

RISK BOUNDS WHEN LEARNING INFINITELY MANY RESPONSE FUNCTIONS BY ORDINARY LINEAR REGRESSION

Plassier, V., Portier, F. And J. Segers

DISCUSSION PAPER | 2020 / 19

RISK BOUNDS WHEN LEARNING INFINITELY MANY RESPONSE FUNCTIONS BY ORDINARY LINEAR REGRESSION

Vincent Plassier

LIDAM/ISBA, UCLouvain
and ENS Paris-Saclay
1348 Louvain-la-Neuve, Belgium
vincent.plassier@ens-paris-saclay.fr

Francois Portier

LTCl, Télécom Paris
Institut polytechnique de Paris
91120 Palaiseau, France
francois.portier@gmail.com

Johan Segers

LIDAM/ISBA, UCLouvain
1348 Louvain-la-Neuve, Belgium
johan.segers@uclouvain.be

June 17, 2020

ABSTRACT

Consider the problem of learning a large number of response functions simultaneously based on the same input variables. The training data consist of a single independent random sample of the input variables drawn from a common distribution together with the associated responses. The input variables are mapped into a high-dimensional linear space, called the feature space, and the response functions are modelled as linear functionals of the mapped features, with coefficients calibrated via ordinary least squares. We provide convergence guarantees on the worst-case excess prediction risk by controlling the convergence rate of the excess risk uniformly in the response function. The dimension of the feature map is allowed to tend to infinity with the sample size. The collection of response functions, although potentially infinite, is supposed to have a finite Vapnik–Chervonenkis dimension. The bound derived can be applied when building multiple surrogate models in a reasonable computing time.

1 Introduction

Context. When the outcome of interest is generated by a black box model which cannot be easily evaluated, a well-spread technique is to build a surrogate model allowing to reproduce the behavior of the true model while being computationally cheaper. This approach, known as *response surface*, is popular in many fields of engineering such as *reliability analysis* (Bucher and Bourgund, 1990), *aerospace science* (Forrester and Keane, 2009), *energy science* (Nguyen et al., 2014), or electromagnetic dosimetry (Azzi et al., 2019), to name a few. In addition, the response surface methodology is useful in applied mathematics, for instance, in *optimization* when the objective function is difficult to evaluate (Jones, 2001) and in *Monte Carlo integration*, where a surrogate function can be used to reduce the variance of the Monte Carlo estimate with the help of control variates (Portier and Segers, 2019). For a general presentation of the response surface methodology, we refer to Myers et al. (2016).

Framework. The statistical framework of a response surface is as follows. Consider a real-valued function f defined on a state space $\mathcal{X}$, that is, $f : \mathcal{X} \rightarrow \mathbb{R}$. In many applications, the function f is a black box and difficult to evaluate. For instance, a request to f might be obtained by running a heavy computer program. In such a situation, one can afford only a few requests to f , that is, one can obtain $f(X_1), \dots, f(X_n)$, where $n \in \mathbb{N}$ and $X_1, \dots, X_n$ is an independent sample of $\mathcal{X}$ -valued random variables with distribution P , called the inputs or the covariates. Based on those evaluations, the goal is to build a *surrogate model*, that is, an approximation of f which is easier to calculate than f itself.

Many different methods can be used to build a surrogate model. The simplest one consists in learning f as a linear combination of the covariates by minimizing the sum of squared errors. A well spread extension is to fit a polynomial function instead of a linear one as presented in Myers et al. (2016) or Konakli and Sudret (2016). Since building a response surface consists in the same task as regression, any regression method might be used in principle. Popular methods include moving least-squares (Breitkopf et al., 2005), Gaussian processes (Frean and Boyle, 2008) or neural nets (Bauer et al., 2019). For surrogate models, the approximation method needs to be sufficiently flexible to fit the black box function f well and simple enough to require only a small amount of computations. It is perhaps due to its connection to the regression framework that the problem of building surrogate models has received little specific attention in the statistical learning literature. Our purpose is to address the question of learning many surrogate models simultaneously from a single, random design.

Learning several models simultaneously. The fundamental question raised in this paper deals with the ability of building several, possibly infinitely many, surrogate models such that (a) they share the same quality and (b) they are constructed with the help of a single input sample. From a theoretical standpoint, it relates to the question of *uniformity* over the tasks: can a broad family of models be learnt with a uniform level of accuracy? From a more practical point of view, by working with the same inputs to solve multiple tasks simultaneously, one benefits from a certain computational advantage, as explained below.

Consider a broad class $\mathcal{F}$ of black box models f . For each such model f , a least-squares estimate is obtained out of the linear span of the components of the d -dimensional feature map $h = (h_1, \dots, h_d)^\top : \mathcal{X} \rightarrow \mathbb{R}^d$, where $^\top$ denotes matrix transposition. The feature map h should be known and easy to evaluate. Define the ordinary least squares estimate

$$\hat{\beta}_f \in \arg \min_{b \in \mathbb{R}^d} \sum_{i=1}^n \{f(X_i) - h(X_i)^\top b\}^2.$$

The surrogate model for f is then defined as $x \mapsto \hat{f}(x) = h(x)^\top \hat{\beta}_f$.

This approach is known as *series estimators* or simply as *least squares estimators* (Härdle, 1990; Györfi et al., 2006). It is quite general as several basis functions might be considered such as polynomials, indicators, spline functions or the Fourier basis. In the regression framework, series estimators have been studied for instance in Newey (1997) and Belloni et al. (2015). Series estimators are a convenient way to include shape constraints on f , as for instance when f is partially linear. They also facilitate the computation of derivatives (Zhou and Wolfe, 2000). In this respect, series estimators can help to build *easy-to-interpret* models, a desirable feature when one is in need of some knowledge about the effects of certain inputs on the black box model f .

One motivation for the use of series estimators is the computational advantage they provide when several models are to be learnt at the same time. When a large number m of such surrogate models f are built with different covariates, running m least-squares algorithms, for instance using the Cholesky decomposition, requires $O(mnd^2)$ operations (Friedman et al., 2001, Section 3.5), assuming that $d = O(n)$. In our framework of a single training sample, however, the Cholesky decomposition needs to be done only once, and the time needed to compute m least squares estimates is rather $O(nd^2 + mnd)$.

Uniform convergence rate in random design with increasing dimension. For a given model $f \in \mathcal{F}$, the error made by a surrogate model $x \mapsto h(x)^\top b$ with coefficient vector $b \in \mathbb{R}^d$ is measured

through the $L^2(P)$ -risk

$$L_f(b) = \int_{\mathcal{X}} \{f(x) - h(x)^\top b\}^2 dP(x) = \mathbb{E}[\{f(X) - h(X)^\top b\}^2],$$

where the $\mathcal{X}$ -valued random variable X has distribution P . We place ourselves in the random design setting and study the risk $L_f(\hat{\beta}_f)$ of the least squares estimator $\hat{\beta}_f$. This risk is a random variable whose randomness stems from the one of the training sample $X_1, \dots, X_n$.

The main result of the paper concerns the *excess risk* $L_f(\hat{\beta}_f) - \min_{b \in \mathbb{R}^d} L_f(b)$ when $n \rightarrow \infty$ and $d \rightarrow \infty$, uniformly over $f \in \mathcal{F}$. That is, we study the convergence rate to zero of the random variable $\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_f) - \min_{b \in \mathbb{R}^d} L_f(b)\}$. A key quantity is the *leverage function* $q : \mathcal{X} \rightarrow [0, \infty)$ defined by

$$\forall x \in \mathcal{X}, \quad q(x) = h(x)^\top G^{-1} h(x),$$

where $G = \mathbb{E}[h(X)h(X)^\top]$ is the $d \times d$ Gram matrix of the feature map. The leverage function is the population version of the *statistical leverage* of a feature vector $h(X_i)$ in the linear regression model. It plays an important role when analyzing regression with random design (Hsu et al., 2014). Note that q does not change if the feature map is composed with an invertible linear transformation. Let $\varepsilon_f = f - h^\top \beta_f$ be the error function, with $\beta_f = \arg \min_{b \in \mathbb{R}^d} L_f(b)$ the risk-minimizing coefficient vector. Our main result, expressed in Corollary 1, is that

$$\sup_{f \in \mathcal{F}} \left\{ L_f(\hat{\beta}_f) - \min_{b \in \mathbb{R}^d} L_f(b) \right\} = O_{\mathbb{P}} \left(\frac{\log n}{n} \sup_{f \in \mathcal{F}} \mathbb{E}[q(X) \varepsilon_f^2(X)] \right), \quad n \rightarrow \infty.$$

Apart from the fact that the obtained bound is invariant under invertible linear transformations of the feature map, this result is remarkable for the three following reasons. First, the dimension d of the feature space is allowed to tend to infinity with the input sample size n at a speed which depends on the leverage function q via Condition 1. Second, in case the class $\mathcal{F}$ contains only a single response function f , a simple analysis leads to a bound for the excess risk that scales as $\mathbb{E}[q(X) \varepsilon_f^2(X)]/n$, see Eq. (6), a bound that matches the one of our main result up to a logarithmic term. Third, the quantity $\mathbb{E}[q(X) \varepsilon_f^2(X)]$ takes over the role of the quantity $\sigma_f^2 d$ in the fixed-design setting, where σ_f^2 is the variance of the error variable in the linear model.

The uniformity in $f \in \mathcal{F}$ is achieved by a decomposition of a quadratic form in terms of a sample mean and a U-statistic in combination with a concentration inequality for the suprema of such statistics. The analysis is focused on the ordinary least squares estimator, whereas the extension to ridge regression as in Hsu et al. (2014) is left for further research.

Paper outline. The mathematical background is presented in Section 2. The main result and a sketch of its proof are presented in Sections 3 and 4, respectively. Detailed proofs are deferred to the appendices.

2 Learning multiple response functions simultaneously

Linear model and ordinary least squares estimator. Consider a collection $\mathcal{F}$ of functions $f : \mathcal{X} \rightarrow \mathbb{R}$ on a probability space $(\mathcal{X}, \mathcal{A}, P)$. We think of $f(x)$ as the real-valued response given an input $x \in \mathcal{X}$. Given an independent random sample $X_1, \dots, X_n$ from P together with the associated responses $f(X_1), \dots, f(X_n)$ for every $f \in \mathcal{F}$, we wish to learn the values $f(x)$ of the response functions $f \in \mathcal{F}$ for new but yet unobserved inputs $x \in \mathcal{X}$. To this end, we map the input space $\mathcal{X}$ into a feature space $\mathbb{R}^d$ via a feature map $h : \mathcal{X} \rightarrow \mathbb{R}^d : x \mapsto h(x) = (h_1(x), \dots, h_d(x))^\top$. One of the feature functions h_j could be the constant function 1, corresponding to an intercept. The response functions are modelled as linear functionals of the mapped features with coefficients estimated by ordinary least squares. The approximation to the response function $f \in \mathcal{F}$ is thus

$$\forall x \in \mathcal{X}, \quad \hat{f}(x) = h(x)^\top \hat{\beta}_f \quad \text{where} \quad \hat{\beta}_f = \arg \min_{b \in \mathbb{R}^d} \sum_{i=1}^n \{f(X_i) - h(X_i)^\top b\}^2.$$

We wish to control the learning error $\hat{f} - f$ uniformly in $f \in \mathcal{F}$.

