infty$, where $\|\xi\|_{\psi_2} := \inf \{C > 0 : \mathbb{E} \exp(|\xi|/C)^2 \leq 2\}$ is the ψ_2 (Orlicz) norm. Moreover, suppose that the covering numbers $N(\mathcal{F}, Q, \varepsilon)$ of the class $\mathcal{F}$ with respect to the norm $L_2(Q)$ ³ satisfy the bound

$$N(\mathcal{F}, Q, \varepsilon) \leq \left(\frac{A\|F\|_{L_2(Q)}}{\varepsilon} \right)^V \quad (3.1)$$

for some constants $A \geq 1$, $V \geq 1$, all $0 < \varepsilon \leq 2\|F\|_{L_2(Q)}$ and all probability measures Q . For instance, VC-subgraph classes are known to satisfy this bound with V being the VC dimension of $\mathcal{F}$ [58, 31]. In this case, it is not difficult to show (see for example the proof of Lemma 3.1 in the Supplementary material [41]) that

$$\mathbb{E} \sup_{\|f-f_*\|_{L_2(\Pi)} \leq D\sqrt{\delta}} \left| \frac{1}{\sqrt{N}} \sum_{j=1}^N \varepsilon_j (f-f_*)(Z_j) \right| \leq \tilde{\omega}(\delta) := C\sqrt{V \log(e^2 A^2 N)} \left(\sqrt{\delta} + \sqrt{\frac{V}{N}} \log(A^2 N) \|F\|_{\psi_2} \right),$$

hence it is easy to check that in this case

$$\bar{\delta} \leq C(\rho) \frac{V \log^{3/2}(e^2 A^2 N) \|F\|_{\psi_2}}{N}.$$

It immediately follows from the discussion following Theorem 2.4 that the excess risk of the estimator $\hat{f}_N''$ satisfies

$$\mathcal{E}(\hat{f}_N'') \leq C(\rho, D) \left(\frac{\rho}{N} + \frac{V \log^{3/2}(e^2 A^2 N) \|F\|_{\psi_2} + s}{N} \right)$$

with probability at least $1 - 20e^{-s}$. Note that we did not need to assume that $h_* := \underset{\text{all measurable } f}{\operatorname{argmin}} \mathbb{E} \ell(Yf(Z))$

belongs to $\mathcal{F}$. Similar results hold for regression problems with Lipschitz losses, such as Huber's loss or quantile loss [3].

3.2 Regression with quadratic loss.

Let $X = (Z, Y) \in S \times \mathbb{R}$ be a random couple with distribution P satisfying $Y = f_*(Z) + \eta$ where the noise variable η is independent of Z and $f_*(z) = \mathbb{E}[Y|Z=z]$ is the regression function. Let $\|\eta\|_{2,1} := \int_0^\infty \sqrt{\Pr(|\eta| > t)} dt$, and observe that $\|\eta\|_{2,1} < \infty$ as $\sup_{f \in \mathcal{F}} \mathbb{E}(Y - f(Z))^4 < \infty$ by assumption. As

³Definition: the covering number $N(\mathcal{F}, Q, \varepsilon)$ is the smallest integer $k \geq 1$ such that there exist $f_1, \dots, f_k \in L_2(Q)$ satisfying $\bigcup_{j=1}^k B(f_j, \varepsilon) \supseteq \mathcal{F}$, where $B(f_j, \varepsilon)$ is the $L_2(Q)$ ball of radius ε centered at f_j .

before, Π will stand for the marginal distribution of Z . Let $\mathcal{F}$ be a given convex class of functions mapping S to $\mathbb{R}$ and such that the regression function f_* belongs to $\mathcal{F}$, so that

$$f_* = \operatorname{argmin}_{f \in \mathcal{F}} \mathbb{E} (Y - f(Z))^2.$$

In this case, the natural choice for the loss function is the quadratic loss $\ell(x) = x^2$ which is not Lipschitz continuous on unbounded domains. Assume that the class $\mathcal{F}$ has a measurable envelope $F(z) := \sup_{f \in \mathcal{F}} |f(z)|$ that satisfies $\|F(Z)\|_{\psi_2} < \infty$. Moreover, suppose that the covering numbers $N(\mathcal{F}, Q, \varepsilon)$ of the class $\mathcal{F}$ with respect to the norm $L_2(Q)$ satisfy the bound

$$N(\mathcal{F}, Q, \varepsilon) \leq \left(\frac{A \|F\|_{L_2(Q)}}{\varepsilon} \right)^V$$

for some constants $A \geq 1$, $V \geq 1$, all $0 < \varepsilon \leq 2\|F\|_{L_2(Q)}$, and all probability measures Q (see remark about VC-subgraph classes following display (3.1)).

Verification of Bernstein's assumption. It follows from Lemma 5.1 in [31] that

$$\mathcal{F}(\delta) \subseteq \{(y - f(z))^2 : f \in \mathcal{F}, \mathbb{E}(f(Z) - f_*(Z))^2 \leq 2\delta\},$$

hence $v(\delta) \leq \sqrt{2\delta}$ so D can be taken to be $\sqrt{2}$ in Assumption 2.

Bound for $\bar{\delta}$. Required estimates follow from the following lemma:

LEMMA 3.1 Under the assumptions made in this section and for $\Delta \geq \sigma(\ell, \mathcal{F})$,

$$\bar{\delta} \leq C(\rho) \frac{V \log^2(A^2 N) (\|F\|_{\psi_2}^2 + \|\eta\|_{2,1}^2)}{N}.$$

Moreover, if the functions in $\mathcal{F}$ are uniformly bounded, the $\log^2(A^2 N)$ can be removed.

The proof is given in Section A.9 of the Supplementary material [41]. An immediate corollary of the lemma, according to the discussion following Theorem 2.4, is that the excess risk of the estimator $\hat{f}_N''$ satisfies the inequality

$$\mathcal{E}(\hat{f}_N'') \leq C(\rho) \left(\frac{\mathcal{O}}{N} + \frac{V \log^2(A^2 N) (\|F\|_{\psi_2}^2 + \|\eta\|_{2,1}^2) + s}{N} \right)$$

with probability at least $1 - 20e^{-s}$, for $0 < s \leq cN^{1/4}$.

4. Proofs of the main results.

In the proofs of the main results, we will rely on the following convenient change of variables. Denote

$$\begin{aligned} \hat{G}_k(z; f) &= \frac{1}{\sqrt{k}} \sum_{j=1}^k \rho' \left(\sqrt{n} \frac{(\bar{\mathcal{L}}_j(f) - \mathcal{L}(f)) - z}{\Delta} \right), \\ G_k(z; f) &= \sqrt{k} \mathbb{E} \rho' \left(\sqrt{n} \frac{(\bar{\mathcal{L}}_1(f) - \mathcal{L}(f)) - z}{\Delta} \right). \end{aligned}$$

In particular, when $\mathcal{O} = 0$, $G_k(z; f) = \mathbb{E}\hat{G}_k(z; f)$. Let $\hat{e}^{(k)}(f)$ and $e^{(k)}(f)$ be defined by the equations

$$\begin{aligned}\hat{G}_k\left(\hat{e}^{(k)}(f); f\right) &= 0, \\ G_k\left(e^{(k)}(f); f\right) &= 0.\end{aligned}$$

Comparing this to the definition of $\widehat{\mathcal{L}}^{(k)}(f)$ (1.2), it is easy to see that $\hat{e}^{(k)}(f) = \widehat{\mathcal{L}}^{(k)}(f) - \mathcal{L}(f)$.

