



Sparse single-index model

Pierre Alquier, Gérard Biau

► To cite this version:

Pierre Alquier, Gérard Biau. Sparse single-index model. Journal of Machine Learning Research, 2013, 14, pp.243–280. hal-00556652v2

HAL Id: hal-00556652

<https://inria.hal.science/hal-00556652v2>

Submitted on 5 Oct 2011

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

SPARSE SINGLE-INDEX MODEL

Pierre Alquier¹

LPMA²

Université Paris Diderot – Paris VII
Boîte 188, 175 rue du Chevaleret
75013 Paris, France

CREST-LS

3 avenue Pierre Larousse
92240 Malakoff, France
`alquier@math.jussieu.fr`

Gérard Biau

LSTA & LPMA³

Université Pierre et Marie Curie – Paris VI
Boîte 158, Tour 15-25, 2ème étage
4 place Jussieu, 75252 Paris Cedex 05, France

DMA⁴

Ecole Normale Supérieure
45 rue d’Ulm
75230 Paris Cedex 05, France
`gerard.biau@upmc.fr`

Abstract

Let $(\mathbf{X}, Y)$ be a random pair taking values in $\mathbb{R}^p \times \mathbb{R}$. In the so-called single-index model, one has $Y = f^*(\theta^{*T} \mathbf{X}) + W$, where f^* is an unknown univariate measurable function, θ^* is an unknown vector in $\mathbb{R}^d$, and W denotes a random noise satisfying $\mathbb{E}[W|\mathbf{X}] = 0$. The single-index model is known to offer a flexible way to model a variety of high-dimensional real-world phenomena. However, despite its relative simplicity, this dimension reduction scheme is faced with severe complications as soon as the underlying dimension becomes

¹Corresponding author.

²Research partially supported by the French “Agence Nationale pour la Recherche” under grant ANR-09-BLAN-0128 “PARCIMONIE”.

³Research partially supported by the French “Agence Nationale pour la Recherche” under grant ANR-09-BLAN-0051-02 “CLARA”.

⁴Research carried out within the INRIA project “CLASSIC” hosted by Ecole Normale Supérieure and CNRS.

larger than the number of observations (“ p larger than n ” paradigm). To circumvent this difficulty, we consider the single-index model estimation problem from a sparsity perspective using a PAC-Bayesian approach. On the theoretical side, we offer a sharp oracle inequality, which is more powerful than the best known oracle inequalities for other common procedures of single-index recovery. The proposed method is implemented by means of the reversible jump Markov chain Monte Carlo technique and its performance is compared with that of standard procedures.

Index Terms — Single-index model, sparsity, regression estimation, PAC-Bayesian, oracle inequality, reversible jump Markov chain Monte Carlo method.

2010 Mathematics Subject Classification: 62G08, 62G05, 62G20.

1 Introduction

Let $\mathcal{D}_n = \{(\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)\}$ be a collection of independent observations, distributed as a generic independent pair $(\mathbf{X}, Y)$ taking values in $\mathbb{R}^p \times \mathbb{R}$ and satisfying $\mathbb{E}Y^2 < \infty$. Throughout, we let $\mathbf{P}$ be the distribution of $(\mathbf{X}, Y)$, so that the sample $\mathcal{D}_n$ is distributed according to $\mathbf{P}^{\otimes n}$. In the regression function estimation problem, the goal is to use the data $\mathcal{D}_n$ in order to construct an estimate $r_n : \mathbb{R}^p \rightarrow \mathbb{R}$ of the regression function $r(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$. In the classical parametric linear model, one assumes

$$Y = \theta^{\star T} \mathbf{X} + W,$$

where $\theta^{\star} = (\theta_1^{\star}, \dots, \theta_p^{\star})^T \in \mathbb{R}^p$ and $\mathbb{E}[W|\mathbf{X}] = 0$. Here

$$r(\mathbf{x}) = \theta^{\star T} \mathbf{x} = \sum_{j=1}^p \theta_j^{\star} x_j$$

is a linear function of the components of $\mathbf{x} = (x_1, \dots, x_p)^T$. More generally, we may define