Sharing the same inputs and the same feature map across multiple response functions brings computational gains. Classical least-squares theory yields

$$\forall f \in \mathcal{F}, \quad \hat{\beta}_f = G_n^{-1} \frac{1}{n} \sum_{i=1}^n h(X_i) f(X_i) \quad \text{where} \quad G_n = \frac{1}{n} \sum_{i=1}^n h(X_i) h(X_i)^\top.$$

In case the empirical Gram matrix G_n is not invertible, a pseudo-inverse is used instead. It follows that the predicted responses are linear in the observed responses:

$$\forall f \in \mathcal{F}, x \in \mathcal{X}, \quad \hat{f}(x) = \frac{1}{n} \sum_{i=1}^n w(x, X_i) f(X_i) \quad \text{where} \quad w(x, X_i) = h(x)^\top G_n^{-1} h(X_i).$$

The weights $w(x, X_i)$ do not depend on the response function f . This invariance is a computational advantage if multiple response functions $f \in \mathcal{F}$ are to be learned simultaneously.

Worst-case excess prediction risk. The response functions are modelled as linear functionals of the mapped features via

$$\forall f \in \mathcal{F}, \forall x \in \mathcal{X}, \quad f(x) = h(x)^\top \beta_f + \varepsilon_f(x). \quad (1)$$

The coefficient vector β_f is defined as the minimizer over $b \in \mathbb{R}^d$ of the prediction risk

$$\forall f \in \mathcal{F}, \forall b \in \mathbb{R}^d, \quad L_f(b) = \mathbb{E}[\{f(X) - h(X)^\top b\}^2].$$

Here, the expectation is taken with respect to a random feature X with distribution P and it is assumed that f and h have finite second moments. Let $G = \mathbb{E}[h(X)h(X)^\top]$ denote the $d \times d$ Gram matrix of the feature map, assumed to be positive definite. Classical least squares theory yields

$$\forall f \in \mathcal{F}, \quad \beta_f = \arg \min_{b \in \mathbb{R}^d} L_f(b) = G^{-1} \mathbb{E}[h(X)f(X)].$$

Since $\varepsilon_f(X) = f(X) - h(X)^\top \beta_f$ is orthogonal to $h(X)$, i.e., $\mathbb{E}[h(X)\varepsilon_f(X)] = 0$, the *excess risk* associated to any other coefficient vector $b \in \mathbb{R}^d$ is

$$L_f(b) - L_f(\beta_f) = \mathbb{E}[\{h(X)^\top(b - \beta_f)\}^2] = (b - \beta_f)^\top G (b - \beta_f).$$

The optimal coefficient vector β_f is unknown, so we estimate it by the least-squares estimator $\hat{\beta}_f$. The expected squared error made by the approximated response function for a new input distributed according to P is

$$\int_{\mathcal{X}} \{\hat{f}(x) - f(x)\}^2 dP(x) = L_f(\hat{\beta}_f) = L_f(\beta_f) + (\hat{\beta}_f - \beta_f)^\top G (\hat{\beta}_f - \beta_f). \quad (2)$$

The right-hand side of (2) decomposes the expected squared error into two parts.

- The first term, $L_f(\beta_f)$, is deterministic. It represents the *modelling error* stemming from the linear model in (1). Given the model, this term is incompressible and does not depend on the learning algorithm nor on the training data.
- The second term on the right-hand side in (2) is random. It represents the *learning error* due to the estimation step and the randomness of the training data.

It is on the learning error that we focus our analysis, with the particularity that we consider the error uniformly in the response function. Our object of interest is thus the *worst-case excess prediction risk*

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_f) - L_f(\beta_f)\} = \sup_{f \in \mathcal{F}} (\hat{\beta}_f - \beta_f)^\top G (\hat{\beta}_f - \beta_f). \quad (3)$$

We are interested in the rate at which this supremum tends to zero as the sample size n and the dimension d of the feature map tend to infinity.

It is to be emphasized that our setting is that of a random design. The excess prediction risk $L_f(\hat{\beta}_f) - L_f(\beta_f)$ is a nonnegative random variable that is constructed out of the random training sample $X_1, \dots, X_n$. As is clear from (2), it incorporates the risk associated to a new and yet unobserved input $x \in \mathcal{X}$, averaged over P . The expression in (3) is thus a supremum over potentially infinitely many random variables, each variable being built on the same inputs.

Excess prediction risk of a single response. Fix a response function $f \in \mathcal{F}$. Let P_n denote the empirical distribution of the training sample $X_1, \dots, X_n$, assigning probability $1/n$ to each observed input. For a real-valued, vector-valued or matrix-valued function g on $\mathcal{X}$, expectations with respect to the unknown sampling distribution P and the empirical distribution P_n are denoted respectively by

$$P(g) = \mathbb{E}[g(X)] = \int_{\mathcal{X}} g(x) dP(x), \quad P_n(g) = \frac{1}{n} \sum_{i=1}^n g(X_i).$$

Both operators are linear: for instance, $P(Ag) = AP(g)$ and $P_n(Ag) = AP_n(g)$ if A is a linear map between Euclidean spaces of suitable dimension. Using this operator notation, we get

$$\begin{aligned} G &= P(hh^\top), & \beta_f &= G^{-1}P(hf), \\ G_n &= P_n(hh^\top), & \hat{\beta}_f &= G_n^{-1}P_n(hf). \end{aligned}$$

Since $f = h^\top \beta_f + \varepsilon_f$, the estimated coefficient vector is

$$\hat{\beta}_f = G_n^{-1}P_n[h(h^\top \beta_f + \varepsilon_f)] = \beta_f + G_n^{-1}P_n(h\varepsilon_f).$$

The excess prediction risk is thus

$$L_f(\hat{\beta}_f) - L_f(\beta_f) = P_n(h\varepsilon_f)^\top G_n^{-1} G G_n^{-1} P_n(h^\top \varepsilon_f). \quad (4)$$

The predicted response functions $\hat{f}$ are linear combinations of the components $h_1, \dots, h_d$ of the feature map h . They only depend on the feature map h through the linear span of the functions $h_1, \dots, h_d$. Therefore, the predicted responses remain unchanged if we compose the feature map with an invertible linear transformation A of $\mathbb{R}^d$. Let $G^{1/2}$ be the unique symmetric square root matrix of G and let $G^{-1/2}$ be its inverse. The whitened feature map is $\tilde{h} = G^{-1/2}h : \mathcal{X} \rightarrow \mathbb{R}^d$. If the random input X has distribution P , then $\mathbb{E}[\tilde{h}(X)\tilde{h}(X)^\top] = P(\tilde{h}\tilde{h}^\top) = I_d$, the $d \times d$ identity matrix. The empirical Gram matrix of the whitened feature map is

$$P_n(\tilde{h}\tilde{h}^\top) = G^{-1/2}P_n(hh^\top)G^{-1/2} = G^{-1/2}G_nG^{-1/2}.$$

Since $h = G^{1/2}\tilde{h}$ and $P_n(\tilde{h}\tilde{h}^\top)^{-1} = G^{1/2}G_n^{-1}G^{1/2}$, the excess prediction risk in (4) becomes

$$L_f(\hat{\beta}_f) - L_f(\beta_f) = |P_n(\tilde{h}\tilde{h}^\top)^{-1}P_n(\tilde{h}\varepsilon_f)|_2^2, \quad (5)$$

where $|y|_2 = (y^\top y)^{1/2}$ denotes the Euclidean norm of a vector $y \in \mathbb{R}^d$.

Since $P(\tilde{h}\tilde{h}^\top) = I_d$, it is reasonable to expect that $P_n(\tilde{h}\tilde{h}^\top)^{-1}$ is approximately equal to I_d , at least if n is large and d is not too large compared to n ; see Lemma 1 below for a precise statement. In that case, the excess prediction risk in (5) is approximately equal to $|P_n(\tilde{h}\varepsilon_f)|_2^2$. The expectation of the latter random variable can be easily calculated: since $P(\tilde{h}\varepsilon_f) = 0$, the terms with $i \neq j$ in the double sum below vanish and we find

$$\mathbb{E}[|P_n(\tilde{h}\varepsilon_f)|_2^2] = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \mathbb{E}[\varepsilon_f(X_i)\tilde{h}(X_i)^\top \tilde{h}(X_j)\varepsilon_f(X_j)] = \frac{1}{n} P(\tilde{h}^\top \tilde{h}\varepsilon_f^2). \quad (6)$$

The excess prediction risk for a single response function f can thus be expected to have an order of magnitude equal to $n^{-1}P(\tilde{h}^\top \tilde{h}\varepsilon_f^2)$. In comparison, in the fixed-design case, where the training sample is considered as non-random, the expected excess risk is equal to $\sigma_f^2 d/n$, where σ_f^2 is the error variance (Hsu et al., 2014). Eq. (6) motivates why in Theorem 1, the convergence rate for the worst-case prediction risk over the whole response family $\mathcal{F}$ involves the quantity $\sup_{f \in \mathcal{F}} P(\tilde{h}^\top \tilde{h}\varepsilon_f^2)$.

3 Convergence rate of the worst-case excess prediction risk

Notation. Consider an asymptotic setting where the size n of the training sample tends to infinity. The feature map may change with n : with a slight change of notation, we write henceforth

$$h_n = (h_{n,1}, \dots, h_{n,d_n})^\top : \mathcal{X} \rightarrow \mathbb{R}^{d_n}.$$

The feature dimension $d_n \geq 1$ depends on n and may tend to infinity. The whitened feature map is

$$\tilde{h}_n = P(h_n h_n^\top)^{-1/2} h_n : \mathcal{X} \rightarrow \mathbb{R}^{d_n},$$

where the $d_n \times d_n$ Gram matrix $P(h_n h_n^\top)$ is supposed to be invertible; otherwise, we can omit some components $h_{n,j}$ without affecting the linear span of the component functions. The least squares coefficient vectors and modelling errors depend on n as well: we write

$$\forall f \in \mathcal{F}, \forall x \in \mathcal{X}, \quad f(x) = h_n(x)^\top \beta_{n,f} + \varepsilon_{n,f}(x) \quad \text{where} \quad \beta_{n,f} = P(h_n h_n^\top)^{-1} P(h_n f).$$

The least squares estimator of $\beta_{n,f}$ is $\hat{\beta}_{n,f} = P_n(h_n h_n^\top)^{-1} P_n(h_n f)$.