Let us explain the main high-level ideas behind the proof. In the classical empirical risk minimization framework, $\widehat{\mathcal{L}}^{(k)}(f)$ is replaced by the empirical mean $P_N \ell(f) = \frac{1}{N} \sum_{j=1}^N \ell(f(X_j))$; in particular, it is linear in $\ell(f)$, meaning that $P_N(\ell(f_1) - \ell(f_2)) = P_N \ell(f_1) - P_N \ell(f_2)$, while $\widehat{\mathcal{L}}^{(k)}(f)$ lacks this property. Imagine that $\widehat{\mathcal{L}}^{(k)}(f)$ was linear in $\ell(f)$. Then, setting $\hat{\delta}_N = \mathcal{L}(\hat{f}_N) - \mathcal{L}(f_*)$, we would be able to write that

$$\begin{aligned}\hat{\delta}_N &= \mathcal{L}(\hat{f}_N) - \mathcal{L}(f_*) = (\mathcal{L}(\hat{f}_N) - \widehat{\mathcal{L}}^{(k)}(\hat{f}_N)) - (\mathcal{L}(f_*) - \widehat{\mathcal{L}}^{(k)}(f_*)) + \underbrace{\widehat{\mathcal{L}}^{(k)}(\hat{f}_N) - \widehat{\mathcal{L}}^{(k)}(f_*)}_{\leq 0} \\ &\leq \sup_{f: \mathcal{E}(f) \leq \hat{\delta}_N} \left| \widehat{\mathcal{L}}^{(k)}(f - f_*) - \mathcal{L}(f - f_*) \right|. \quad (4.1)\end{aligned}$$

It would then suffice to find a good upper bound for the supremum on the right side of (4.1) and solve the resulting inequality to get an upper bound for $\hat{\delta}_N$. However, this argument does not work directly due to the lack of linearity. Instead, we use Bahadur-type representation of the $\hat{e}^{(k)}(f)$ to introduce linearity into the problem. Specifically, we will show that $\hat{e}^{(k)}(f) = -\frac{\hat{G}_k(0; f)}{\partial_z \hat{G}_k(0; f)} + r_N(f)$ where $r_N(f)$ is a small remainder term and $\partial_z G_k(0; f)$ is the partial derivative of $G_k(z; f)$ with respect to z evaluated in $z = 0$. The process $\hat{G}_k(0; f)$ is “almost” linear in $\ell(f)$, the only obstacle being the nonlinearity due to ρ' . Mimicking (4.1), we can write that

$$\begin{aligned}\hat{\delta}_N &= \hat{e}^{(k)}(\hat{f}_N) - \hat{e}^{(k)}(f_*) + \underbrace{\widehat{\mathcal{L}}^{(k)}(\hat{f}_N) - \widehat{\mathcal{L}}^{(k)}(f_*)}_{\leq 0} \leq \hat{e}^{(k)}(\hat{f}_N) - \hat{e}^{(k)}(f_*) \\ &= \left| \frac{\hat{G}_k(0; \hat{f}_N)}{\partial_z \hat{G}_k(0; \hat{f}_N)} - \frac{\hat{G}_k(0; f_*)}{\partial_z \hat{G}_k(0; f_*)} \right| + r'_N(\hat{f}_N, f_*) \leq \sup_{f: \mathcal{E}(f) \leq \hat{\delta}_N} \left(\left| \frac{\hat{G}_k(0; f)}{\partial_z \hat{G}_k(0; f)} - \frac{\hat{G}_k(0; f_*)}{\partial_z \hat{G}_k(0; f_*)} \right| + r'_N(f, f_*) \right)\end{aligned}$$

for appropriately defined $r'_N(\cdot, \cdot)$. The difference $\frac{\hat{G}_k(0; f)}{\partial_z \hat{G}_k(0; f)} - \frac{\hat{G}_k(0; f_*)}{\partial_z \hat{G}_k(0; f_*)}$ can be tackled with the techniques commonly used to estimate suprema of the empirical processes; in particular, symmetrization and contraction inequalities for Rademacher sums [37] are used to remove the additional nonlinearity in the definition of $\hat{G}_k(z, f)$ introduced by ρ' . At that point, one only needs to carefully estimate the remainder term r'_N .

4.1 Technical tools.

We summarize the key results that our proofs rely on.

LEMMA 4.1 Let ρ satisfy Assumption 1. Then for any random variable Y with $\mathbb{E}Y^2 < \infty$,

$$\text{Var}(\rho'(Y)) \leq \text{Var}(Y).$$

Proof. See Lemma 5.3 in [47]. $\square$

LEMMA 4.2 For any function h of with bounded third derivative and a sequence of i.i.d. random variables $\xi_1, \dots, \xi_n$ such that $\mathbb{E}\xi_1 = 0$ and $\mathbb{E}|\xi_1|^3 < \infty$,

$$\left| \mathbb{E}h\left(\sum_{j=1}^n \xi_j\right) - \mathbb{E}h\left(\sum_{j=1}^n Z_j\right) \right| \leq Cn \|h'''\|_\infty \mathbb{E}|\xi_1|^3,$$

where $C > 0$ is an absolute constant and $Z_1, \dots, Z_n$ are i.i.d. centered normal random variables such that $\text{Var}(Z_1) = \text{Var}(\xi_1)$.

Proof. This bound follows from a standard application of Lindeberg's replacement method; see chapter 11 in [49]. $\square$

LEMMA 4.3 Assume that $\mathbb{E}|f(X) - \mathbb{E}f(X)|^2 < \infty$ for all $f \in \mathcal{F}$ and that ρ satisfies Assumption 1. Then for all $f \in \mathcal{F}$ and $z \in \mathbb{R}$ satisfying $|z| \leq \frac{\Delta}{\sqrt{n}} \frac{1}{2}$,

$$\left| \mathbb{E}\rho' \left(\frac{\sqrt{n}(\bar{\theta}_j(f) - Pf) - z}{\Delta} \right) - \mathbb{E}\rho' \left(\frac{W(f) - \sqrt{nz}}{\Delta} \right) \right| \leq 2G_f(n, \Delta).$$

Proof. See Lemma 4.2 in [47]. $\square$

Given N i.i.d. random variables $X_1, \dots, X_N \in \mathcal{S}$, let $\|f - g\|_{L_\infty(\Pi_N)} := \max_{1 \leq j \leq N} |f(X_j) - g(X_j)|$. Moreover, define

$$\Gamma_{n,\infty}(\mathcal{F}) := \mathbb{E}\gamma_2^2(\mathcal{F}; L_\infty(\Pi_N)),$$

where $\gamma_2(\mathcal{F}, L_\infty(\Pi_N))$ is Talagrand's generic chaining complexity [54].

LEMMA 4.4 Let $\sigma^2 := \sup_{f \in \mathcal{F}} \mathbb{E}f^2(X)$. Then there exists a universal constant $C > 0$ such that

$$\mathbb{E} \sup_{f \in \mathcal{F}} \left| \frac{1}{N} \sum_{j=1}^N f^2(X_j) - \mathbb{E}f^2(X) \right| \leq C \left(\sigma \sqrt{\frac{\Gamma_{N,\infty}(\mathcal{F})}{N}} \sqrt{\frac{\Gamma_{N,\infty}(\mathcal{F})}{N}} \right).$$

Proof. See Theorem 3.16 in [31]. $\square$

The following form of Talagrand's concentration inequality is due to Klein and Rio (see Section 12.5 in [12]).

LEMMA 4.5 Let $\{Z_j(f), f \in \mathcal{F}\}$, $j = 1, \dots, N$ be independent (not necessarily identically distributed) separable stochastic processes indexed by class $\mathcal{F}$ and such that $|Z_j(f) - \mathbb{E}Z_j(f)| \leq M$ a.s. for all $1 \leq j \leq N$ and $f \in \mathcal{F}$. Then the following inequality holds with probability at least $1 - e^{-s}$:

$$\sup_{f \in \mathcal{F}} \left(\sum_{j=1}^N (Z_j(f) - \mathbb{E}Z_j(f)) \right) \leq 2\mathbb{E} \sup_{f \in \mathcal{F}} \left(\sum_{j=1}^N (Z_j(f) - \mathbb{E}Z_j(f)) \right) + V(\mathcal{F})\sqrt{2s} + \frac{4Ms}{3}, \quad (4.2)$$

where $V^2(\mathcal{F}) = \sup_{f \in \mathcal{F}} \sum_{j=1}^N \text{Var}(Z_j(f))$.

It is easy to see, applying (4.2) to processes $\{-Z_j(f), f \in \mathcal{F}\}$, that

$$\inf_{f \in \mathcal{F}} \left(\sum_{j=1}^N (Z_j(f) - \mathbb{E}Z_j(f)) \right) \geq -2\mathbb{E} \sup_{f \in \mathcal{F}} \left(\sum_{j=1}^N (\mathbb{E}Z_j(f) - Z_j(f)) \right) - V(\mathcal{F})\sqrt{2s} - \frac{4Ms}{3}$$

with probability at least $1 - e^{-s}$.

4.2 Proof of Theorems 2.2 and 2.3.

We will provide detailed proofs for the estimator $\hat{f}_N$ that is based on disjoint subsamples indexed by $G_1, \dots, G_k$. The bounds for its permutation-invariant version $\hat{f}_N^U$ follow exactly the same steps where all applications of the Talagrand's concentration inequality (Lemma 4.5) should be replaced by its version (B.3) for nondegenerate U-statistics stated in Section B of the Supplementary material [41].