$$Y = f^{\star}(\theta^{\star T} \mathbf{X}) + W, \tag{1.1}$$

where $f^{\star}$ is an unknown univariate measurable function. This is the celebrated single-index model, which is recognized as a particularly useful variation of the linear formulation and can easily be interpreted: The model changes only in the direction $\theta^{\star}$, and the way it changes in this direction is described by the function $f^{\star}$. This model has applications to a variety of fields, such as discrete choice analysis in econometrics and dose response

models in biometrics, where high-dimensional regression models are often employed. There are too many references to be included here, but the monographs of McCullagh and Nelder [33] and Horowitz [27] together with the references [25, 28, 21, 18, 29] will provide the reader with good introductions to the general subject area.

One of the main advantages of the single-index model is its supposed ability to deal with the problem of high dimension (Bellman [6]). It is known that estimating the regression function is especially difficult whenever the dimension p of $\mathbf{X}$ becomes large. As a matter of fact, the optimal mean square convergence rate $n^{-2k/(2k+p)}$ for the estimation of a k -times differentiable regression function converges to zero dramatically slowly if the dimension p is large compared to k . This leads to an unsatisfactory accuracy of estimation for moderate sample sizes, and one possibility to circumvent this problem is to impose additional assumptions on the regression function. Thus, in particular, if $r(\mathbf{x}) = f^*(\theta^{*T}\mathbf{x})$ holds for every $\mathbf{x} \in \mathbb{R}^p$, then the underlying structural dimension of the model is 1 (instead of p) and the estimation of r can hopefully be performed easier. In this regard, Gaïffas and Lecué show in [22] that the optimal rate of convergence over the single-index model class is $n^{-2k/(2k+1)}$ (instead of $n^{-2k/(2k+p)}$), thereby answering a conjecture of Stone [40].

Nevertheless, practical estimation of the link function f^* and the index θ^* still requires a degree of statistical smoothing. Perhaps the most common approach to reach this goal is to use a nonparametric smoother (for instance, a kernel or a local polynomial method) to construct an approximation $\hat{f}_n$ of f^* , then substitute $\hat{f}_n$ into an empirical version $R_n(\theta)$ of the mean square error $R(\theta) = \mathbb{E}[Y - f(\theta^T\mathbf{X})]^2$, and finally choose $\hat{\theta}_n$ to minimize $R_n(\theta)$ (see e.g. Härdle, Hall and Ichimura [25] and Delecroix, Hristache and Patilea [21] where the procedure is discussed in detail). The rationale behind this type of two-stage approach, which is asymptotic in spirit, is that it produces a $\sqrt{n}$ -consistent estimate of θ , thereby devolving the difficulty to the simpler problem of computing a good estimate for the one-dimensional function f^* . However, the relative simplicity of this strategy is accompanied by severe difficulties (overfitting) when the dimension p becomes larger than the number of observations n . Estimation in this setting (called “ p larger than n ” paradigm) is generally acknowledged as an important challenge in contemporary statistics, see e.g. the recent monograph of Bühlmann and van de Geer [9]. In fact, this drawback considerably reduces the ability of the single-index model to behave as an effective dimension reduction technique.

On the other hand, there is empirical evidence that many signals in high-dimensional spaces admit a sparse representation. As an example, wavelet

coefficients of images often exhibit exponential decay, and a relatively small subset of all wavelet coefficients allow for a good approximation of the original image. Such signals have few nonzero coefficients and can therefore be described as sparse in the signal domain (see for instance [8]). Similarly, recent advances in high-throughput technologies—such as array comparative genomic hybridization—indicate that, despite the huge dimensionality of problems, only a small number of genes may play a role in determining the outcome and be required to create good predictors ([43] for instance). Sparse estimation is playing an increasingly important role in the statistics and machine learning communities, and several methods have recently been developed in both fields, which rely upon the notion of sparsity (e.g. penalty methods like the Lasso and Dantzig selector, see [41, 11, 10, 7] and the references therein).