In this setting, the *leverage function* $q_n : \mathcal{X} \rightarrow [0, \infty)$ is defined by

$$\forall x \in \mathcal{X}, \quad q_n(x) = h_n(x)^\top P(h_n h_n^\top)^{-1} h_n(x) = \tilde{h}_n(x)^\top \tilde{h}_n(x) = |\tilde{h}_n(x)|_2^2.$$

The name of q_n is derived from the notion of leverage of a design point in multiple linear regression. Note that q_n does not change if the feature map is composed with an invertible linear transformation. We always have $P(q_n) = \text{tr}[P(\tilde{h}_n \tilde{h}_n^\top)] = d_n$, where $\text{tr}(A)$ denotes the trace of a square matrix A .

Let $\mathbb{P}$ denote the probability measure on the probability space on which the random inputs $X_1, \dots, X_n$, taking values in $\mathcal{X}$, are defined. For any sequence $(Y_n)_n$ of real-valued random variables on that space and for any positive sequence $(a_n)_n$, the expression $Y_n = O_{\mathbb{P}}(a_n)$ as $n \rightarrow \infty$ signifies that Y_n/a_n is bounded in probability, that is, for every $\epsilon > 0$ there exists $K > 0$ such that $\limsup_{n \rightarrow \infty} \mathbb{P}(|Y_n| > a_n K) \leq \epsilon$. Similarly, the expression $Y_n = o_{\mathbb{P}}(a_n)$ as $n \rightarrow \infty$ signifies that Y_n/a_n converges to zero in probability, that is, $\lim_{n \rightarrow \infty} \mathbb{P}(|Y_n| > a_n \epsilon) = 0$ for every $\epsilon > 0$. Our aim is to determine a positive sequence a_n such that $a_n \rightarrow 0$ and, under reasonable assumptions,

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_{n,f}) - L_f(\beta_{n,f})\} = O_{\mathbb{P}}(a_n), \quad n \rightarrow \infty.$$

Lastly, the supremum norm of a function $g : \mathcal{X} \rightarrow \mathbb{R}$ is denoted by $\|g\|_\infty = \sup_{x \in \mathcal{X}} |g(x)|$.

Conditions. The conditions under which the main result holds concern the leverage function q_n and the response functions $\mathcal{F}$.

Condition 1. *One of the two alternative conditions hold:*

(a) $P(q_n^2) = o(n)$ and $\log(\|q_n\|_\infty) = O(\log n)$ as $n \rightarrow \infty$;

(b) $\|q_n\|_\infty \log(2d_n) = o(n)$ as $n \rightarrow \infty$.

Since $P(q_n) = d_n$ and $P(q_n^2) \geq [P(q_n)]^2$, condition (a) implies that $d_n = o(n^{1/2})$ as $n \rightarrow \infty$. As $P(q_n^2) \leq \|q_n\|_\infty P(q_n) = \|q_n\|_\infty d_n$, a sufficient condition for (a) moreover is that $\|q_n\|_\infty = o(n/d_n)$, which is the leverage condition in Portier and Segers (2019). Here, we have just assumed that the speed at which $\|q_n\|_\infty$ tends to infinity is at most polynomial in n . Condition (b) implies that $d_n \log(2d_n) = o(n)$ as $n \rightarrow \infty$ but, compared to (a), imposes a stronger condition on $\|q_n\|_\infty$.

Condition 2. *The collection $\mathcal{F}$ of response functions admits a uniformly bounded envelope $F : \mathcal{X} \rightarrow [0, \infty)$, i.e., $|f(x)| \leq F(x)$ for any $f \in \mathcal{F}$ and any $x \in \mathcal{X}$, and $\|F\|_\infty$ is finite. In addition, $\mathcal{F}$ is suppose to be at most countably infinite.*

As $\|q_n\|_\infty$ and $\|F\|_\infty$ are finite, the collection of error functions $\{\varepsilon_{n,f} : f \in \mathcal{F}\}$ is uniformly bounded, see Lemma 2 in the appendices.

The assumption in Condition 2 that $\mathcal{F}$ is countable assures that suprema over random variables indexed by $f \in \mathcal{F}$ are measurable. Otherwise, probabilities involving such suprema would need to be replaced by outer probabilities (van der Vaart and Wellner, 1996, Part 1). In practice, the countability assumption is harmless insofar as $\mathcal{F}$ can usually be approximated by a countable dense subfamily anyway without affecting the value of the supremum (van der Vaart and Wellner, 1996, Section 2.3.3).

Covering numbers capture the complexity of a subset of a metric space and play a central role in a number of areas in information theory and statistics, including nonparametric function estimation, density estimation, empirical processes, and machine learning.

Definition 1 (Covering number). For a subset $\mathcal{F}$ of a metric space $(\mathcal{Y}, \rho)$, the η -covering number $\mathcal{N}(\mathcal{F}, \rho, \eta)$ is the smallest number of open ρ -balls of radius $\eta > 0$ required to cover $\mathcal{F}$, i.e.,

$$\mathcal{N}(\mathcal{F}, \rho, \eta) = \min \left\{ p \geq 1 : \exists f_1, \dots, f_p \in \mathcal{Y}, \mathcal{F} \subset \bigcup_{i=1}^p B_\rho(f_i, \eta) \right\},$$

where $B_\rho(f, \eta) = \{g \in \mathcal{Y} : \rho(g, f) < \eta\}$ for $f \in \mathcal{Y}$ and $\eta > 0$.

Definition 2 (VC-class). A class $\mathcal{F}$ of real functions on a measurable space $(\mathcal{X}, \mathcal{A})$ is called a VC-class (Vapnik–Chervonenkis) of parameters $(v, A) \in (0, \infty)^2$ with respect to the envelope F if for any $0 < \eta < 1$ and any probability measure Q on $(\mathcal{X}, \mathcal{A})$, we have

$$\mathcal{N}(\mathcal{F}, L^2(Q), \eta \|F\|_{L^2(Q)}) \leq (A/\eta)^v.$$

In Definition 2, we view $\mathcal{F}$ as a subset of the metric space $L^2(Q) \equiv L^2(\mathcal{X}, \mathcal{A}, Q)$ of Q -square-integrable functions $f : \mathcal{X} \rightarrow \mathbb{R}$ equipped with the metric $\rho(f, g) = \|f - g\|_{L^2(Q)}$, where $\|h\|_{L^2(Q)} = [Q(h^2)]^{1/2}$ for measurable $h : \mathcal{X} \rightarrow \mathbb{R}$.

Condition 3. With respect to the envelope F , the collection $\mathcal{F}$ is VC with parameters (v, A) .

Main result. The maximal error standard deviation is

$$\sigma_n = \sup_{f \in \mathcal{F}} [P(\varepsilon_{n,f}^2)]^{1/2},$$

which under Condition 2 is bounded by $[P(F^2)]^{1/2}$ since $P(\varepsilon_{n,f}^2) \leq P(f^2) \leq P(F^2)$ for every $f \in \mathcal{F}$ by orthogonality of $\varepsilon_{n,f}$ with h_n . The growth rate of the worst-case excess prediction risk will be expressed in terms of

$$\gamma_n = \sup_{f \in \mathcal{F}} [P(q_n \varepsilon_{n,f}^2)]^{1/2}, \quad L_n = \|q_n\|_\infty^{1/2} \sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty. \quad (7)$$

Clearly, $\gamma_n \leq L_n$. Write $a \vee b = \max(a, b)$ for real a and b .

Theorem 1 (Convergence rate of worst-case excess prediction risk). If Conditions 1, 2, and 3 hold and if $\log(1/\sigma_n) = O(\log n)$ as $n \rightarrow \infty$, then

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_f) - L_f(\beta_f)\} = O_{\mathbb{P}} \left(\left(\gamma_n^2 \vee \frac{L_n^2 \log n}{\sqrt{n}} \right) \frac{\log n}{n} \right), \quad n \rightarrow \infty.$$

The quantity γ_n^2 is related to the expectation of the functions $q_n \varepsilon_{n,f}^2$ while L_n^2 is related to their supremum norm. It is reasonable to hope that the latter will not be too large in comparison to the former. This motivates the growth rate (8) assumed in the next corollary.

Corollary 1. In Theorem 1, if, in addition, also

$$\|q_n\|_\infty \sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}^2\|_\infty = O \left(\frac{\sqrt{n}}{\log n} \sup_{f \in \mathcal{F}} P(q_n \varepsilon_{n,f}^2) \right), \quad n \rightarrow \infty, \quad (8)$$

then

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_f) - L_f(\beta_f)\} = O_{\mathbb{P}} \left(\frac{\log n}{n} \sup_{f \in \mathcal{F}} P(q_n \varepsilon_{n,f}^2) \right), \quad n \rightarrow \infty. \quad (9)$$

Apart from the factor $\log n$, the convergence rate in (9) corresponds to the one for a single response function f as discussed in the paragraph around Eq. (6). The additional factor $\log n$ stems from a concentration inequality for U-statistics in combination with bounds on the covering numbers of function classes derived from $\mathcal{F}$.

4 Sketch of proof of Theorem 1

Let $\lambda_{\min}(A)$ and $\lambda_{\max}(A)$ denote the smallest and largest eigenvalue, respectively, of the symmetric matrix A . Consider the matrix norm $|A|_2 = \sup\{|Ay|_2 : y \in \mathbb{R}^d, |y|_2 \leq 1\}$ for $A \in \mathbb{R}^{d \times d}$. If A is symmetric and positive semi-definite, then $|A|_2 = \lambda_{\max}(A)$. If, moreover, A is positive definite, then $\lambda_{\max}(A^{-1}) = \{\lambda_{\min}(A)\}^{-1}$. In view of (5), the worst-case excess prediction risk is bounded by

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_f) - L_f(\beta_f)\} \leq [\lambda_{\min}\{P_n(\bar{h}_n \bar{h}_n^\top)\}]^{-2} \cdot \sup_{f \in \mathcal{F}} |P_n(\bar{h}_n \varepsilon_{n,f})|_2^2. \quad (10)$$

Under reasonable conditions permitting $d_n \rightarrow \infty$, the smallest eigenvalue of $P_n(\bar{h}_n \bar{h}_n^\top)$ remains bounded away from zero with high probability.