Let $J \subset \{1, \dots, k\}$ of cardinality $|J| \geq k - \mathcal{O}$ be the set containing all j such that the subsample $\{X_i, i \in G_j\}$ does not include outliers. Clearly, $\{X_i : i \in G_j, j \in J\}$ are still conditionally i.i.d. as the partitioning scheme is independent of the data. Moreover, set $N_J := \sum_{j \in J} |G_j|$, and note that, since $\mathcal{O} < k/2$,

$$N_J \geq n|J| \geq \frac{N}{2}.$$

Consider stochastic process $R_N(f)$ defined as

$$R_N(f) = \hat{G}_k(0; f) + \partial_z G_k(0; f) \cdot \hat{e}^{(k)}(f), \quad (4.3)$$

where $\partial_z G_k(0; f) := \partial_z G_k(z; f)|_{z=0}$. Whenever $\partial_z G_k(0; f) \neq 0$ (this assumption will be justified by Lemma 4.6 below), we can solve (4.3) for $\hat{e}^{(k)}(f)$ to obtain

$$\hat{e}^{(k)}(f) = -\frac{\hat{G}_k(0; f)}{\partial_z G_k(0; f)} + \frac{R_N(f)}{\partial_z G_k(0; f)}, \quad (4.4)$$

which can be viewed as a Bahadur-type representation of $\hat{e}^{(k)}(f)$. Setting $f := \hat{f}_N$ and recalling that $\hat{e}^{(k)}(f) = \widehat{\mathcal{L}}^{(k)}(f) - \mathcal{L}(f)$, we deduce that

$$\widehat{\mathcal{L}}^{(k)}(\hat{f}_N) = \mathcal{L}(\hat{f}_N) - \frac{\hat{G}_k(0; \hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)} + \frac{R_N(\hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)}.$$

By the definition (1.4) of $\hat{f}_N$, $\widehat{\mathcal{L}}^{(k)}(\hat{f}_N) \leq \widehat{\mathcal{L}}^{(k)}(f_*)$, hence

$$\mathcal{L}(\hat{f}_N) - \frac{\hat{G}_k(0; \hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)} + \frac{R_N(\hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)} \leq \mathcal{L}(f_*) - \frac{\hat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} + \frac{R_N(f_*)}{\partial_z G_k(0; f_*)}.$$

Rearranging the terms, it is easy to see that

$$\widehat{\delta}_N = \mathcal{L}(\hat{f}_N) - \mathcal{L}(f_*) \leq \left| \frac{\hat{G}_k(0; \hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)} - \frac{\hat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \right| + 2 \sup_{f \in \mathcal{F}(\widehat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right|. \quad (4.5)$$

REMARK 4.1 Similar argument also implies, in view of the inequality $\mathcal{L}(f_*) \leq \mathcal{L}(\hat{f}_N)$, that

$$\widehat{\mathcal{L}}^{(k)}(f_*) + \frac{\hat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} - \frac{R_N(f_*)}{\partial_z G_k(0; f_*)} \leq \widehat{\mathcal{L}}^{(k)}(\hat{f}_N) + \frac{\hat{G}_k(0; \hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)} - \frac{R_N(\hat{f}_N)}{\partial_z G_k(0; \hat{f}_N)},$$

hence

$$\widehat{\mathcal{L}}^{(k)}(f_*) - \widehat{\mathcal{L}}^{(k)}(\widehat{f}_N) \leq \left| \frac{\widehat{G}_k(0; \widehat{f}_N)}{\partial_z G(0; \widehat{f}_N)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G(0; f_*)} \right| + 2 \sup_{f \in \mathcal{F}(\widehat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right|.$$

It follows from (4.5) that in order to estimate the excess risk of $\widehat{f}_N$, it suffices to obtain the upper bounds for

$$A_1 := \left| \frac{\widehat{G}_k(0; \widehat{f}_N)}{\partial_z G_k(0; \widehat{f}_N)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \right| \text{ and } A_2 := \sup_{f \in \mathcal{F}(\widehat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right|.$$

Observe that

$$\begin{aligned} & \frac{\widehat{G}_k(0; \widehat{f}_N)}{\partial_z G_k(0; \widehat{f}_N)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \\ &= \frac{\widehat{G}_k(0; \widehat{f}_N) - \widehat{G}_k(0; f_*)}{\partial_z G_k(0; \widehat{f}_N)} + \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \left(\partial_z G_k(0; f_*) - \partial_z G_k(0; \widehat{f}_N) \right). \end{aligned}$$

Since ρ'' is Lipschitz continuous by assumption,

$$\begin{aligned} & \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \left(\partial_z G_k(0; f_*) - \partial_z G_k(0; \widehat{f}_N) \right) \right| \\ &= \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \frac{\sqrt{nk}}{\Delta} \mathbb{E} \left(\rho'' \left(\frac{\sqrt{n} \overline{\mathcal{L}}_1(f_*) - \mathcal{L}(f_*)}{\Delta} \right) - \rho'' \left(\frac{\sqrt{n} \overline{\mathcal{L}}_1(\widehat{f}_N) - \mathcal{L}(\widehat{f}_N)}{\Delta} \right) \right) \right| \\ &\leq L(\rho'') \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \frac{\sqrt{nk}}{\Delta^2} \text{Var}^{1/2} \left(\ell(\widehat{f}_N(X)) - \ell(f_*(X)) \right) \right| \\ &= C(\rho) \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \right| \frac{\sqrt{nk}}{\Delta^2} \nu(\widehat{\delta}_N). \quad (4.6) \end{aligned}$$

The following two lemmas are required to proceed.

LEMMA 4.6 There exist $C(\rho) > 0$ such that for any $f \in \mathcal{F}$,

$$|\partial_z G_k(0; f)| \geq \frac{\sqrt{kn}}{2\sqrt{2\pi}\Delta} \left(\min \left(\frac{\Delta}{\sqrt{\text{Var}(\ell(f(X)))}}, 2\sqrt{\log 2} \right) - \frac{C(\rho)}{\sqrt{n}} \mathbb{E} \left| \frac{\ell(f(X)) - P\ell(f)}{\Delta} \right|^3 \right).$$

Proof. See Section A.1. □

In particular, the bound of Lemma 4.6 implies that for n large enough,

$$\inf_{f \in \mathcal{F}} |\partial_z G_k(0; f)| \geq \frac{1}{4\sqrt{2\pi}} \frac{\sqrt{kn}}{\max(\Delta, \sigma(\ell, \mathcal{F}))} = \frac{1}{4\sqrt{2\pi}} \frac{\sqrt{kn}}{\Delta}. \quad (4.7)$$

It is also easy to deduce from the proof of Lemma 4.6 that for small n and $\Delta > \sigma(\ell, \mathcal{F})$, $\inf_{f \in \mathcal{F}} |\partial_z G_k(0; f)| \geq c(\rho) \frac{\sqrt{kn}}{\Delta}$ for some positive $c(\rho)$.

LEMMA 4.7 For any $f \in \mathcal{F}$,

$$\widehat{G}_k(0; f) \leq 2 \left(\sqrt{k} G_f(n, \Delta) + \frac{\sigma(\ell, f)}{\Delta} \sqrt{s} + \frac{2s}{\sqrt{k}} + \frac{\mathcal{O}}{\sqrt{k}} \right)$$

with probability at least $1 - 2e^{-s}$, where $C > 0$ is an absolute constant.

Proof.

See Section A.2.