In the present document, we consider the single-index model (1.1) from a sparsity perspective, i.e., we assume that θ^* has only a few coordinates different from 0. In the dimension reduction scenario we have in mind, the ambient dimension p can be very large, much larger than the sample size n , but we believe that the representation is sparse, i.e., that very few coordinates of θ^* are nonzero. This assumption is helpful at least for two reasons: If p is large and the number of nonzero coordinates is small enough, then the model is easier to interpret and its efficient estimation becomes possible. Our setting is close in spirit of the approach of Cohen, Daubechies, DeVore, Kerkycharian and Picard [16], who study approximation from queries of functions of the form $f(\theta^T \mathbf{x})$, where θ is approximately sparse (in the sense that it belongs to a weak- ℓ_p space). However, these authors do not provide any statistical study of their model. Our *modus operandi* will rather rely on the so-called PAC-Bayesian approach, originally developed in the classification context by Shawe-Taylor and Williamson [39], McAllester [32] and Catoni [12, 13]. This strategy was further investigated for regression by Audibert [4] and Alquier [1] and, more recently, worked out in the sparsity framework by Dalalyan and Tsybakov [19, 20] and Alquier and Lounici [2]. The main message of [19, 20, 2] is that aggregation with a properly chosen prior is able to deal nicely with the sparsity issue. Contrary to procedures such as the Lasso, the Dantzig selector and other penalized least square methods, which achieve fast rates under rather restrictive assumptions on the Gram matrix associated to the predictors, PAC-Bayesian aggregation requires only minimal assumptions on the model. Besides, it is computationally feasible even for a large p and exhibits good statistical performance.

The paper is organized as follows. In Section 2, we first set out some notation and introduce the single-index estimation procedure. Then we state our main

result (Theorem 2.1), which offers a sparsity oracle inequality more powerful than the best known oracle inequalities for other common procedures of single-index recovery. Section 3 is devoted to the practical implementation of the estimate via a reversible jump Markov chain Monte Carlo (MCMC) algorithm, and to numerical experiments on both simulated and real-life data sets. In order to preserve clarity, proofs have been postponed to Section 4 and the description of the MCMC method in its full length is given in the Appendix Section 5.

Note finally that our techniques extend to the case of multiple-index models, of the form

$$Y = f^*(\theta_1^{*T} \mathbf{X}, \dots, \theta_m^{*T} \mathbf{X}) + W,$$

where the underlying structural dimension m is supposed to be larger than 1 but substantially smaller than p . However, to keep things simple, we let $m = 1$ and leave the reader the opportunity to adapt the results to the more general situation $m \geq 1$.

2 Sparse single-index estimation

We start this section with some notation and basic requirements.

2.1 Notation

Throughout the document, we suppose that the recorded data $\mathcal{D}_n$ is generated according to the single-index model (1.1). More precisely, for each $i = 1, \dots, n$,

$$Y_i = f^*(\theta^{*T} \mathbf{X}_i) + W_i,$$

where f^* is a univariate measurable function, θ^* is a p -variate vector, and $W_1, \dots, W_n$ are independent copies of W . We emphasize that it is implicitly assumed that the observations are drawn according to the true model under study. Therefore, this casts our problem in a nonparametric setting rather than in a classical PAC-Bayesian one.

Recall that, in model (1.1), $\mathbb{E}[W|\mathbf{X}] = 0$ and, consequently, that $\mathbb{E}W = 0$. However, the distribution of W (in particular, the variance) may depend on $\mathbf{X}$. We shall not precisely specify this dependence, and will rather require the following condition on the distribution of W .

Assumption N. There exist two positive constants σ and L such that, for all integers $k \geq 2$,

$$\mathbb{E} [|W|^k | \mathbf{X}] \leq \frac{k!}{2} \sigma^2 L^{k-2}.$$

Observe that Assumption **N** holds in particular if $W = \Phi(\mathbf{X})\varepsilon$, where ε is a standard Gaussian random variable independent of $\mathbf{X}$ and $\Phi(\mathbf{X})$ is almost surely bounded.

Let $\|\theta\|_1$ denote the ℓ_1 -norm of the vector $\theta = (\theta_1, \dots, \theta_p)^T$, i.e., $\|\theta\|_1 = \sum_{j=1}^p |\theta_j|$. Without loss of generality, it will be assumed throughout the document that the index θ^* belongs to $\mathcal{S}_{1,+}^p$, where $\mathcal{S}_{1,+}^p$ is the set of all $\theta \in \mathbb{R}^p$ with $\|\theta\|_1 = 1$ and $\theta_{j(\theta)} > 0$, where $j(\theta)$ is the smallest $j \in \{1, \dots, p\}$ such that $\theta_j \neq 0$.