Lemma 1. *Suppose one of the following two conditions holds:*

- (a) $P(q_n^2) = o(n)$ as $n \rightarrow \infty$;
- (b) $\|q_n\|_\infty \log(2d_n) = o(n)$ as $n \rightarrow \infty$.

Then $P_n(\bar{h}_n \bar{h}_n^\top)$ is invertible with probability tending to one and $\lambda_{\min}\{P_n(\bar{h}_n \bar{h}_n^\top)\} \geq 1 + o_{\mathbb{P}}(1)$ as $n \rightarrow \infty$.

The proof of Lemma 1 in case (a) builds upon Portier and Segers (2019, Lemmas 2 and 3), while the one in case (b) is based upon Leluc et al. (2019, Lemma A.2), relying on a matrix Chernoff inequality due to Tropp (2015, Theorem 5.1.1). The proof is given in Section A in the appendices.

In view of the bound (10) in combination with Lemma 1, it is sufficient to show the claimed convergence rate with the excess risk $L_f(\hat{\beta}_f) - L_f(\beta_f)$ replaced by $|P_n(\bar{h}_n \varepsilon_{n,f})|_2^2$. For $f \in \mathcal{F}$, define $g_{n,f} : \mathcal{X}^2 \rightarrow \mathbb{R}$ by

$$\forall (x, y) \in \mathcal{X}^2, \quad g_{n,f}(x, y) = \varepsilon_{n,f}(x) \bar{h}_n(x)^\top \bar{h}_n(y) \varepsilon_{n,f}(y).$$

Note that $g_{n,f}(x, x) = q_n(x) \varepsilon_{n,f}^2(x)$ for $x \in \mathcal{X}$. The quantity of interest is bounded by

$$n^2 |P_n(\bar{h}_n \varepsilon_{n,f})|_2^2 \leq n P(q_n \varepsilon_{n,f}^2) + \left| \sum_{i=1}^n \{(q_n \varepsilon_{n,f}^2)(X_i) - P(q_n \varepsilon_{n,f}^2)\} \right| + \left| \sum_{1 \leq i \neq j \leq n} g_{n,f}(X_i, X_j) \right|. \quad (11)$$

We will bound the supremum over $f \in \mathcal{F}$ of the sum on the right-hand side of (11) by the sum of the suprema of the three terms.

- The first supremum is just $n \gamma_n^2$, which is dominated by the stated convergence rate.
- The second supremum is a supremum over a sum of independent and identically distributed random variables with mean zero and finite variance. We will find its rate of convergence to zero via Proposition 2 below.
- The third supremum involves a U-statistic of order two with a degenerate kernel: by orthogonality of $\bar{h}_n$ and $\varepsilon_{n,f}$, we have $\mathbb{E}[g_{n,f}(X, x)] = \mathbb{E}[\varepsilon_{n,f}(X) \bar{h}_n(X)^\top] \bar{h}_n(x) \varepsilon_{n,f}(x) = 0$ and similarly $\mathbb{E}[g_{n,f}(x, X)] = 0$ for every $x \in \mathcal{X}$. We determine its convergence rate via a concentration inequality due to Major (2006) quoted as Theorem 2 in the appendices.

For the suprema over $f \in \mathcal{F}$ of the second and third terms on the right-hand side of (11), we will need to control the covering numbers of the classes $\mathcal{G}_n$ and $\mathcal{G}_n^{(d)}$ (“d” for “diagonal”) given by

$$\mathcal{G}_n = \{g_{n,f} : f \in \mathcal{F}\}, \quad \mathcal{G}_n^{(d)} = \{q_n \varepsilon_{n,f}^2 : f \in \mathcal{F}\}. \quad (12)$$

We will find bounds on their covering numbers in terms of those of the ones of the collection of response functions $\mathcal{F}$. The bounds are of potentially independent interest. The proof of Proposition 1 is given in Section B in the appendices.

Proposition 1 (Preservation VC-class). *Let $\mathcal{F}$ be a VC-class with parameters (v, A) with respect to the envelope function F . Assume that the associated residuals $\varepsilon_{n,f}$ are uniformly bounded, i.e., there exists $M_n > 0$ such that $\sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty \leq M_n$. Then $\mathcal{G}_n$ and $\mathcal{G}_n^{(d)}$ defined in (12) are VC-classes with respect to the envelope $M_n^2 \|q_n\|_\infty$ with parameters $(4v, 4A_n)$ and $(2v, 2A_n)$ respectively, where*

$$A_n := 8A \|F\|_\infty \|q_n\|_\infty / M_n. \quad (13)$$

As $\mathcal{G}_n^{(d)}$ is a VC-class of functions by Proposition 1, we will be able to bound the supremum over $f \in \mathcal{F}$ of the second term on the right-hand side in (11) by an application of the following result.

Proposition 2. *On the probability space $(\mathcal{X}, \mathcal{A}, P)$, let $\mathcal{H}_n$ be a sequence of VC-classes of parameters (w_n, B_n) with respect to envelopes $U_n \geq \sup_{h \in \mathcal{H}_n} \|h\|_\infty$. Suppose the following conditions hold:*

- (i) $\tau_n^2 \geq \sup_{h \in \mathcal{H}_n} \text{var}_P(h)$ and $2\tau_n \leq U_n$;
- (ii) $w_n \geq 1$ and $B_n \geq 3\sqrt{e}$;
- (iii) $w_n U_n^2 \log(B_n U_n / \tau_n) = O(n\tau_n^2)$, $n \rightarrow \infty$.

Then, for P_n the empirical distribution of an independent random sample $X_1, \dots, X_n$ from P , we have

$$\sup_{h \in \mathcal{H}_n} |P_n\{h - P(h)\}| = O_{\mathbb{P}} \left(\frac{\tau_n}{\sqrt{n}} \sqrt{w_n \log \left(\frac{B_n U_n}{\tau_n} \right)} \right), \quad n \rightarrow \infty.$$

The proof of Proposition 2 is given in Section C in the appendices. In the proof, we bound the expectation of the supremum by combining a well-known symmetrization inequality (Devroye et al., 2013, Theorem 12.4) with Proposition 2.1 in Giné and Guillou (2001), and we will find a rate on the deviation of the supremum around its expectation by Corollary 3.4 in Talagrand (1994).

Following the above plan, the proof of Theorem 1 is given in detail in Section D in the appendices.

A Auxiliary lemmas

Proof of Lemma 1. The lemma states an asymptotic lower bound for the smallest eigenvalue of the empirical Gram matrix $P_n(\tilde{h}_n \tilde{h}_n^\top)$ under two alternative conditions, (a) or (b). The two cases are treated in the next two paragraphs.

Under assumption (a). Lemma 3 in Portier and Segers (2019) states that $P_n(h_n h_n^\top)$ and thus $P_n(\tilde{h}_n \tilde{h}_n^\top)$ fails to be invertible with probability at most $n^{-1}P(q_n^2)$. This probability tends to zero by assumption.

Recall the spectral norm $|\cdot|_2$ and let $|A|_F = (\sum_{i,j} A_{ij}^2)^{1/2}$ denote the Frobenius norm of a matrix A . Lemma 2 in Portier and Segers (2019) states that $\mathbb{E}[|P_n(\tilde{h}_n \tilde{h}_n^\top) - I_{d_n}|_F^2]$ is bounded by $n^{-1}P(q_n^2)$ and thus converges to zero as $n \rightarrow \infty$. But then the same is true for $\mathbb{E}[|P_n(\tilde{h}_n \tilde{h}_n^\top) - I_{d_n}|_2^2]$, since $|A|_2 \leq |A|_F$ for any square matrix A . It follows that $|P_n(\tilde{h}_n \tilde{h}_n^\top) - I_{d_n}|_2 = o_{\mathbb{P}}(1)$ as $n \rightarrow \infty$.

On the event that $P_n(\tilde{h}_n \tilde{h}_n^\top)$ is invertible, we have

$$\begin{aligned} |P_n(\tilde{h}_n \tilde{h}_n^\top)^{-1}|_2 &= |I_{d_n} + P_n(\tilde{h}_n \tilde{h}_n^\top)^{-1}\{I_{d_n} - P_n(\tilde{h}_n \tilde{h}_n^\top)\}|_2 \\ &\leq 1 + |P_n(\tilde{h}_n \tilde{h}_n^\top)^{-1}|_2 \cdot |P_n(\tilde{h}_n \tilde{h}_n^\top) - I_{d_n}|_2 \end{aligned}$$

from which

$$\frac{1}{\lambda_{\min}\{P_n(\tilde{h}_n \tilde{h}_n^\top)\}} = |P_n(\tilde{h}_n \tilde{h}_n^\top)^{-1}|_2 \leq \frac{1}{1 - |P_n(\tilde{h}_n \tilde{h}_n^\top) - I_{d_n}|_2} = 1 + o_{\mathbb{P}}(1), \quad n \rightarrow \infty.$$

Under assumption (b). Lemma A.2 in Leluc et al. (2019), which is based on Theorem 5.1.1 in Tropp (2015), states that for $0 < \delta < 1$ and for n sufficiently large such that $n > 2\|q_n\|_\infty \log(d_n/\delta)$, we have

$$\mathbb{P}\left[\lambda_{\min}\{P_n(\tilde{h}_n \tilde{h}_n^\top)\} \leq 1 - \sqrt{(2/n)\|q_n\|_\infty \log(d_n/\delta)}\right] \leq \delta.$$

By assumption, $\|q_n\|_\infty \log(d_n/\delta) = o(n)$ as $n \rightarrow \infty$, for any $0 < \delta < 1$. It follows that $\lambda_{\min}\{P_n(\tilde{h}_n \tilde{h}_n^\top)\} \geq 1 - o_{\mathbb{P}}(1)$ as $n \rightarrow \infty$. $\square$

Lemma 2. *If Conditions 1 and 2 hold, then*

$$\sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty \leq \|F\|_\infty + [\|q_n\|_\infty P(F^2)]^{1/2} \leq (1 + \|q_n\|_\infty^{1/2})\|F\|_\infty.$$

Proof of Lemma 2. Let $f \in \mathcal{F}$. We have $f = h_n^\top \beta_{n,f} + \varepsilon_{n,f}$ with $\beta_{n,f} = P(h_n h_n^\top)^{-1} P(h_n f)$ and $P(h_n \varepsilon_{n,f}) = 0$. Since $\tilde{h}_n = P(h_n h_n^\top)^{-1/2} h_n$, we get $f = \tilde{h}_n^\top P(\tilde{h}_n f) + \varepsilon_{n,f}$. Now $\tilde{h}_n$ and $\varepsilon_{n,f}$ are orthogonal while $P(\tilde{h}_n \tilde{h}_n^\top) = I_{d_n}$, so that

$$P(f^2) = |P(\tilde{h}_n f)|_2^2 + P(\varepsilon_{n,f}^2) \geq |P(\tilde{h}_n f)|_2^2.$$

It follows that

$$[\tilde{h}_n^\top P(\tilde{h}_n f)]^2 = q_n |P(\tilde{h}_n f)|_2^2 \leq q_n P(f^2) \leq q_n P(F^2).$$

But then

$$|\varepsilon_{n,f}| \leq |f| + |\tilde{h}_n^\top P(\tilde{h}_n f)| \leq |F| + [q_n P(F^2)]^{1/2}.$$

Since $P(F^2) \leq \|F\|_\infty^2$, the result follows. $\square$

B Proof of Proposition 1

The idea of the proof is to create a grid of functions on $\mathcal{X}$ based on a covering of $\mathcal{F}$ to cover $\mathcal{E}_n = \{\varepsilon_{n,f} : f \in \mathcal{F}\}$. From this grid, we will deduce coverings of $\mathcal{G}_n$ and $\mathcal{G}_n^{(d)}$.