□

Lemma 4.7 and (4.7) imply, together with (4.6), that

$$\begin{aligned} & \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \left(\partial_z G_k(0; f_*) - \partial_z G_k(0; \widehat{f}_N) \right) \right| \\ & \leq C(\rho) \frac{\widetilde{\Delta}^2}{\Delta^2} \left(\frac{\sigma(\ell, f_*)}{\Delta} \sqrt{\frac{s}{N}} + \frac{G_{f_*}(n, \Delta)}{\sqrt{n}} + \sqrt{n} \frac{s}{N} + \sqrt{n} \frac{\mathcal{O}}{N} \right) v(\widehat{\delta}_N) \end{aligned}$$

on event Θ_1 of probability at least $1 - 2e^{-s}$. As $\widetilde{\Delta} \geq \sigma(\ell, \mathcal{F})$ by assumption, we deduce that

$$\begin{aligned} & \left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \left(\partial_z G_k(0; f_*) - \partial_z G_k(0; \widehat{f}_N) \right) \right| \\ & \leq C(\rho) v(\widehat{\delta}_N) \left(\sqrt{\frac{s}{N}} + \frac{G_{f_*}(n, \Delta)}{\sqrt{n}} + \sqrt{n} \frac{s}{N} + \sqrt{n} \frac{\mathcal{O}}{N} \right). \end{aligned}$$

Define

$$\bar{\delta}_1 := \min \left\{ \delta > 0 : C_1(\rho) \left(\sqrt{\frac{s}{N}} + \frac{G_{f_*}(n, \Delta)}{\sqrt{n}} + \sqrt{n} \frac{s}{N} + \sqrt{n} \frac{\mathcal{O}}{N} \right) \frac{\widetilde{v}(\delta)}{\delta} \leq \frac{1}{7} \right\}, \quad (4.8)$$

where $C_1(\rho)$ is sufficiently large. It is easy to see that on event $\Theta_1 \cap \{\widehat{\delta}_N > \bar{\delta}_1\}$,

$$\left| \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*) \partial_z G_k(0; \widehat{f}_N)} \left(\partial_z G_k(0; f_*) - \partial_z G_k(0; \widehat{f}_N) \right) \right| \leq \frac{\widehat{\delta}_N}{7}, \quad (4.9)$$

for appropriately chosen $C_1(\rho)$.

Our next goal is to obtain an upper bound for $\left| \frac{\widehat{G}_k(0; \widehat{f}_N) - \widehat{G}_k(0; f_*)}{\partial_z G_k(0; \widehat{f}_N)} \right|$. To this end, we will need to control the local oscillations of the process $\widehat{G}_k(0; f)$. Specifically, we are interested in the bounds on the random variable $\sup_{f \in \mathcal{F}(\delta)} |\widehat{G}_k(0; f) - \widehat{G}_k(0; f_*)|$. The following technical lemma is important for the analysis.

LEMMA 4.8 Let $(\xi_1, \eta_1), \dots, (\xi_n, \eta_n)$ be a sequence of independent identically distributed random couples such that $\mathbb{E}\xi_1 = 0$, $\mathbb{E}\eta_1 = 0$, and $\mathbb{E}|\xi_1|^2 + \mathbb{E}|\eta_1|^2 < \infty$. Let F be an odd, smooth function with

bounded derivatives up to fourth order. Then

$$\left| \mathbb{E}F\left(\sum_{j=1}^n \xi_j\right) - \mathbb{E}F\left(\sum_{j=1}^n \eta_j\right) \right| \leq \max_{\alpha \in [0,1]} \sqrt{n} \text{Var}^{1/2}(\xi_1 - \eta_1) \left(\mathbb{E} \left| F' \left(S_n^\eta + \alpha (S_n^\xi - S_n^\eta) \right) \right|^2 \right)^{1/2}.$$

Moreover, if $\mathbb{E}|\xi_1|^4 + \mathbb{E}|\eta_1|^4 < \infty$, then

$$\left| \mathbb{E}F\left(\sum_{j=1}^n \xi_j\right) - \mathbb{E}F\left(\sum_{j=1}^n \eta_j\right) \right| \leq C(F) \cdot n \left(\text{Var}^{1/2}(\xi_1 - \eta_1) \left(R_4^2 + \sqrt{n-1} R_4^3 \right) + (\mathbb{E}|\xi_1 - \eta_1|^4)^{1/4} R_4^3 \right),$$

where $R_4 = (\max(\mathbb{E}|\xi_1|^4, \mathbb{E}|\eta_1|^4))^{1/4}$ and $C(F) > 0$ is a constant that depends only on F .

Proof. See Section A.3. □

Now we are ready to state the bound for the local oscillations of the process $\widehat{G}_k(0; f)$. Let

$$U(\delta, s) := \frac{2}{\Delta} \left(8\sqrt{2}\omega(\delta) + \nu(\delta) \sqrt{\frac{s}{2}} \right) + \frac{32s}{3\sqrt{k}}.$$

Moreover, if $\tilde{\omega}(\delta)$ and $\tilde{\nu}(\delta)$ are upper bounds for $\omega(\delta)$ and $\nu(\delta)$ and are of concave type, then

$$\tilde{U}(\delta, s) := \frac{2}{\Delta} \left(c(\gamma) \tilde{\omega}(\delta) + \tilde{\nu}(\delta) \sqrt{\frac{s}{2}} \right) + \frac{32s}{\sqrt{k}}, \quad (4.10)$$

where $c(\gamma) > 0$ depends only on γ , is also an upper bound for $U(\delta, s)$ of strictly concave type. Moreover, define

$$\begin{aligned} R_4(\ell, \mathcal{F}) &:= \sup_{f \in \mathcal{F}} \mathbb{E}^{1/4} \left(\ell(f(X)) - \mathbb{E} \ell(f(X)) \right)^4, \\ \nu_4(\delta) &:= \sup_{f \in \mathcal{F}(\delta)} \mathbb{E}^{1/4} \left(\ell(f(X)) - \ell(f_*(X)) - \mathbb{E}(\ell(f(X)) - \ell(f_*(X))) \right)^4, \\ \mathfrak{B}(\ell, \mathcal{F}) &:= \frac{R_4(\ell, \mathcal{F})}{\sigma(\ell, \mathcal{F})}, \\ \tilde{B}(\delta) &:= \begin{cases} \frac{\tilde{\nu}(\delta)}{\Delta} \frac{1}{M_\Delta}, & R_4(\ell, \mathcal{F}) = \infty, \\ \frac{\mathfrak{B}^3(\ell, \mathcal{F})}{\sqrt{n}} \left(\frac{\tilde{\nu}(\delta)}{\Delta} \frac{1}{M_\Delta^2} + \frac{\tilde{\nu}_4(\delta)}{\Delta} \frac{1}{M_\Delta^3 \sqrt{n}} \right), & R_4(\ell, \mathcal{F}) < \infty, \end{cases} \end{aligned}$$

where $\tilde{\nu}_4(\delta)$ upper bounds $\nu_4(\delta)$ and is of concave type. Below, we will use a crude bound $\nu_4(\delta) \leq 2R_4(\ell, \mathcal{F})$, but additional improvements are possible if better estimates of $\nu_4(\delta)$ are available.

LEMMA 4.9 With probability at least $1 - e^{-2s}$,

$$\sup_{f \in \mathcal{F}(\delta)} \left| \widehat{G}_k(0; f) - \widehat{G}_k(0; f_*) \right| \leq U(\delta, s) + C(\rho) \sqrt{k} \tilde{B}(\delta) + 4 \frac{\mathcal{O}}{\sqrt{k}},$$

where $C(\rho) > 0$ is constant that depends only on ρ .

Proof. See Section A.4. □

Next, we state the “uniform version” of Lemma 4.9.

LEMMA 4.10 With probability at least $1 - e^{-s}$, for all $\delta \geq \delta_{\min}$ simultaneously,

$$\sup_{f \in \mathcal{F}(\delta)} \left| \widehat{G}_k(0; f) - \widehat{G}_k(0; f_*) \right| \leq C(\rho) \delta \left(\frac{\widetilde{U}(\delta_{\min}, s)}{\delta_{\min}} + \sqrt{k} \frac{\widetilde{B}(\delta_{\min})}{\delta_{\min}} \right) + 4 \frac{\mathcal{O}}{\sqrt{k}},$$

where $C(\rho) > 0$ is constant that depends only on ρ .

Proof. See Section A.5. □

It follows from Lemma 4.10 and inequality (4.7) that on event Θ_2 of probability at least $1 - e^{-s}$, for all $\delta \geq \delta_{\min}$ simultaneously,

$$\sup_{f \in \mathcal{F}(\delta)} \left| \frac{\widehat{G}_k(0; f) - \widehat{G}_k(0; f_*)}{\partial_z G_k(0; f)} \right| \leq C(\rho) \delta \left(\frac{\widetilde{\Delta}}{\sqrt{N}} \frac{\widetilde{U}(\delta_{\min}, s)}{\delta_{\min}} + \frac{\widetilde{\Delta}}{\sqrt{n}} \frac{\widetilde{B}(\delta_{\min})}{\delta_{\min}} \right) + 4 \widetilde{\Delta} \sqrt{n} \frac{\mathcal{O}}{N}.$$

Define

$$\begin{aligned} \bar{\delta}_2 &:= \min \left\{ \delta > 0 : C_2(\rho) \frac{\widetilde{\Delta}}{\sqrt{N}} \frac{\widetilde{U}(\delta, s)}{\delta} \leq \frac{1}{7} \right\}, \\ \bar{\delta}_3 &:= \min \left\{ \delta > 0 : C_3(\rho) \frac{\widetilde{\Delta}}{\sqrt{n}} \frac{\widetilde{B}(\delta)}{\delta} \leq \frac{1}{7} \right\}, \end{aligned}$$

where $C_2(\rho)$, $C_3(\rho)$ are sufficiently large constants. Then, on event $\Theta_2 \cap \{\widehat{\delta}_N > \max(\bar{\delta}_2, \bar{\delta}_3)\}$,

$$\sup_{f \in \mathcal{F}(\widehat{\delta}_N)} \left| \frac{\widehat{G}_k(0; f) - \widehat{G}_k(0; f_*)}{\partial_z G_k(0; f)} \right| \leq \frac{2 \widehat{\delta}_N}{7} + 4 \widetilde{\Delta} \sqrt{n} \frac{\mathcal{O}}{N} \quad (4.11)$$

for appropriately chosen $C_2(\rho), C_3(\rho)$.