We will also require that the random variable $\mathbf{X}$ is almost surely bounded by a constant which, without loss of generality, can be taken equal to 1. Moreover, it will also be assumed that the link function f^* is bounded by some known positive constant C . Thus, denoting by $\|\mathbf{X}\|_\infty$ the supremum norm of $\mathbf{X}$ and by $\|f^*\|_\infty$ the functional supremum norm of f^* over $[-1, 1]$, we set:

Assumption B. The condition $\|\mathbf{X}\|_\infty \leq 1$ holds almost surely and there exists a positive constant C larger than 1 such that $\|f^*\|_\infty \leq C$.

Remark 2.1 *To keep a sufficient degree of clarity, no attempt was made to optimize the constants. In particular, the requirement $C \geq 1$ is purely technical. It can easily be removed by replacing C by $\max[C, 1]$ throughout the document.*

In order to approximate the link function f^* , we shall use the vector space $\mathcal{F}$ spanned by a given countable dictionary of measurable functions $\{\varphi_j\}_{j=1}^\infty$. Put differently, the approximation space $\mathcal{F}$ is the set of (finite) linear combinations of functions of the dictionary. Each φ_j of the collection is assumed to be defined on $[-1, 1]$ and to take values in $[-1, 1]$. To avoid getting into too much technicalities, we will also assume that each φ_j is differentiable and such that, for some positive constant ℓ , $\|\varphi_j'\|_\infty \leq \ell j$. This assumption is satisfied by the (non-normalized) trigonometric system

$$\varphi_1(t) = 1, \varphi_{2j}(t) = \cos(\pi j t), \varphi_{2j+1}(t) = \sin(\pi j t), \quad j = 1, 2, \dots$$

Finally, for any measurable $f : \mathbb{R}^p \rightarrow \mathbb{R}$ and $\theta \in \mathcal{S}_{1,+}^p$, we let

$$R(\theta, f) = \mathbb{E} \left[(Y - f(\theta^T \mathbf{X}))^2 \right]$$

and denote by

$$R_n(\theta, f) = \frac{1}{n} \sum_{i=1}^n (Y_i - f(\theta^T \mathbf{X}_i))^2$$

the empirical counterpart of $R(\theta, f)$ based on the sample $\mathcal{D}_n$.

2.2 Estimation procedure

We are now in a position to describe our estimation procedure. The method which is presented here is inspired by the approach developed by Catoni in [12, 13]. It strongly relies on the choice of a probability measure π on $\mathcal{S}_{1,+}^p \times \mathcal{F}$, called the prior, which in our framework should enforce the sparsity properties of the target regression function. With this objective in mind, we first let

$$d\pi(\theta, f) = d\mu(\theta)d\nu(f),$$

i.e., we assume that the distribution over the indexes is independent of the distribution over the link functions. With respect to the parameter θ , we put

$$d\mu(\theta) = \frac{\sum_{i=1}^p 10^{-i} \sum_{I \subset \{1, \dots, p\}, |I|=i} \binom{p}{i}^{-1} d\mu_I(\theta)}{1 - (\frac{1}{10})^p}, \quad (2.1)$$

where $|I|$ denotes the cardinality of I and $d\mu_I(\theta)$ is the uniform probability measure on the set

$$\mathcal{S}_{1,+}^p(I) = \{\theta = (\theta_1, \dots, \theta_p) \in \mathcal{S}_{1,+}^p : \theta_j = 0 \text{ if and only if } j \notin I\}.$$

We see that $\mathcal{S}_{1,+}^p(I)$ may be interpreted as the set of “active” coordinates in the single-index regression of Y on $\mathbf{X}$, and note that the prior on $\mathcal{S}_{1,+}^p$ is a convex combination of uniform probability measures on the subsets $\mathcal{S}_{1,+}^p(I)$. The weights of this combination depend only on the size of the active coordinate subset I . As such, the value $|I|$ characterizes the sparsity of the model: The smaller $|I|$, the smaller the number of variables involved in the model. The factor 10^{-i} penalizes models of high dimension, in accordance with the sparsity idea.