Step 1: covering of $\mathcal{E}_n$. By assumption, the class $\mathcal{F}$ is VC of parameters (v, A) with respect to an envelope F . That means that for any $0 < \eta < 1$ and for any probability measure Q on $\mathcal{X}$, we have

$$\mathcal{N}(\mathcal{F}, L^2(Q), \eta \|F\|_{L^2(Q)}) \leq \left(\frac{A}{\eta}\right)^v.$$

Moreover, when $\eta \geq 1$, the ball centered at the constant function equal to zero and with radius $\|F\|_{L^2(Q)}$ is enough to cover $\mathcal{F}$. Thus, for any positive choice of η , the covering number is bounded from above by

$$\forall \eta \in (0, \infty), \quad \mathcal{N}(\mathcal{F}, L^2(Q), \eta) \leq \left(\frac{A\|F\|_{L^2(Q)}}{\eta}\right)^v \vee 1.$$

Fix $\eta > 0$, write $\eta_P = \eta/(4\|q_n\|_\infty)$ and $\eta_Q = \eta/4$, and define the covering number

$$N_P = \mathcal{N}(\mathcal{F}, L^2(P), \eta_P/2), \quad N_Q = \mathcal{N}(\mathcal{F}, L^2(Q), \eta_Q/2) \quad (14)$$

associated to the open balls

$$B_P(f, \delta) = \{g \in L^2(P) : \|g - f\|_{L^2(P)} < \delta\}, \quad B_Q(f, \delta) = \{g \in L^2(Q) : \|g - f\|_{L^2(Q)} < \delta\},$$

for $\delta > 0$. The balls in the definition of the covering numbers in (14) have their centers in $L^2(P)$ and $L^2(Q)$ but not necessarily in $\mathcal{F}$. At the price of doubling the radii, the triangle inequality permits us to find functions $f_1^{(P)}, \dots, f_{N_P}^{(P)}$ and $f_1^{(Q)}, \dots, f_{N_Q}^{(Q)}$ in $\mathcal{F}$ such that

$$\mathcal{F} \subset \bigcup_{i=1}^{N_P} B_P(f_i^{(P)}, \eta_P), \quad \mathcal{F} \subset \bigcup_{j=1}^{N_Q} B_Q(f_j^{(Q)}, \eta_Q). \quad (15)$$

Therefore, the class $\mathcal{F}$ is covered by the union of the intersections between the balls, that is to say

$$\mathcal{F} \subset \bigcup_{\substack{1 \leq i \leq N_P \\ 1 \leq j \leq N_Q}} \left(B_P(f_i^{(P)}, \eta_P) \cap B_Q(f_j^{(Q)}, \eta_Q) \right). \quad (16)$$

Define the support of this covering as

$$\mathcal{S} = \left\{ (i, j) \in \{1, \dots, N_P\} \times \{1, \dots, N_Q\} : B_P(f_i^{(P)}, \eta_P) \cap B_Q(f_j^{(Q)}, \eta_Q) \neq \emptyset \right\}.$$

For every $(i, j) \in \mathcal{S}$, we fix an arbitrary function $f_{i,j} \in B_P(f_i^{(P)}, \eta_P) \cap B_Q(f_j^{(Q)}, \eta_Q)$.

Let $f \in \mathcal{F}$ and let $(i, j) \in \mathcal{S}$ be such that $f \in B_P(f_i^{(P)}, \eta_P) \cap B_Q(f_j^{(Q)}, \eta_Q)$. We will show that $\|\varepsilon_{n,f} - \varepsilon_{n,f_{i,j}}\|_{L^2(Q)} \leq \eta$. Since f and $f_{i,j}$ belong to the same intersection in (16), we have

$$\|f - f_{i,j}\|_{L^2(P)} \leq 2\eta_P, \quad \|f - f_{i,j}\|_{L^2(Q)} \leq 2\eta_Q. \quad (17)$$

The residual functions can be expressed in terms of the whitened feature map $\mathbf{h}_n$ via

$$\varepsilon_{n,f} = f - \mathbf{h}_n^\top P(\mathbf{h}_n f), \quad \varepsilon_{n,f_{i,j}} = f_{i,j} - \mathbf{h}_n^\top P(\mathbf{h}_n f_{i,j}).$$

By the triangle inequality, we find

$$\|\varepsilon_{n,f} - \varepsilon_{n,f_{i,j}}\|_{L^2(Q)} \leq \|f - f_{i,j}\|_{L^2(Q)} + \|\mathbf{h}_n^\top P[\mathbf{h}_n(f - f_{i,j})]\|_{L^2(Q)}. \quad (18)$$

Recall $M_n \geq \sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty$, a constant envelope for the class $\mathcal{E}_n$, and recall the leverage function $q_n = |\mathbf{h}_n|_2^2$. By the Cauchy-Schwarz inequality,

$$\begin{aligned} \|\mathbf{h}_n^\top P[\mathbf{h}_n(f - f_{i,j})]\|_{L^2(Q)}^2 &= \int_{y \in \mathcal{X}} \left\{ \mathbf{h}_n(y)^\top \int_{x \in \mathcal{X}} \mathbf{h}_n(x)(f - f_{i,j})(x) dP(x) \right\}^2 dQ(y) \\ &\leq \int_{y \in \mathcal{X}} |\mathbf{h}_n(y)|_2^2 \left| \int_{x \in \mathcal{X}} \mathbf{h}_n(x)(f - f_{i,j})(x) dP(x) \right|_2^2 dQ(y) \\ &\leq \int_{y \in \mathcal{X}} |\mathbf{h}_n(y)|_2^2 \int_{x \in \mathcal{X}} |\mathbf{h}_n(x)|_2^2 (f - f_{i,j})^2(x) dP(x) dQ(y) \\ &\leq \|q_n\|_\infty^2 \|f - f_{i,j}\|_{L^2(P)}^2. \end{aligned} \quad (19)$$

The combination of (17), (18) and (19) yields

$$\|\varepsilon_{n,f} - \varepsilon_{n,f_{i,j}}\|_{L^2(Q)} \leq 2\eta_Q + 2\|q_n\|_\infty \eta_P = \eta/2 + \eta/2 = \eta.$$

We have thus constructed the covering

$$\mathcal{E}_n \subset \bigcup_{(i,j) \in \mathcal{S}} B_Q(\varepsilon_{n,f_{i,j}}, \eta)$$

of $\mathcal{E}_n$ with $L^2(Q)$ balls of radius at most η . The covering number of $\mathcal{E}_n$ is thus bounded by

$$\mathcal{N}(\mathcal{E}_n, L^2(Q), \eta) \leq \#\mathcal{S} \leq \mathcal{N}(\mathcal{F}, L^2(P), \eta_P/2) \cdot \mathcal{N}(\mathcal{F}, L^2(Q), \eta_Q/2).$$

Using the definition of a VC-class, we get that

$$\begin{aligned} & \mathcal{N}(\mathcal{E}_n, L^2(Q), \eta) \\ & \leq \left(\frac{2A\|F\|_{L^2(P)}}{\eta_P} \vee 1 \right)^v \cdot \left(\frac{2A\|F\|_{L^2(Q)}}{\eta_Q} \vee 1 \right)^v \\ & \leq \left(\frac{64A^2\|F\|_\infty^2\|q_n\|_\infty}{\eta^2} \right)^v \vee \left(\frac{8A\|F\|_\infty\|q_n\|_\infty}{\eta} \right)^v \vee 1 \\ & = \left(\frac{64A^2\|F\|_\infty^2\|q_n\|_\infty}{\eta^2} \right)^v \mathbf{1}_{\eta \leq t_-} + \left(\frac{8A\|F\|_\infty\|q_n\|_\infty}{\eta} \right)^v \mathbf{1}_{t_- < \eta < t_+} + \mathbf{1}_{t_+ \leq \eta} \\ & \leq \left(\frac{A_n M_n}{\eta} \right)^{2v} \vee 1, \end{aligned}$$

with $t_- = 8A\|F\|_\infty$, $t_+ = 8A\|F\|_\infty\|q_n\|_\infty$ and

$$A_n = 8A\|F\|_\infty\|q_n\|_\infty/M_n. \quad (20)$$

We deduce that the residual class $\mathcal{E}_n$ is VC of parameters $(2v, A_n)$ with respect to the envelope M_n .