Finally, we provide an upper bound for the process $R_N(f)$ defined via

$$R_N(f) = \widehat{G}_k(0; f) + \partial_z G_k(0; f) \cdot \widehat{e}^{(k)}(f).$$

LEMMA 4.11 Assume that conditions of Theorem 2.1 hold, and let $\delta_{\min} > 0$ be fixed. Then for all $s > 0$, $\delta \geq \delta_{\min}$, positive integers n and k such that

$$\delta \frac{\widetilde{U}(\delta_{\min}, s)}{\delta_{\min} \sqrt{k}} + \sup_{f \in \mathcal{F}} G_f(n, \Delta) + \frac{s + \mathcal{O}}{k} \leq c(\rho), \quad (4.12)$$

the following inequality holds with probability at least $1 - 7e^{-s}$, uniformly over all δ satisfying (4.12):

$$\begin{aligned} \sup_{f \in \mathcal{F}(\delta)} |R_N(f)| &\leq C(\rho) \sqrt{N} \frac{\widetilde{\Delta}^2}{\Delta^2} \left(n^{1/2} \delta^2 \left(\frac{\widetilde{U}(\delta_{\min}, s)}{\delta_{\min} \sqrt{N}} \right)^2 \vee \frac{\sigma^2(\ell, f_*)}{\Delta^2} \frac{n^{1/2} s}{N} \right. \\ &\quad \left. \vee n^{1/2} \left(\sup_{f \in \mathcal{F}} \frac{G_f(n, \Delta)}{\sqrt{n}} \right)^2 \vee n^{3/2} \frac{s^2}{N^2} \vee n^{3/2} \frac{\mathcal{O}^2}{N^2} \right). \end{aligned}$$

Moreover, the bound of Theorem 2.1 holds on the same event.

Proof. See Section A.6. □

Recall that

$$\bar{\delta}_2 = \min \left\{ \delta > 0 : C_2(\rho) \frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\delta, s)}{\delta} \leq \frac{1}{7} \right\},$$

where $C_2(\rho)$ is a large enough constant. Let Θ_3 be the event of probability at least $1 - 7e^{-s}$ on which Lemma 4.11 holds with $\delta_{\min} = \bar{\delta}_2$, and consider the event $\Theta_3 \cap \{\hat{\delta}_N > \bar{\delta}_2\}$. We will now show that on this event, Lemma 4.11 applies with $\delta = \hat{\delta}_N$. Indeed, the bound of Theorem 2.1 is valid on Θ_3 , hence the inequality (2.4) implies that on Θ_3 , $\hat{\delta}_N \leq C(\rho) \frac{\tilde{\Delta}}{\sqrt{n}}$, and it is straightforward to check that condition (4.12) of Lemma 4.11 holds with $\delta_{\min} = \bar{\delta}_2$ and $\delta = \hat{\delta}_N$. It follows from inequality (4.7) that on event $\Theta_3 \cap \{\hat{\delta}_N \geq \bar{\delta}_2\}$,

$$\begin{aligned} \sup_{f \in \mathcal{F}(\hat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right| &\leq C(\rho) \frac{\tilde{\Delta}^2}{\Delta^2} \left(\frac{n^{1/2}}{\tilde{\Delta}} \hat{\delta}_N^2 \left(\frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\delta_2, s)}{\delta_2} \right)^2 \vee \tilde{\Delta} \frac{\sigma^2(\ell, f_*)}{\Delta^2} \frac{n^{1/2} s}{N} \right. \\ &\quad \left. \vee n^{1/2} \tilde{\Delta} \left(\sup_{f \in \mathcal{F}} \frac{G_f(n, \Delta)}{\sqrt{n}} \right)^2 \vee n^{3/2} \tilde{\Delta} \frac{s^2 + \mathcal{O}^2}{N^2} \right). \end{aligned}$$

Consider the expression

$$C(\rho) \frac{\tilde{\Delta}^2}{\Delta^2} \frac{n^{1/2}}{\tilde{\Delta}} \hat{\delta}_N^2 \left(\frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\delta_2, s)}{\delta_2} \right)^2 = C(\rho) \frac{\tilde{\Delta}^2}{\Delta^2} \left(\frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\delta_2, s)}{\delta_2} \right)^2 \hat{\delta}_N \cdot \frac{n^{1/2} \hat{\delta}_N}{\tilde{\Delta}},$$

and observe that whenever Theorem 2.1 holds, $\frac{n^{1/2} \hat{\delta}_N}{\tilde{\Delta}} \leq c(\rho)$, hence the latter is bounded from above by

$$\hat{\delta}_N \cdot C(\rho) \frac{\tilde{\Delta}^2}{\Delta^2} \left(\frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\bar{\delta}_2, s)}{\bar{\delta}_2} \right)^2 \leq \frac{\hat{\delta}_N}{7}$$

whenever $\Delta \geq \sigma(\ell, \mathcal{F})$ (so that $\tilde{\Delta} = \Delta$) and $C_2(\rho)$ in the definition of $\bar{\delta}_2$ is large enough. Moreover,

$$C(\rho) \frac{\tilde{\Delta}^3}{\Delta^3} \frac{\sigma^2(\ell, f_*)}{\Delta} \frac{n^{1/2} s}{N} \leq C'(\rho) \cdot \sigma(\ell, f_*) \sqrt{n} \frac{s}{N} \leq C'(\rho) \tilde{\Delta} \sqrt{n} \frac{s}{N}$$

if $\tilde{\Delta} \geq \sigma(\ell, f_*)$. As $\frac{s + \mathcal{O}}{k} \leq c$ under the conditions of Theorem 2.1, $n^{3/2} \tilde{\Delta} \frac{s^2 + \mathcal{O}^2}{N^2} \leq C \tilde{\Delta} \sqrt{n} \frac{s + \mathcal{O}}{N}$. Combining the inequalities obtained above, we deduce on event $\Theta_3 \cap \{\hat{\delta}_N \geq \bar{\delta}_2\}$,

$$2 \sup_{f \in \mathcal{F}(\hat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right| \leq \frac{2\hat{\delta}_N}{7} + C(\rho) \tilde{\Delta} \left(\sqrt{n} \frac{s + \mathcal{O}}{N} \vee \frac{\sup_{f \in \mathcal{F}} (G_f(n, \Delta))^2}{\sqrt{n}} \right)$$

whenever $\tilde{\Delta} \geq \sigma(\ell, \mathcal{F})$. Finally, define

$$\bar{\delta}_4 := C_4(\rho) \tilde{\Delta} \left(\sqrt{n} \frac{s + \mathcal{O}}{N} \vee \frac{\sup_{f \in \mathcal{F}} (G_f(n, \Delta))^2}{\sqrt{n}} \right),$$

where $C_4(\rho)$ is sufficiently large. Then on event $\Theta_3 \cap \{\hat{\delta}_N \geq \max(\bar{\delta}_2, 7\bar{\delta}_4)\}$,

$$2 \sup_{f \in \mathcal{F}(\hat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right| + 4\tilde{\Delta} \sqrt{n} \frac{\mathcal{O}}{N} \leq \frac{2\hat{\delta}_N}{7} + \frac{\hat{\delta}_N}{7} = \frac{3\hat{\delta}_N}{7}. \quad (4.13)$$