The choice of the prior ν on $\mathcal{F}$ is more involved. To begin with, we define, for any positive integer $M \leq n$ and all $\Lambda > 0$,

$$\mathcal{B}_M(\Lambda) = \left\{ (\beta_1, \dots, \beta_M) \in \mathbb{R}^M : \sum_{j=1}^M j|\beta_j| \leq \Lambda \text{ and } \beta_M \neq 0 \right\}.$$

Next, we let $\mathcal{F}_M(\Lambda) \subset \mathcal{F}$ be the image of $\mathcal{B}_M(\Lambda)$ by the map

$$\begin{aligned} \Phi_M : \mathbb{R}^M &\rightarrow \mathcal{F} \\ (\beta_1, \dots, \beta_M) &\mapsto \sum_{j=1}^M \beta_j \varphi_j. \end{aligned}$$

It is worth pointing out that, roughly, Sobolev spaces are well approximated by $\mathcal{F}_M(\Lambda)$ as M grows (more on this in Subsection 2.3). Finally, we define

$\nu_M(df)$ on the set $\mathcal{F}_M(C+1)$ as the image of the uniform measure on $\mathcal{B}_M(C+1)$ induced by the map Φ_M , and take

$$d\nu(f) = \frac{\sum_{M=1}^n 10^{-M} d\nu_M(f)}{1 - (\frac{1}{10})^n}. \quad (2.2)$$

Some comments are in order here. First, we note that the prior π is defined on $\mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$ endowed with its canonical Borel σ -field. The choice of $C+1$ instead of C in the definition of the prior support is essentially technical. This bound ensures that when the target f^* belongs to $\mathcal{F}_n(C)$, then a small ball around it is contained in $\mathcal{F}_n(C+1)$. It could be safely replaced by $C + u_n$, where $\{u_n\}_{n=1}^\infty$ is any positive sequence vanishing sufficiently slowly as $n \rightarrow \infty$. Next, the integer M should be interpreted as a measure of the “dimension” of the function f —the larger M , the more complex the function—and the prior ν adapts again to the sparsity idea by penalizing large-dimensional functions f . The coefficients 10^{-i} and 10^{-M} which appear in (2.1) and (2.2) show that more complex models have a geometrically decreasing influence. Note however that the value 10, which has been chosen because of its good practical results, is somehow arbitrary. It could be, in all generality, replaced by a more general coefficient α at the price of a more technical analysis. Finally, we observe that, for each $f = \sum_{j=1}^M \beta_j \varphi_j \in \mathcal{F}_M(C+1)$,

$$\|f\|_\infty \leq \sum_{j=1}^M |\beta_j| \leq C+1.$$

Now, let λ be a positive real number, called the inverse temperature parameter hereafter. The estimates $\hat{\theta}_\lambda$ and $\hat{f}_\lambda$ of θ^* and f^* , respectively, are simply obtained by randomly drawing

$$(\hat{\theta}_\lambda, \hat{f}_\lambda) \sim \hat{\rho}_\lambda,$$

where $\hat{\rho}_\lambda$ is the so-called Gibbs posterior distribution over $\mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$, defined by the probability density

$$\frac{d\hat{\rho}_\lambda}{d\pi}(\theta, f) = \frac{\exp[-\lambda R_n(\theta, f)]}{\int \exp[-\lambda R_n(\theta, f)] d\pi(\theta, f)}.$$

[The notation $d\hat{\rho}_\lambda/d\pi$ means the density of $\hat{\rho}_\lambda$ with respect to π .] The estimate $(\hat{\theta}_\lambda, \hat{f}_\lambda)$ has a simple interpretation. Firstly, the level of significance

of each pair (θ, f) is assessed via its least square error performance on the data $\mathcal{D}_n$. Secondly, a Gibbs distribution with respect to the prior π enforcing those pairs (θ, f) with the most empirical significance is assigned on the space $\mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$. Finally, the resulting estimate is just a random realization (conditional to the data) of this Gibbs posterior distribution.