Step 2: covering of $\mathcal{G}_n$. Consider two functions $f, \tilde{f} \in \mathcal{F}$ and a probability measure Q on $\mathcal{X}^2$ with marginals Q_1, Q_2 on $\mathcal{X}$, that is, $Q_1(B) = Q(B \times \mathcal{X})$ and $Q_2(B) = Q(\mathcal{X} \times B)$ for measurable $B \subset \mathcal{X}$. By definition, every function in $\mathcal{G}_n$ is written as $(x, y) \mapsto \tilde{h}_n(x)^\top \tilde{h}_n(y) \varepsilon_{n,f}(x) \varepsilon_{n,f}(y)$. The Cauchy–Schwarz inequality gives

$$\forall (x, y) \in \mathcal{X}^2, \quad |\tilde{h}_n(x)^\top \tilde{h}_n(y)|^2 \leq |\tilde{h}_n(x)|_2^2 |\tilde{h}_n(y)|_2^2 = q_n(x)q_n(y) \leq \|q_n\|_\infty^2. \quad (21)$$

Furthermore, since $(a+b)^2 \leq 2a^2 + 2b^2$ for $a, b \in \mathbb{R}$, we have

$$\begin{aligned} & \{\varepsilon_{n,f}(x) \varepsilon_{n,f}(y) - \varepsilon_{n,\tilde{f}}(x) \varepsilon_{n,\tilde{f}}(y)\}^2 \\ & \leq 2\varepsilon_{n,f}^2(x) \{\varepsilon_{n,f}(y) - \varepsilon_{n,\tilde{f}}(y)\}^2 + 2\{\varepsilon_{n,f}(x) - \varepsilon_{n,\tilde{f}}(x)\}^2 \varepsilon_{n,\tilde{f}}^2(y) \\ & \leq 2M_n^2 [\{\varepsilon_{n,f}(y) - \varepsilon_{n,\tilde{f}}(y)\}^2 + \{\varepsilon_{n,f}(x) - \varepsilon_{n,\tilde{f}}(x)\}^2] \end{aligned}$$

for all $(x, y) \in \mathcal{X}^2$. We obtain

$$\begin{aligned} & \|g_{n,f} - g_{n,\tilde{f}}\|_{L^2(Q)}^2 \\ & = \int_{(x,y) \in \mathcal{X}^2} |\tilde{h}_n(x)^\top \tilde{h}_n(y)|^2 \{\varepsilon_{n,f}(x) \varepsilon_{n,f}(y) - \varepsilon_{n,\tilde{f}}(x) \varepsilon_{n,\tilde{f}}(y)\}^2 dQ(x, y) \\ & \leq 2\|q_n\|_\infty^2 M_n^2 \left(\|\varepsilon_{n,f} - \varepsilon_{n,\tilde{f}}\|_{L^2(Q_1)}^2 + \|\varepsilon_{n,f} - \varepsilon_{n,\tilde{f}}\|_{L^2(Q_2)}^2 \right). \end{aligned} \quad (22)$$

Fix $\eta > 0$. Following the approach in Step 1, we can for $\ell = 1, 2$ construct a covering of $\mathcal{E}_n$ by $L^2(Q_\ell)$ -balls of at most radius η . The centers of the balls are of the form

$$\forall \ell = 1, 2, \forall k = 1, \dots, m_\ell, \quad \varepsilon_{n,f_k^{(\ell)}} \in \mathcal{E}_n,$$

where $f_k^{(\ell)}$ belongs to $\mathcal{F}$ and where m_ℓ is the number of such balls needed, a number which is bounded by $(A_n M_n / \eta)^{2v} \vee 1$ for A_n defined in (20). Consider the intersections

$$\forall i = 1, \dots, m_1, \forall j = 1, \dots, m_2, \quad B(i, j) = B_{Q_1}(\varepsilon_{n, f_i^{(1)}}(\eta)) \cap B_{Q_2}(\varepsilon_{n, f_j^{(2)}}(\eta)).$$

The set $\mathcal{E}_n$ is covered by the union of all those intersections $B(i, j)$. For each $(i, j) \in \{1, \dots, m_1\} \times \{1, \dots, m_2\}$ such that $\mathcal{E}_n$ intersects $B(i, j)$, pick an arbitrary $f_{i,j} \in \mathcal{F}$ such that $\varepsilon_{n, f_{i,j}} \in \mathcal{E}_n \cap B(i, j)$. Note that the functions $f_{i,j}$ are different from the ones denoted in the same way in Step 1.

Let $f \in \mathcal{F}$ and let (i, j) be such that $\varepsilon_{n, f}$ and $\varepsilon_{n, f_{i,j}}$ belong to the same intersection $B(i, j)$. Since the diameters of the two balls in the definition of $B(i, j)$ are bounded by 2η in view of the triangle inequality, we find

$$\forall \ell = 1, 2, \quad \|\varepsilon_{n, f} - \varepsilon_{n, f_{i,j}}\|_{L^2(Q_\ell)} \leq 2\eta.$$

By (22), it follows that

$$\|g_{n, f} - g_{n, f_{i,j}}\|_{L^2(Q)}^2 \leq 2\|q_n\|_\infty^2 M_n^2 [(2\eta)^2 + (2\eta)^2] = 16\|q_n\|_\infty^2 M_n^2 \eta^2.$$

We find that $\mathcal{G}_n$ is covered by the union of the balls $B_Q(g_{n, f_{i,j}}, 4\|q_n\|_\infty M_n \eta)$. Its covering number is thus bounded by

$$\mathcal{N}(\mathcal{G}_n, L^2(Q), 4\|q_n\|_\infty M_n \eta) \leq m_1 m_2 \leq \left(\frac{A_n M_n}{\eta}\right)^{4v} \vee 1.$$

Rescaling $M_n \eta' = 4\eta$, we get

$$\mathcal{N}(\mathcal{G}_n, L^2(Q), \|q_n\|_\infty M_n^2 \eta') \leq \left(\frac{4A_n}{\eta'}\right)^{4v} \mathbb{1}_{\eta' < 4} + \mathbb{1}_{\eta' \geq 4} \leq \left(\frac{4A_n}{\eta'}\right)^{4v} \vee 1.$$

In view of (21), the functions in $\mathcal{G}_n$ are uniformly bounded by $\|q_n\|_\infty M_n^2$. Since Q was an arbitrary probability measure on $\mathcal{X}^2$, we conclude that $\mathcal{G}_n$ is a VC-class with parameters $(4v, 4A_n)$ with respect to the constant envelope $\|q_n\|_\infty M_n^2$.

Step 3: covering of $\mathcal{G}_n^{(d)}$. Let Q be a probability measure on $\mathcal{X}$ and let $\eta > 0$. In Step 1, we found functions $f_{i,j} \in \mathcal{F}$ such that $\mathcal{E}_n$ is covered by the balls $B_Q(\varepsilon_{n, f_{i,j}}(\eta))$, and we needed at most $(A_n M_n / \eta)^{2v}$ of such functions. For $f \in \mathcal{F}$ we can thus find (i, j) such that

$$\|\varepsilon_{n, f} - \varepsilon_{n, f_{i,j}}\|_{L^2(Q)} \leq \eta$$

and thus

$$\begin{aligned} \|q_n \varepsilon_{n, f}^2 - q_n \varepsilon_{n, f_{i,j}}^2\|_{L^2(Q)} &= \int_{x \in \mathcal{X}} \left\{ q_n(x) \varepsilon_{n, f}^2(x) - q_n(x) \varepsilon_{n, f_{i,j}}^2(x) \right\}^2 dQ(x) \\ &\leq 4\|q_n\|_\infty^2 M_n^2 \|\varepsilon_{n, f} - \varepsilon_{n, f_{i,j}}\|_{L^2(Q)}^2 \\ &\leq 4\|q_n\|_\infty^2 M_n^2 \eta^2. \end{aligned}$$

It follows that the number of $L^2(Q)$ balls of radius $\sqrt{4\|q_n\|_\infty^2 M_n^2 \eta^2}$ needed to cover $\mathcal{G}_n^{(d)}$ is bounded by the number of functions $f_{i,j}$ in the construction in Step 1, and so

$$\begin{aligned} \mathcal{N}(\mathcal{G}_n^{(d)}, L^2(Q), 2\|q_n\|_\infty M_n \eta) &\leq \mathcal{N}(\mathcal{E}_n, L^2(Q), \eta)^2 \\ &\leq \left(\frac{A_n M_n}{\eta}\right)^{2v} \mathbb{1}_{\eta < M_n} + \mathbb{1}_{\eta \geq M_n}. \end{aligned}$$

Upon rescaling $M_n \eta' = 2\eta$, we find

$$\begin{aligned} \mathcal{N}(\mathcal{G}_n^{(d)}, L^2(Q), \|q_n\|_\infty M_n^2 \eta') &\leq \left(\frac{2A_n}{\eta'}\right)^{2v} \mathbb{1}_{\eta' < 2} + \mathbb{1}_{\eta' \geq 2} \\ &\leq \left(\frac{2A_n}{\eta'}\right)^{2v} \vee 1. \end{aligned}$$

We conclude that $\mathcal{G}_n^{(d)}$ is a VC-class with parameters $(2v, 2A_n)$ with respect to the constant envelope $\|q_n\|_\infty M_n^2$. The proof of Proposition 1 is complete.

C Proof of Proposition 2

In this proof, we consider $Z_n := \sup_{h \in \mathcal{H}_n} |P_n(h) - P(h)|$. Thanks to the triangle inequality, we get

$$Z_n \leq \mathbb{E}(Z_n) + |Z_n - \mathbb{E}(Z_n)|. \quad (23)$$

We treat the two terms on the right-hand side of (23) in Steps 1 and 2, respectively. The convergence rate for Z_n then follows by adding both bounds in Step 3.

Step 1. Let $(\eta_i)_i$ denote a sequence of independent Rademacher variables, that is, $\mathbb{P}(\eta_i = +1) = \mathbb{P}(\eta_i = -1) = 1/2$ for all i , and such that $(\eta_i)_i$ and $(X_i)_i$ are independent. The symmetrization inequality (Devroye et al., 2013, Theorem 12.4) gives

$$n \mathbb{E}(Z_n) \leq 2 \mathbb{E} \left[\sup_{h \in \mathcal{H}_n} \left| \sum_{i=1}^n \eta_i \{h(X_i) - P(h)\} \right| \right].$$

Further, Proposition 2.1 in Giné and Guillou (2001) shows the existence of a universal constant $C > 0$ such that

$$\mathbb{E} \left[\sup_{\mathcal{H}_n} \left| \sum_{i=1}^n \eta_i \{h(X_i) - P(h)\} \right| \right] \leq C \left(2w_n U_n \log(2\theta_n) + \tau_n \sqrt{w_n n \log(2\theta_n)} \right),$$

where we wrote

$$\theta_n = B_n U_n / \tau_n.$$

By the asymptotic relation in (iii), the latter inequality is bounded from above by $\tilde{C} r_n$ where

$$\begin{aligned} r_n &= \tau_n \sqrt{w_n n \log(\theta_n)}, \\ \tilde{C} &= 2C \left[1 + \sup_n \sqrt{w_n U_n^2 \log(2\theta_n) / (n \tau_n^2)} \right]. \end{aligned} \quad (24)$$

We find that $n \mathbb{E}(Z_n) \leq \tilde{C} r_n$.