Note that the expression above takes care of the term $4\tilde{\Delta} \sqrt{n} \frac{\mathcal{O}}{N}$ that appeared in (4.11). Combining (4.9), (4.11), (4.13), we deduce that on event $\Theta_1 \cap \Theta_2 \cap \Theta_3 \cap \{\hat{\delta}_N \geq \max(\bar{\delta}_1, \bar{\delta}_2, \bar{\delta}_3, 7\bar{\delta}_4)\}$,

$$\hat{\delta}_N \leq \frac{6}{7} \hat{\delta}_N,$$

leading to a contradiction, hence on event $\Theta_1 \cap \Theta_2 \cap \Theta_3$ of probability at least $1 - 10e^{-s}$,

$$\hat{\delta}_N \leq \max(\bar{\delta}_1, \bar{\delta}_2, \bar{\delta}_3, 7\bar{\delta}_4). \quad (4.14)$$

Recall the definition (4.8) of $\bar{\delta}_1$. If condition 2 (“Bernstein condition”) holds, then $\tilde{v}(\delta) \leq D\sqrt{\delta}$ for small enough δ , in which case

$$\bar{\delta}_1 \leq C(\rho) D^2 \left(\frac{s + \mathcal{O}}{N} + \frac{G_{f_*}^2(n, \Delta)}{n} \right),$$

where we used the fact that $\frac{s}{k} \leq c$ by assumption. Together with the bound (2.1) for $G_{f_*}(n, \Delta)$, we deduce that, under the assumption that $R_4(\ell, \mathcal{F}) < \infty$,

$$\bar{\delta}_1 \leq C(\rho) D^2 \left(\frac{s + \mathcal{O}}{N} + \frac{(\mathbb{E}|f_*(X) - \mathbb{E}f_*(X)|^3)^2}{\Delta^6 n^2} \right).$$

Since $\Delta = \sigma(\ell, \mathcal{F}) M_\Delta$, $\frac{\mathbb{E}|f_*(X) - \mathbb{E}f_*(X)|^3}{\Delta^3} \leq \frac{\sup_{f \in \mathcal{F}} \mathbb{E}|f(X) - \mathbb{E}f(X)|^3}{\sigma^3(\ell, \mathcal{F}) M_\Delta^3} \leq \frac{\mathfrak{B}^3(\ell, \mathcal{F})}{M_\Delta^3}$, where

$$\mathfrak{B}(\ell, \mathcal{F}) = \frac{\sup_{f \in \mathcal{F}} \mathbb{E}^{1/4}(\ell(f(X)) - \mathbb{E}\ell(f(X)))^4}{\sigma(\ell, \mathcal{F})},$$

hence

$$\bar{\delta}_1 \leq C(\rho) D^2 \left(\frac{s + \mathcal{O}}{N} + \frac{\mathfrak{B}^6(\ell, \mathcal{F})}{n^2 M_\Delta^6} \right). \quad (4.15)$$

At the same time, if only $\sigma(\ell, \mathcal{F}) < \infty$, we similarly obtain that

$$\bar{\delta}_1 \leq C(\rho) D^2 \left(\frac{s + \mathcal{O}}{N} + \frac{1}{M_\Delta^4 n} \right). \quad (4.16)$$

Next we will estimate $\bar{\delta}_3$. Recall that, when $R_4(\ell, \mathcal{F}) < \infty$,

$$\tilde{B}(\delta) = \frac{\mathfrak{B}^3(\ell, \mathcal{F})}{\sqrt{n}} \left(\frac{\tilde{v}(\delta)}{\Delta} \frac{1}{M_\Delta^2} + \frac{\tilde{v}_4(\delta)}{\Delta} \frac{1}{M_\Delta^3 \sqrt{n}} \right).$$

For sufficiently small δ (namely, for which condition 2 holds) and $\Delta \geq \sigma(\ell, \mathcal{F})$,

$$\frac{\tilde{\Delta}}{\sqrt{n}} \tilde{B}(\delta) \leq \frac{\mathfrak{B}^3(\ell, \mathcal{F})}{n} \left(\frac{\tilde{v}(\delta)}{M_\Delta^2} + \frac{R_4(\ell, \mathcal{F})}{M_\Delta^3 \sqrt{n}} \right) \leq \frac{\mathfrak{B}^3(\ell, \mathcal{F})}{n} \left(D \frac{\sqrt{\delta}}{M_\Delta^2} + \sigma(\ell, \mathcal{F}) \frac{\mathfrak{B}(\ell, \mathcal{F})}{M_\Delta^3 \sqrt{n}} \right)$$

and

$$\bar{\delta}_3 \leq C(\rho) \left(D^2 \frac{\mathfrak{B}^6(\ell, \mathcal{F})}{n^2 M_\Delta^4} + \sigma(\ell, \mathcal{F}) \frac{\mathfrak{B}^4(\ell, \mathcal{F})}{n^{3/2} M_\Delta^3} \right). \quad (4.17)$$

At the same time, if only the second moments are finite, $\tilde{B}(\delta) = \frac{\tilde{v}(\delta)}{\Delta} \frac{1}{M_\Delta}$, and it is easy to deduce that in this case,

$$\bar{\delta}_3 \leq C(\rho) \frac{D^2}{M_\Delta^2 n}. \quad (4.18)$$

Next, we obtain a simpler bound for $\bar{\delta}_4$: as $\Delta \geq \sigma(\ell, \mathcal{F})$ by assumption, $\tilde{\Delta} = \Delta = \sigma(\ell, \mathcal{F}) M_\Delta$, and the estimate (2.1) for $G_{f_*}(n, \Delta)$ implies (if $R_4(\ell, \mathcal{F}) < \infty$) that

$$\bar{\delta}_4 \leq C(\rho) \sigma(\ell, \mathcal{F}) \left(\sqrt{n} M_\Delta \frac{s + \mathcal{O}}{N} + \frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_\Delta^5 n^{3/2}} \right). \quad (4.19)$$

If only $\sigma(\ell, \mathcal{F}) < \infty$, we similarly deduce from (2.1) that

$$\bar{\delta}_4 \leq C(\rho) \sigma(\ell, \mathcal{F}) \left(\sqrt{n} M_\Delta \cdot \frac{s + \mathcal{O}}{N} + \frac{1}{M_\Delta^3 \sqrt{n}} \right). \quad (4.20)$$

Finally, recall that $\tilde{U}(\delta, s) = \frac{2}{\Delta} (c(\gamma) \tilde{w}(\delta) + \tilde{v}(\delta) \sqrt{\frac{s}{2}}) + \frac{32s}{\sqrt{k}}$ and $\bar{\delta}_2 = \min \left\{ \delta > 0 : C_2(\rho) \frac{\tilde{\Delta}}{\sqrt{N}} \frac{\tilde{U}(\delta, s)}{\delta} \leq \frac{1}{\gamma} \right\}$, hence

$$\bar{\delta}_2 \leq \bar{\delta} \bigvee C(\rho) D^2 \frac{s}{N} \bigvee C(\rho) \sigma(\ell, \mathcal{F}) \frac{s \sqrt{n} M_\Delta}{N}, \quad (4.21)$$

where $\bar{\delta}$ was defined in (2.5). Combining inequalities (4.15), (4.21), (4.17), (4.19) and (4.14), we obtain the final form of the bound under the stronger assumption $R_4(\ell, \mathcal{F}) < \infty$. Similarly, the combination of (4.16), (4.21), (4.18), (4.20) and (4.14) yields the bound under the weaker assumption $\sigma(\ell, \mathcal{F}) < \infty$.

4.3 Proof of Theorem 2.4.

Recall that $\hat{\mathcal{E}}_N(f_*) := \widehat{\mathcal{L}}^{(k)}(f_*) - \widehat{\mathcal{L}}^{(k)}(\hat{f}_N^*)$ is the “empirical excess risk” of f_* , and let $\hat{\delta}_N := \mathcal{E}(\hat{f}_N^*)$. It follows from Remark 4.1 that (using the notation used in the proof of Theorems 2.2 and 2.3)

$$\hat{\mathcal{E}}_N(f_*) \leq \left| \frac{\hat{G}_k(0; \hat{f}_N^*)}{\partial_z G(0; \hat{f}_N^*)} - \frac{\hat{G}_k(0; f_*)}{\partial_z G(0; f_*)} \right| + 2 \sup_{f \in \mathcal{F}(\hat{\delta}_N)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right|.$$

On the event of Theorem 2.3 of probability at least $1 - 10e^{-s}$,

$$\mathcal{E}(\hat{f}_N^*) \leq \delta' := \bar{\delta} + C(\rho) (D^2 \sigma(\ell, \mathcal{F}) \sqrt{n} M_\Delta) \left(\frac{\mathfrak{B}^6(\ell, \mathcal{F})}{M_\Delta^4 n^2} + \frac{s + \mathcal{O}}{N} \right),$$

hence on this event

$$\widehat{\mathcal{E}}_N(f_*) \leq \sup_{f \in \widehat{\mathcal{F}}(\delta')} \left| \frac{\widehat{G}_k(0; f)}{\partial_z G(0; f)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G(0; f_*)} \right| + 2 \sup_{f \in \widehat{\mathcal{F}}(\delta')} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right| \leq \frac{6}{7} \delta',$$

where the last inequality again follows from main steps in the proof of Theorem 2.3; note that similar result holds if δ' is replaced by its analogue from Theorem 2.3. Consider the set $\widehat{\mathcal{F}}(\delta') = \{f \in \mathcal{F} : \widehat{\mathcal{E}}_N(f) \leq \delta'\}$.