2.3 Sparsity oracle inequality

For any $I \subset \{1, \dots, p\}$ and any positive integer $M \leq n$, we set

$$(\theta_{I,M}^*, f_{I,M}^*) \in \arg \min_{(\theta, f) \in \mathcal{S}_{1,+}^p(I) \times \mathcal{F}_M(C)} R(\theta, f).$$

At this stage, it is very important to note that, for each M , the infimum $f_{I,M}^*$ is defined on $\mathcal{F}_M(C)$, whereas the prior charges a slightly bigger set, namely $\mathcal{F}_M(C+1)$.

The main result of the paper is the following theorem. Here and everywhere, the wording “with probability $1 - \delta$ ” means the probability evaluated with respect to the distribution $\mathbf{P}^{\otimes n}$ of the data $\mathcal{D}_n$ and the conditional probability measure $\hat{\rho}_\lambda$. Recall that ℓ is a positive constant such that $\|\varphi'_j\|_\infty \leq \ell j$.

Theorem 2.1 *Assume that Assumption **N** and Assumption **B** hold. Set $w = 8(2C+1) \max[L, 2C+1]$ and take*

$$\lambda = \frac{n}{w + 2[(2C+1)^2 + 4\sigma^2]}. \quad (2.3)$$

Then, for all $\delta \in]0, 1[$, with probability at least $1 - \delta$ we have

$$R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \leq \Xi \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \left\{ R(\theta_{I,M}^*, f_{I,M}^*) - R(\theta^*, f^*) + \frac{M \log(Cn) + |I| \log(pn) + \log\left(\frac{2}{\delta}\right)}{n} \right\},$$

where Ξ is a positive constant, function of L , C , σ and ℓ only.

Remark 2.2 *Interestingly enough, analysis of the estimate $(\hat{\theta}_\lambda, \hat{f}_\lambda)$ is still possible when Assumption **N** is not satisfied. Indeed, even if Bernstein’s inequality (see Lemma 4.1) is not valid, a recent paper by Seldin, Cesa-Bianchi, Laviolette, Auer, Shawe-Taylor and Peters [38] provides us with a nice alternative inequality assuming less restrictive assumptions. However, we would*

then suffer a loss in the upper bound of Theorem 2.1. It is also interesting to note that recent results by Audibert and Catoni [5] allow the study of PAC-Bayesian estimates without Assumption **N**. However, the results of these authors are valid for linear models only, and it is therefore not clear to what extent their technique can be transposed to our setting.

Theorem 2.1 can be given a simple interpretation. Indeed, we see that if there is a “small” I and a “small” M such that $R(\theta_{I,M}^*, f_{I,M}^*)$ is close to $R(\theta^*, f^*)$, then $R(\hat{\theta}_\lambda, \hat{f}_\lambda)$ is also close to $R(\theta^*, f^*)$ up to terms of order $1/n$. However, if no such I or M exists, then one of the terms $M \log(Cn)/n$ and $|I| \log(pn)/n$ starts to dominate, thereby deteriorating the general quality of the bound. A good approximation with a “small” I is typically possible when θ^* is sparse or, at least, when it can be approximated by a sparse parameter. On the other hand, a good approximation with a “small” M is possible if f^* has a sufficient degree of regularity.

To illustrate the latter remark, assume for instance that $\{\varphi_j\}_{j=1}^\infty$ is the (non-normalized) trigonometric system and suppose that the target f^* belongs to the Sobolev ellipsoid, defined by

$$\mathcal{W}\left(k, \frac{6C^2}{\pi^2}\right) = \left\{ f \in L_2([-1, 1]) : f = \sum_{j=1}^{\infty} \beta_j \varphi_j \text{ and } \sum_{j=1}^{\infty} j^{2k} \beta_j^2 \leq \frac{6C^2}{\pi^2} \right\}$$

for some unknown regularity parameter $k \geq 2$ (see, e.g., Tsybakov [42]). Observe that, in this context, the approximation sets $\mathcal{F}_M(C+1)$ take the form

$$\begin{aligned} & \mathcal{F}_M(C+1) \\ &= \left\{ f \in L_2([-1, 1]) : f = \sum_{j=1}^M \beta_j \varphi_j, \sum_{j=1}^M j |\beta_j| \leq C+1 \text{ and } \beta_M \neq 0 \right\}. \end{aligned}$$

It is important to note that the regularity parameter k is assumed to be unknown, and this casts our results in the so-called adaptive setting. The following additional assumption will be needed:

Assumption D. The random variable $\theta^{*T} \mathbf{X}$ has a probability density on $[-1, 1]$, bounded from above by a positive constant B .