Step 2. Corollary 3.4 in Talagrand (1994) states the existence of universal constants $K, L > 0$ such that

$$\forall t > 0, \quad \mathbb{P}(n|Z_n - \mathbb{E}(Z_n)| \geq r_n t) \leq K \exp \left(-\frac{r_n t}{2K U_n} \log \left(1 + 2t \frac{r_n U_n}{R_n} \right) \right) \quad (25)$$

where r_n is defined in (24) and where

$$R_n = \left(\sqrt{n} \tau_n + 2L U_n \sqrt{w_n \log(2\theta_n)} \right)^2.$$

Since $\log(1+x) \geq x/(1+x/2)$ for all $x \geq 0$, we get

$$\frac{r_n t}{2K U_n} \log \left(1 + 2t \frac{r_n U_n}{R_n} \right) \geq \frac{r_n^2 t^2}{K R_n \left(1 + \frac{r_n U_n t}{R_n} \right)}. \quad (26)$$

Eventually, we deduce from (25) and (26) that

$$\forall t > 0, \quad \mathbb{P}(n|Z_n - \mathbb{E}(Z_n)| \geq r_n t) \leq K \exp \left(-\frac{r_n^2 t^2}{K R_n \left(1 + \frac{r_n U_n t}{R_n} \right)} \right).$$

For any $\delta \in (0, 1]$, solving in t the equation $K \exp(-r_n^2 t^2 / \{K R_n (1 + r_n U_n t / R_n)\}) = \delta$ gives

$$t_n^{(\delta)} = K(U_n / 2r_n) \log(K/\delta) + \sqrt{K^2 (U_n / 2r_n)^2 \log(K/\delta)^2 + K(R_n / r_n^2) \log(K/\delta)}. \quad (27)$$

The first term U_n/r_n in (27) can be rewritten as

$$\begin{aligned}\frac{U_n}{r_n} &= \frac{U_n}{\tau_n \sqrt{w_n n \log(\theta_n)}} \\ &= \frac{U_n \sqrt{w_n \log(\theta_n)}}{\sqrt{n} \tau_n} \{w_n \log(\theta_n)\}^{-1}.\end{aligned}$$

while the second argument R_n/r_n^2 in (27) is bounded from above by

$$\begin{aligned}\frac{R_n}{r_n^2} &\leq \frac{4 \max \{n\tau_n^2, 4L^2 U_n^2 w_n \log(2\theta_n)\}}{n\tau_n^2 w_n \log(\theta_n)} \\ &\leq 4 \max \left\{ [w_n \log(\theta_n)]^{-1}, \frac{4L^2 U_n^2 \log(2\theta_n)}{n\tau_n^2 \log(\theta_n)} \right\}.\end{aligned}$$

By the assumption (iii), we obtain $\sup_n U_n/r_n < \infty$ and $\sup_n R_n/r_n^2 < \infty$. For any $\delta \in (0, 1]$, it follows that $\sup_n t_n^{(\delta)} < \infty$.

Step 3. Combining the bound (23) with the rates found in Steps 1 and 2 for the two terms in the bound, we obtain

$$\forall \delta \in (0, 1], \quad \mathbb{P} \left(nZ_n > (\sup_n t_n^{(\delta)} + \tilde{C})r_n \right) \leq \delta.$$

Divide by n and plug in the definition (24) of r_n to conclude the proof of Proposition 2.

D Proof of Theorem 1

We recall and introduce the following quantities:

$$\begin{aligned}\sigma_n &= \sup_{f \in \mathcal{F}} [P(\varepsilon_{n,f}^2)]^{1/2}, \quad M_n = \sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty, \\ \gamma_n &= \sup_{f \in \mathcal{F}} [P(q_n \varepsilon_{n,f}^2)]^{1/2}, \quad L_n = \|q_n\|_\infty^{1/2} M_n.\end{aligned}\tag{28}$$

For all $f \in \mathcal{F}$, we have $P(\varepsilon_{n,f}^2) \leq P(f^2) \leq P(F^2)$. Since $P(q_n) = d_n$, we find

$$\begin{aligned}\sigma_n &\leq M_n \wedge [P(F^2)]^{1/2} \leq M_n \wedge \|F\|_\infty, \\ \gamma_n &\leq d_n^{1/2} M_n \wedge \|q_n\|_\infty^{1/2} \sigma_n \leq L_n.\end{aligned}\tag{29}$$

Step 1. We follow the plan laid out in Section 4 in the paper. In view of (10) and Lemma 1, we have

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_{n,f}) - L_f(\beta_f)\} \leq \{1 + \mathbb{O}_{\mathbb{P}}(1)\} \sup_{f \in \mathcal{F}} |P_n(\hat{h}\varepsilon_{n,f})|_2^2.$$

The inequality (11) provides a bound on $n^2 |P_n(\hat{h}\varepsilon_{n,f})|_2^2$ consisting of a sum of three terms. The first term is just $nP(q_n \varepsilon_{n,f}^2) \leq n\gamma_n^2$. The second and third terms require a separate analysis (Steps 2 and 3). Finally, we collect the bounds to arrive at the stated rate (Step 4).

Step 2. We apply Proposition 2 to the class $\mathcal{G}_n^{(d)} = \{q_n \varepsilon_{n,f}^2 : f \in \mathcal{F}\}$ introduced in (12). We get

$$\sup_{g \in \mathcal{G}_n^{(d)}} \left| \sum_{i=1}^n \{g(X_i) - P(g)\} \right| = \mathcal{O} \left(\sqrt{n\tau^2 w \log(BU/\tau)} \right), \quad n \rightarrow \infty, \tag{30}$$

where the positive quantities τ, U, w, B depend on n and must satisfy the following conditions:

- (i) $\tau^2 \geq \sup_{f \in \mathcal{F}} \text{var}[(q_n \varepsilon_{n,f}^2)(X_1)]$;

- (ii) $U \geq \sup_{f \in \mathcal{F}} \|q_n \varepsilon_{n,f}^2\|_\infty$ and $U \geq 2\tau$;
- (iii) (w, B) are VC parameters of $\mathcal{G}_n^{(d)}$ with respect to the envelope U , and $w \geq 1$ and $B \geq 3\sqrt{e}$;
- (iv) $w(U/\tau)^2 \log(BU/\tau) = O(n)$.

Recall $M_n = \sup_{f \in \mathcal{F}} \|\varepsilon_{n,f}\|_\infty$ and $L_n = \|q_n\|_\infty^{1/2} M_n$ in (28). We have $\sup_{f \in \mathcal{F}} \|q_n \varepsilon_{n,f}^2\|_\infty \leq L_n^2$ and therefore we put

$$U = 2L_n^2. \quad (31)$$

Further, we have

$$\text{var}[(q_n \varepsilon_{n,f}^2)(X_1)] \leq \mathbb{E}[q_n^2(X_1) \varepsilon_{n,f}^4(X_1)] \leq \|q_n\|_\infty M_n^2 P(q_n \varepsilon_{n,f}^2) \leq L_n^2 \gamma_n^2.$$

To meet (i) and at the same time have an upper bound on U/τ , we set

$$\tau = L_n \gamma_n \vee \frac{L_n^2 \sqrt{\log n}}{\sqrt{n}}. \quad (32)$$

Since $\gamma_n \leq L_n$, conditions (i) and (ii) are met.

Next, Proposition 1 says that, for any probability measure Q on $\mathcal{X}$ and any $0 < \eta < 1$,

$$\mathcal{N}(\mathcal{G}_n^{(d)}, L^2(Q), \eta) \leq \left(\frac{2A_n M_n^2 \|q_n\|_\infty}{\eta} \right)^{2v} = \left(\frac{A_n U}{\eta} \right)^{2v}$$

with A_n defined in (13). We can therefore set $w = \max(2v, 1)$, while for B it is sufficient to ensure that $B \geq \max(A_n, 3\sqrt{e})$. But

$$A_n \leq C_{\mathcal{F}} \|q_n\|_\infty / M_n \quad (33)$$

for some constant $C_{\mathcal{F}}$ that depends only on $\mathcal{F}$. Therefore, condition (iii) is met for

$$B = (C_{\mathcal{F}} \|q_n\|_\infty / M_n) \vee (3\sqrt{e}).$$

As $M_n \geq \sigma_n$, the conditions in the main theorem imply that $\log(B) = O(\log n)$ as $n \rightarrow \infty$. By definition of U and τ in (31) and (32) respectively, we have $2 \leq U/\tau \leq 2(n/\log n)^{1/2}$. It follows that $\log(U/\tau) = O(\log n)$ as $n \rightarrow \infty$ and thus that (iv) is met since

$$(U/\tau)^2 \log(BU/\tau) \leq 4(n/\log n) O(\log n) = O(n), \quad n \rightarrow \infty.$$

We can now develop the rate obtained in (30). As $\log(BU/\tau) = O(\log n)$ as $n \rightarrow \infty$, we get

$$\sup_{g \in \mathcal{G}_n^{(d)}} \left| \sum_{i=1}^n \{g(X_i) - P(g)\} \right| = O \left(\left(L_n \gamma_n \vee \frac{L_n^2 \sqrt{\log n}}{\sqrt{n}} \right) \sqrt{n \log n} \right), \quad n \rightarrow \infty. \quad (34)$$

In the computation of the rate of the excess prediction risk, we must divide the rate above by n^2 .

Step 3. The term $\sum_{1 \leq i \neq j \leq n} g_{n,f}(X_i, X_j)$ is a degenerate U -statistic of order two. Note indeed that $\mathbb{E}[g_{n,f}(X, x)] = \mathbb{E}[g_{n,f}(x, X)] = 0$ for any $x \in \mathcal{X}$, where the random variable X has distribution P . We apply a special case of Theorem 2 in Major (2006), cited for convenience as Theorem 2 below. The functions $g_{n,f}$ are uniformly bounded by

$$\sup_{x, y \in \mathcal{X}} |g_{n,f}(x, y)| \leq \|q_n\|_\infty M_n^2 = L_n^2.$$

Consider the class of rescaled functions

$$\tilde{\mathcal{G}}_n = \{g_{n,f}/L_n^2 : f \in \mathcal{F}\}.$$

In view of Proposition 1, its covering numbers are bounded by

$$\mathcal{N}(\tilde{\mathcal{G}}_n, L^2(Q), \eta) \leq (4A_n)^{4v} \eta^{-4v}$$

for any probability measure Q on $\mathcal{X}^2$ and for any $\eta \in (0, 1]$. In the terminology of Major (2006, p. 490), this means that $\tilde{\mathcal{G}}_n$ is an L^2 -dense class of functions with parameter D and exponent w defined by

$$D = (4A_n)^{4v}, \quad w = 4v \vee 1.$$

For u , we take the maximum with 1 in order to apply Theorem 2 later on.