First, observe that on the event $\mathcal{E}_1$ of Theorem 2.3, $f_* \in \widehat{\mathcal{F}}(\delta')$ as implied by the previous display. We will next show that $\widehat{\mathcal{F}}(\delta') \subseteq \mathcal{F}(7\delta')$ on the event $\mathcal{E}_1$ of Theorem 2.3, meaning that for any $f \in \widehat{\mathcal{F}}(\delta')$, $\mathcal{E}(f) \leq 7\delta'$. Indeed, let $f \in \widehat{\mathcal{F}}(\delta')$ be such that $\mathcal{E}(f) = \sigma$. Then (4.4) implies that

$$\begin{aligned} \mathcal{L}(f) - \mathcal{L}(f_*) &\leq \widehat{\mathcal{L}}^{(k)}(f) - \widehat{\mathcal{L}}^{(k)}(f_*) + \left| \frac{\widehat{G}_k(0; f)}{\partial_z G_k(0; f)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \right| + \left| \frac{R_N(f)}{\partial_z G_k(0; f)} + \frac{R_N(f_*)}{\partial_z G_k(0; f_*)} \right| \\ &\leq \widehat{\mathcal{E}}_N(f) + \sup_{f \in \widehat{\mathcal{F}}(\sigma)} \left| \frac{\widehat{G}_k(0; f)}{\partial_z G_k(0; f)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \right| + 2 \sup_{f \in \widehat{\mathcal{F}}(\sigma)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right|. \end{aligned}$$

Again, it follows from the arguments used in proof of Theorem 2.3 that on event $\mathcal{E}_1$ of probability at least $1 - 10e^{-s}$,

$$\sup_{f \in \widehat{\mathcal{F}}(\sigma)} \left| \frac{\widehat{G}_k(0; f)}{\partial_z G_k(0; f)} - \frac{\widehat{G}_k(0; f_*)}{\partial_z G_k(0; f_*)} \right| + 2 \sup_{f \in \widehat{\mathcal{F}}(\sigma)} \left| \frac{R_N(f)}{\partial_z G_k(0; f)} \right| \leq \frac{6}{7} \max(\delta', \sigma).$$

Consequently, $\sigma \leq \delta' + \frac{6}{7} \max(\delta', \sigma)$ on this event, implying that $\sigma \leq 7\delta'$. Next, Assumption 2 yields that

$$\begin{aligned} &\sup_{f \in \widehat{\mathcal{F}}(\delta')} \text{Var} \left(\ell(f(X)) - \ell(\widehat{f}_N) \right) \\ &\leq 2 \left(\sup_{f \in \widehat{\mathcal{F}}(\delta')} \text{Var}(\ell(f(X)) - \ell(f_*(X))) + \text{Var}(\ell(\widehat{f}_N(X)) - \ell(f_*(X))) \right) \leq 2D(\sqrt{7} + 1)\delta' \end{aligned}$$

on $\mathcal{E}_1$. It remains to apply Theorem 2.3, conditionally on $\mathcal{E}_1$, to the class

$$\widehat{\mathcal{F}}(\delta') - \widehat{f}_N := \{f - \widehat{f}_N, f \in \widehat{\mathcal{F}}(\delta')\}.$$

To this end, we need to verify the assumption of Theorem 2.1 that translates into the requirement

$$c\Delta_2 \geq \frac{1}{\sqrt{k_2}} \mathbb{E} \sup_{f \in \widehat{\mathcal{F}}(7\delta')} \frac{1}{\sqrt{|S_2|}} \sum_{j=1}^{|S_2|} (\ell(f(X_j)) - \ell(f_*(X_j)) - P(\ell(f) - \ell(f_*))).$$

As $\delta' > \bar{\delta}$ and $|S_2| \geq \lfloor N/2 \rfloor$, we have the inequality

$$\mathbb{E} \sup_{f \in \widehat{\mathcal{F}}(7\delta')} \frac{1}{\sqrt{|S_2|}} \sum_{j=1}^{|S_2|} (\ell(f(X_j)) - \ell(f_*(X_j)) - P(\ell(f) - \ell(f_*))) \leq C\delta' \sqrt{N},$$

hence it suffices to check that $\Delta_2 = DM_{\Delta_2} \sqrt{7\delta'} \geq C\delta' \sqrt{\frac{N}{k_2}}$. The latter is equivalent to $\delta' \leq CD^2 M_{\Delta_2}^2 \frac{k_2}{N}$ that holds by assumption. Result now follows easily as we assumed that the subsamples S_1 and S_2 used to construct $\widehat{f}_N$ and $\widehat{f}_N'$ are disjoint.

Funding

Stanislav Minsker gratefully acknowledges support by the National Science Foundation [DMS-1712956 and CCF-1908905].

Data Availability Statement

The data underlying this article are available in UCI Machine Learning Repository [52] at <https://archive.ics.uci.edu/ml/datasets/Communities+and+Crime+Unnormalized>. Artificially generated data underlying this article can be generated using the code that is available on github at <https://github.com/TimotheeMathieu/Excess-risk-bounds-in-robust-empirical-risk-minimization/>.

REFERENCES

- [1] ALISTARH, D., ALLEN-ZHU, Z. & LI, J. (2018) Byzantine stochastic gradient descent. in *Advances in Neural Information Processing Systems*, pp. 4613–4623.
- [2] ALON, N., MATIAS, Y. & SZEGEDY, M. (1996) The space complexity of approximating the frequency moments. in *Proceedings of the twenty-eighth annual ACM symposium on Theory of computing*, pp. 20–29. ACM.
- [3] ALQUIER, P., COTTET, V. & LECUÉ, G. (2019) Estimation bounds and sharp oracle inequalities of regularized procedures with Lipschitz loss functions. *Annals of Statistics*, **47**(4), 2117–2144.
- [4] ANTHONY, M. & BARTLETT, P. L. (2009) *Neural network learning: Theoretical foundations*. Cambridge University Press.
- [5] AUDIBERT, J.-Y. & CATONI, O. (2011) Robust linear least squares regression. *The Annals of Statistics*, **39**(5), 2766–2794.
- [6] BAHADUR, R. R. (1966) A note on quantiles in large samples. *The Annals of Mathematical Statistics*, **37**(3), 577–580.
- [7] BARTLETT, P. L., BOUSQUET, O. & MENDELSON, S. (2005) Local Rademacher complexities. *The Annals of Statistics*, **33**(4), 1497–1537.
- [8] BARTLETT, P. L., JORDAN, M. I. & MCAULIFFE, J. D. (2004) Large margin classifiers: convex loss, low noise, and convergence rates. in *Advances in Neural Information Processing Systems*, pp. 1173–1180.
- [9] ——— (2006) Convexity, classification, and risk bounds. *Journal of the American Statistical Association*, **101**(473), 138–156.
- [10] BARTLETT, P. L. & MENDELSON, S. (2006) Empirical minimization. *Probability Theory and Related Fields*, **135**(3), 311–334.
- [11] BARTLETT, P. L., MENDELSON, S. & NEEMAN, J. (2012) ℓ_1 -regularized linear regression: persistence and oracle inequalities. *Probability theory and related fields*, **154**(1-2), 193–224.
- [12] BOUCHERON, S., LUGOSI, G. & MASSART, P. (2013) *Concentration inequalities: A nonasymptotic theory of independence*. Oxford University Press.
- [13] BROWNLEES, C., JOLY, E. & LUGOSI, G. (2015) Empirical risk minimization for heavy-tailed losses. *The Annals of Statistics*, **43**(6), 2507–2536.
- [14] CATONI, O. (2012) Challenging the empirical mean and empirical variance: a deviation study. *Annales de l’Institut Henri Poincaré, Probabilités et Statistiques*, **48**(4), 1148–1185.
- [15] CHEN, L. H. & SHAO, Q.-M. (2001) A non-uniform Berry–Esseen bound via Stein’s method. *Probability theory and related fields*, **120**(2), 236–254.
- [16] CHEN, Y., SU, L. & XU, J. (2017) Distributed statistical machine learning in adversarial settings: Byzantine gradient descent. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, **1**(2), 44.
- [17] CHERAPANAMJERI, Y., HOPKINS, S. B., KATHURIA, T., RAGHAVENDRA, P. & TRIPURANENI, N. (2020) Algorithms for heavy-tailed statistics: Regression, covariance estimation, and beyond. in *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing*, pp. 601–609.