Last, we let I^* be the set I such that $\theta^* \in \mathcal{S}_{1,+}^p(I)$ and set $\|\theta^*\|_0 = |I^*|$.

Corollary 2.1 *Assume that Assumption **N**, Assumption **B** and Assumption **D** hold. Suppose also that f^* belongs to the Sobolev ellipsoid $\mathcal{W}(k, 6C^2/\pi^2)$,*

where the real number $k \geq 2$ is an (unknown) regularity parameter. Set $w = 8(2C + 1) \max[L, 2C + 1]$ and take λ as in (2.3). Then, for all $\delta \in]0, 1[$, with probability at least $1 - \delta$ we have

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \Xi' \left\{ \left(\frac{\log(Cn)}{n} \right)^{\frac{2k}{2k+1}} + \frac{\|\theta^*\|_0 \log(pn)}{n} + \frac{\log\left(\frac{2}{\delta}\right)}{n} \right\}, \end{aligned} \quad (2.4)$$

where Ξ' is a positive constant, function of L, C, σ, ℓ and B only.

As far as we are aware, all existing methods achieving rates of convergence similar to the ones provided by Corollary 2.1 are valid in an asymptotic setting only (p fixed and $n \rightarrow \infty$). The strength of Corollary 2.1 is to provide a finite sample bound and to show that our estimate still behaves well in a nonasymptotic situation if the intrinsic dimension (i.e., the sparsity) is small with respect to n . To understand this remark, just assume that p is a function of n such that $p \rightarrow \infty$ as $n \rightarrow \infty$. Whereas a classical asymptotic approach cannot say anything clever about this situation, our bounds still provide some information, provided the model is sparse enough (i.e., $\|\theta^*\|_0$ is sufficiently small with respect to n).

We see that, asymptotically (p fixed and $n \rightarrow \infty$), the leading term on the right-hand side of inequality (2.4) is $(\log(n)/n)^{\frac{2k}{2k+1}}$. This is the minimax rate of convergence over a Sobolev class, up to a $\log(n)$ factor. However, when n is “small” and θ^* is not sparse (i.e., $\|\theta^*\|_0$ is not “small”), the term $\|\theta^*\|_0 \log(pn)/n$ starts to emerge and cannot be neglected. Put differently, in large dimension, the estimation of θ^* itself is a problem—this phenomenon is not taken into account by asymptotic studies.

It is worth mentioning that the approach developed in the present article does not offer any guarantee on the point of view of variable (feature) selection. To reach this objective, an interesting route to follow is the sufficient dimension reduction (SDR) method proposed by Chen, Zou and Cook [15], which can be applied to the single-index model to estimate consistently the parameter θ^* and perform variable selection in a sparsity framework. Note however that such results require strong assumptions on the distribution of the data.

Finally, it should be stressed that the choice of λ in Theorem 2.1 and Corollary 2.1 is not the best possible and may eventually be improved, at the price of a more technical analysis however.

3 Implementation and numerical results

A series of experiments was conducted, both on simulated and real-life data sets, in order to assess the practical capabilities of the proposed method and compare its performance with that of standard procedures. Prior to analysis, we first need to discuss its concrete implementation, which has been carried out via a Markov Chain Monte Carlo (MCMC) method.