By the Cauchy–Schwarz inequality, we have $g_{n,f}(x, y)^2 \leq q_n(x)\varepsilon_{n,f}^2(x)q_n(y)\varepsilon_{n,f}^2(y)$ for all $(x, y) \in \mathcal{X}^2$ and $f \in \mathcal{F}$. It follows that, for independent random variables X_1 and X_2 with common distribution P , we have

$$\forall f \in \mathcal{F}, \quad (\mathbb{E}[g_{n,f}^2(X_1, X_2)])^{1/2} \leq P(q_n\varepsilon_{n,f}^2) \leq \gamma_n^2.$$

Put

$$v = \frac{\gamma_n^2}{L_n^2} \vee \frac{\log n}{\sqrt{n}}.$$

Then $v^2 \in (0, 1]$ is an upper bound on the second moment of the functions in $\tilde{\mathcal{G}}_n$.

Theorem 2 yields that

$$\mathbb{P} \left(\sup_{f \in \mathcal{F}} \left| \sum_{1 \leq i \neq j \leq n} g_{n,f}(X_i, X_j) \right| \geq 2nvL_n^2 y \right) \leq CD \exp(-\alpha y), \quad (35)$$

for all $y \in [y_{n,-}, y_{n,+}]$, where

$$\begin{aligned} y_{n,-} &= K \left(w + \left(\frac{\log D}{\log n} \right)_+ \right)^{3/2} \log \left(\frac{2}{v} \right), \\ y_{n,+} &= nv^2, \end{aligned} \quad (36)$$

and where α , C and K are universal constants.

By definition, we have $1/v \leq \sqrt{n}/\log n$ and thus $\log(1/v) = O(\log n)$ as $n \rightarrow \infty$. Furthermore,

$$\log(\|q_n\|_\infty) + |\log M_n| = O(\log n), \quad n \rightarrow \infty, \quad (37)$$

by assumption and by Lemma 2, whence $\log(D) = O(\log n)$ as $n \rightarrow \infty$ in view of (33). It follows that

$$y_{n,-} = O(\log n) = o(y_{n,+}), \quad n \rightarrow \infty,$$

and thus $y_{n,-} < y_{n,+}$ for all sufficiently large n .

Let $0 < \delta < 1$ and define

$$y(\delta) = \frac{1}{\alpha} \log \left(\frac{CD}{\delta} \right) \vee y_{n,-}.$$

Then $y(\delta) \in [y_{n,-}, y_{n,+}]$ for all sufficiently large n and moreover

$$CD \exp \{-\alpha y(\delta)\} \leq \delta.$$

Since moreover $y(\delta) = O(\log n)$ as $n \rightarrow \infty$, we get from (35) and $vL_n^2 = \gamma_n^2 \vee L_n^2 n^{-1/2} \log n$ that

$$\sup_{f \in \mathcal{F}} \left| \sum_{1 \leq i \neq j \leq n} g_{n,f}(X_i, X_j) \right| = O_{\mathbb{P}} \left(\left(\gamma_n^2 \vee \frac{L_n^2 \log n}{\sqrt{n}} \right) n \log n \right), \quad n \rightarrow \infty. \quad (38)$$

Step 4. Assemble the bounds for the suprema over $f \in \mathcal{F}$ of the three terms on the right-hand side of (11). From (34) and 38, dividing by n^2 , we get

$$\begin{aligned} \sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_{n,f}) - L_f(\beta_f)\} &= O\left(\frac{\gamma_n^2}{n}\right) + O_{\mathbb{P}}\left(\left(L_n \gamma_n \vee \frac{L_n^2 \sqrt{\log n}}{\sqrt{n}}\right) \frac{\sqrt{\log n}}{n\sqrt{n}}\right) \\ &\quad + O_{\mathbb{P}}\left(\left(\gamma_n^2 \vee \frac{L_n^2 \log n}{\sqrt{n}}\right) \frac{\log n}{n}\right), \quad n \rightarrow \infty. \end{aligned}$$

As $\gamma_n \leq L_n$, the first and second terms are dominated by the third one. The convergence rate of the excess prediction risk is thus

$$\sup_{f \in \mathcal{F}} \{L_f(\hat{\beta}_{n,f}) - L_f(\beta_f)\} = O_{\mathbb{P}}\left(\left(\gamma_n^2 \vee \frac{L_n^2 \log n}{\sqrt{n}}\right) \frac{\log n}{n}\right), \quad n \rightarrow \infty. \quad (39)$$

This completes the proof of Theorem 1.

E A concentration inequality for degenerate U-statistics

The main term in the proof of Theorem 1 concerned a degenerate U-statistic of order two, see Step 3 in Section D. We dealt with it via a special case of the concentration inequality in Theorem 2 in Major (2006), stated next.

Theorem 2 (Special case of Theorem 2 in Major (2006)). *Let $(\mathcal{X}, \mathcal{A}, P)$ be a probability space and let $\mathcal{G}$ be an at most countably infinite collection of measurable functions $g : \mathcal{X}^2 \rightarrow [-1, 1]$ such that $\int_{\mathcal{X}} g(x, z) dP(z) = \int_{\mathcal{X}} g(z, x) dP(z) = 0$ for every $x \in \mathcal{X}$. Assume that $\mathcal{G}$ is a countable VC-class of parameters (w, B) , with $w \geq 1$ and $B > 0$. Let $v \in (0, 1]$ be such that $\sup_{g \in \mathcal{G}} \mathbb{E}[g^2(X_1, X_2)] \leq v^2$. Let $X_1, \dots, X_n$ be an independent random sample from P . There exist universal positive constants α, C and K such that*

$$\forall y \in [y_-, y_+], \quad \mathbb{P}\left(\sup_{g \in \mathcal{G}} \left|\sum_{1 \leq i \neq j \leq n} g(X_i, X_j)\right| \geq 2nv\right) \leq CB^w e^{-\alpha y}$$

where

$$y_- = K[w + (w \log B / \log n)_+]^{3/2} \log(2/v), \quad y_+ = nv^2.$$

Acknowledgments and Disclosure of Funding

The authors thank very particularly Rémi Leluc for his valuable feedback and gratefully acknowledge funding from FNRS-F.R.S. grant CDR J.0146.19.

References

- Azzi, S., Y. Huang, B. Sudret, and J. Wiart (2019). Surrogate modeling of stochastic functions – application to computational electromagnetic dosimetry. *International Journal for Uncertainty Quantification* 9(4).
- Bauer, B., F. Heimrich, M. Kohler, and A. Krzyżak (2019). On estimation of surrogate models for multivariate computer experiments. *Annals of the Institute of Statistical Mathematics* 71(1), 107–136.
- Belloni, A., V. Chernozhukov, D. Chetverikov, and K. Kato (2015). Some new asymptotic theory for least squares series: Pointwise and uniform results. *Journal of Econometrics* 186(2), 345–366.
- Breitkopf, P., H. Naceur, A. Rassineux, and P. Villon (2005). Moving least squares response surface approximation: formulation and metal forming applications. *Computers & Structures* 83(17-18), 1411–1428.

- Bucher, C. G. and U. Bourgund (1990). A fast and efficient response surface approach for structural reliability problems. *Structural Safety* 7(1), 57–66.
- Devroye, L., L. Györfi, and G. Lugosi (2013). *A Probabilistic Theory of Pattern Recognition*, Volume 31. New York: Springer Science+Business Media.
- Forrester, A. I. and A. J. Keane (2009). Recent advances in surrogate-based optimization. *Progress in Aerospace Sciences* 45(1–3), 50–79.
- Frean, M. and P. Boyle (2008). Using Gaussian processes to optimize expensive functions. In *Australasian Joint Conference on Artificial Intelligence*, pp. 258–267. Springer.
- Friedman, J., T. Hastie, and R. Tibshirani (2001). *The Elements of Statistical Learning*. New York: Springer Science+Business Media.
- Giné, E. and A. Guillaou (2001). On consistency of kernel density estimators for randomly censored data: rates holding uniformly over adaptive intervals. *Annales de l’IHP Probabilités et Statistiques* 37(4), 503–522.
- Györfi, L., M. Kohler, A. Krzyzak, and H. Walk (2006). *A Distribution-Free Theory of Nonparametric Regression*. New York: Springer Science & Business Media.
- Härdle, W. (1990). *Applied Nonparametric Regression*. Cambridge: Cambridge University Press.
- Hsu, D., S. M. Kakade, and T. Zhang (2014). Random design analysis of ridge regression. *Foundations of Computational Mathematics* 14(3), 569–600.
- Jones, D. R. (2001). A taxonomy of global optimization methods based on response surfaces. *Journal of Global Optimization* 21(4), 345–383.
- Konakli, K. and B. Sudret (2016). Polynomial meta-models with canonical low-rank approximations: Numerical insights and comparison to sparse polynomial chaos expansions. *Journal of Computational Physics* 321, 1144–1169.
- Leluc, R., F. Portier, and J. Segers (2019). Control variate selection for monte carlo integration. *arXiv preprint arXiv:1906.10920*.
- Major, P. (2006). An estimate on the supremum of a nice class of stochastic integrals and U-statistics. *Probability Theory and Related Fields* 134(3), 489–537.
- Myers, R. H., D. C. Montgomery, and C. M. Anderson-Cook (2016). *Response Surface Methodology: Process and Product Optimization Using Designed Experiments* (Fourth ed.). Hoboken, NJ: John Wiley & Sons.
- Newey, W. K. (1997). Convergence rates and asymptotic normality for series estimators. *Journal of Econometrics* 79(1), 147–168.
- Nguyen, A.-T., S. Reiter, and P. Rigo (2014). A review on simulation-based optimization methods applied to building performance analysis. *Applied Energy* 113, 1043–1058.
- Portier, F. and J. Segers (2019). Monte Carlo integration with a growing number of control variates. *Journal of Applied Probability* 56(4), 1168–1186.
- Talagrand, M. (1994). Sharper bounds for Gaussian and empirical processes. *The Annals of Probability* 22(1), 28–76.
- Tropp, J. A. (2015). An introduction to matrix concentration inequalities. *Foundations and Trends® in Machine Learning* 8(1-2), 1–230. arXiv:1501.01571.
- van der Vaart, A. W. and J. A. Wellner (1996). *Weak Convergence and Empirical Process. With Applications to Statistics*. New York: Springer-Verlag.
- Zhou, S. and D. A. Wolfe (2000). On derivative estimation in spline regression. *Statistica Sinica* 10, 93–108.