- [18] CHINOT, G., LECUÉ, G. & LERASLE, M. (2019) Robust statistical learning with Lipschitz and convex loss functions. *Probability Theory and related fields*, pp. 1–44.
- [19] DEVROYE, L., LERASLE, M., LUGOSI, G. & OLIVEIRA, R. I. (2016) Sub-Gaussian mean estimators. *The Annals of Statistics*, **44**(6), 2695–2725.
- [20] DIAKONIKOLAS, I., KAMATH, G., KANE, D., LI, J., MOITRA, A. & STEWART, A. (2019) Robust estimators in high-dimensions without the computational intractability. *SIAM Journal on Computing*, **48**(2), 742–864.
- [21] DIAKONIKOLAS, I., KAMATH, G., KANE, D., LI, J., STEINHARDT, J. & STEWART, A. (2019) Sever: A robust meta-algorithm for stochastic optimization. *Proceedings of the International Conference on Machine Learning*, pp. 1596–1606.
- [22] DIAKONIKOLAS, I., KAMATH, G., KANE, D. M., LI, J., MOITRA, A. & STEWART, A. (2017) Being robust (in high dimensions) can be practical. *Proceedings of the International Conference on Machine Learning*, pp. 999–1008.
- [23] DIAKONIKOLAS, I. & KANE, D. M. (2019) Recent advances in algorithmic high-dimensional robust statistics. *arXiv preprint arXiv:1911.05911*.
- [24] DUDLEY, R. M. (2014) *Uniform central limit theorems*, vol. 142. Cambridge University Press.
- [25] HOEFFDING, W. (1948) A Class of Statistics with Asymptotically Normal Distribution. *Annals of Mathematical Statistics*, **19**(3), 293–325.
- [26] ——— (1963) Probability inequalities for sums of bounded random variables. *Journal of the American statistical association*, **58**(301), 13–30.
- [27] HOLLAND, M. J. & IKEDA, K. (2017) Robust regression using biased objectives. *Machine Learning*, **106**(9–10), 1643–1679.
- [28] ——— (2019) Efficient learning with robust gradient descent. *Machine Learning*, **108**(8–9), 1523–1560.
- [29] HOPKINS, S. B. (2020) Mean estimation with sub-Gaussian rates in polynomial time. *Annals of Statistics*, **48**(2), 1193–1213.
- [30] HUBER, P. J. & RONCHETTI, E. M. (2009) *Robust statistics; 2nd ed.*, Wiley Series in Probability and Statistics. Wiley, Hoboken, NJ.
- [31] KOLTCHINSKII, V. (2011) *Oracle inequalities in empirical risk minimization and sparse recovery problems*, vol. 2033 of *Lecture Notes in Mathematics*. Springer, École d’Été de Probabilités de Saint-Flour 2008.
- [32] KOLTCHINSKII, V. & MENDELSON, S. (2015) Bounding the smallest singular value of a random matrix without concentration. *International Mathematics Research Notices*, **2015**(23), 12991–13008.
- [33] LAI, K. A., RAO, A. B. & VEMPALA, S. (2016) Agnostic estimation of mean and covariance. in *2016 IEEE 57th Annual Symposium on Foundations of Computer Science (FOCS)*, pp. 665–674. IEEE.
- [34] LECUÉ, G. & LERASLE, M. (2020) Robust machine learning by median-of-means: theory and practice. *Annals of Statistics*, **48**(2), 906–931.
- [35] LECUÉ, G., LERASLE, M. & MATHIEU, T. (2020) Robust classification via MOM minimization. *Machine Learning*, **109**(8), 1635–1665.
- [36] LECUÉ, G. & MENDELSON, S. (2010) Sharper lower bounds on the performance of the empirical risk minimization algorithm. *Bernoulli*, pp. 605–613.
- [37] LEDOUX, M. & TALAGRAND, M. (1991) *Probability in Banach spaces: isoperimetry and processes*. Springer-Verlag, Berlin.
- [38] LERASLE, M. & OLIVEIRA, R. I. (2011) Robust empirical mean estimators. *arXiv preprint arXiv:1112.3914*.
- [39] LUGOSI, G. & MENDELSON, S. (2019a) Regularization, sparse recovery, and median-of-means tournaments. *Bernoulli*, **25**(3), 2075–2106.
- [40] ——— (2019b) Risk minimization by median-of-means tournaments. *Journal of the European Mathematical Society*, **22**(3), 925–965.
- [41] MATHIEU, T. & MINSKER, S. (2021) Excess risk bounds in robust empirical risk minimization: Supplementary Material. *Information and Inference: A Journal of the IMA*.
- [42] MAYZLIN, D., DOVER, Y. & CHEVALIER, J. (2014) Promotional reviews: An empirical investigation of

- online review manipulation. *American Economic Review*, **104**(8), 2421–55.
- [43] MENDELSON, S. (2008) Lower bounds for the empirical minimization algorithm. *IEEE Transactions on Information Theory*, **54**(8), 3797–3803.
 - [44] ——— (2014) Learning without concentration. in *Conference on Learning Theory*, pp. 25–39.
 - [45] ——— (2016) Upper bounds on product and multiplier empirical processes. *Stochastic Processes and their Applications*, **126**(12), 3652–3680.
 - [46] MINSKER, S. (2019a) Distributed statistical estimation and rates of convergence in normal approximation. *Electronic Journal of Statistics*, **13**(2), 5213–5252.
 - [47] ——— (2019b) Uniform bounds for robust mean estimators. *arXiv preprint arXiv:1812.03523*.
 - [48] NEMIROVSKI, A. & YUDIN, D. (1983) *Problem complexity and method efficiency in optimization*. John Wiley & Sons.
 - [49] O'DONNELL, R. (2014) *Analysis of Boolean functions*. Cambridge University Press.
 - [50] PEDREGOSA, F., VAROQUAUX, G., GRAMFORT, A., MICHEL, V., THIRION, B., GRISEL, O., BLONDEL, M., PRETTENHOFER, P., WEISS, R. & DUBOURG, V. (2011) Scikit-learn: Machine learning in Python. *Journal of machine learning research*, **12**(Oct), 2825–2830.
 - [51] PRASAD, A., SUGGALA, A. S., BALAKRISHNAN, S. & RAVIKUMAR, P. (2020) Robust estimation via robust gradient estimation. *Journal of the Royal Statistical Society Series B*, **82**(3), 601–627.
 - [52] REDMOND, M. (2011) Communities and Crime Unnormalized Data Set. Creating using 1990 US Census and other sources. Available at <https://archive.ics.uci.edu/ml/datasets/Communities+and+Crime+Unnormalized>.
 - [53] RON SCHMELZER, F. (2019) The Achilles' heel Of AI. <https://www.forbes.com/sites/cognitiveworld/2019/03/07/the-achilles-heel-of-ai>.
 - [54] TALAGRAND, M. (2014) *Upper and lower bounds for stochastic processes: modern methods and classical problems*, vol. 60. Springer Science & Business Media.
 - [55] TSYBAKOV, A. B. (2004) Optimal aggregation of classifiers in statistical learning. *The Annals of Statistics*, **32**(1), 135–166.
 - [56] VAN DE GEER, S. (2016) Estimation and testing under sparsity. *Lecture Notes in Mathematics*, **2159**.
 - [57] VAN DE GEER, S. A. & VAN DE GEER, S. (2000) *Empirical Processes in M-estimation*, vol. 6. Cambridge University Press.
 - [58] VAN DER VAART, A. W. & WELLNER, J. A. (2000) *Weak convergence and empirical Processes: with applications to statistics*. Springer.
 - [59] YIN, D., CHEN, Y., KANNAN, R. & BARTLETT, P. (2018) Byzantine-robust distributed learning: towards optimal statistical rates. in *International Conference on Machine Learning*, pp. 5650–5659.