3.1 Implementation via reversible jump MCMC

The use of MCMC methods has become a popular way to compute Bayesian estimates. For an introduction to the domain, one should refer to the comprehensive monograph of Marin and Robert [30] and the references therein. Importantly, in this computational framework, an adaptation of the well-known Hastings-Metropolis algorithm to the case where the posterior distribution gives mass to several models of different dimensions was proposed by Green [23] under the name Reversible Jump MCMC (RJMCMC) method. In the PAC-Bayesian setting, MCMC procedures were first considered by Catoni [12], whereas Dalalyan and Tsybakov [19, 20] and Alquier and Lounici [2] explore their practical implementation in the sparse context using Langevin Monte Carlo and RJMCMC, respectively. Regarding the single-index model, MCMC algorithms were used to compute Bayesian estimates by Antoniadis, Grégoire and McKeague [3] and, more recently, by Wang [44], who develop a fully Bayesian method to analyse the single-index model. Our implementation technique is close in spirit to the one of Wang [44].

As a starting point for the approximate computation of our estimate, we used the RJMCMC method of Green [23], which is in fact an adaptation of the Hastings-Metropolis algorithm to the case where the objective posterior probability distribution (here, $\hat{\rho}_\lambda$) assigns mass to several different models. The idea is to start from an initial given pair $(\theta^{(0)}, f^{(0)}) \in \mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$ and then, at each step, to iteratively compute $(\theta^{(t+1)}, f^{(t+1)})$ from $(\theta^{(t)}, f^{(t)})$ via the following chain of rules:

- Sample a random pair $(\tau^{(t)}, h^{(t)})$ according to some proposal conditional density $k_t(\cdot | (\theta^{(t)}, f^{(t)}))$ with respect to the prior π ;
- Take

$$(\theta^{(t+1)}, f^{(t+1)}) = \begin{cases} (\tau^{(t)}, h^{(t)}) & \text{with probability } \alpha_t \\ (\theta^{(t)}, f^{(t)}) & \text{with probability } 1 - \alpha_t, \end{cases}$$

where

$$\alpha_t = \min \left(1, \frac{\frac{d\hat{\rho}_\lambda}{d\pi}(\tau^{(t)}, h^{(t)}) \times k_t((\theta^{(t)}, f^{(t)}) | (\tau^{(t)}, h^{(t)}))}{\frac{d\hat{\rho}_\lambda}{d\pi}(\theta^{(t)}, f^{(t)}) \times k_t((\tau^{(t)}, h^{(t)}) | (\theta^{(t)}, f^{(t)}))} \right).$$

This protocol ensures that the sequence $\{(\theta^{(t)}, f^{(t)})\}_{t=0}^\infty$ is a Markov chain with invariant probability distribution $\hat{\rho}_\lambda$ (see e.g. Marin and Robert [30]). A usual choice is to take $k_t \equiv k$, so that the Markov chain is homogeneous. However, in our context, it is more convenient to let $k_t = k_1$ if t is odd and $k_t = k_2$ if t is even. Roughly, the effect of k_1 is to modify the index $\theta^{(t)}$ while k_2 will essentially act on the link function $f^{(t)}$. While the ideas underlying the proposal densities k_1 and k_2 are quite simple, a precise description in its full length turns out to be more technical. Thus, in order to preserve the readability of the paper, the explicit construction of k_1 and k_2 has been postponed to the Appendix Section 5.

From a theoretical point of view, it is clear that the implementation of our method requires knowledge of the constant C (the upper bound on $\|f^*\|_\infty$). A too small C will result in a smaller model, which is unable to perform a good approximation. On the other hand, a larger C induces a poor bound in Theorem 2.1. In practice, however, the influence of C turns out to be secondary compared to the impact of the parameter λ . Indeed, it was found empirically that a very large choice of C (e.g., $C = 10^{100}$) does not deteriorate the overall quality of the results, as soon as λ is appropriately chosen. This is the approach that was followed in the experimental testing process.

Besides, the time for the Markov chains to converge depends strongly on the ambient dimension p and the starting point of the simulations. When the dimension is small (typically, $p \leq 10$), the chains converge fast and any value may be chosen as a starting point. In this case, we let the MCMC run 1000 steps and obtained satisfying results. On the other hand, when the dimension is larger (typically, $p > 10$), the convergence is very slow, in the sense that $R_n(\theta^{(t)}, f^{(t)})$ takes a very long time to stabilize. However, using as a starting point for the chains the preliminary estimate $\hat{\theta}_{\text{HHI}}$ (see below) significantly reduces the number of steps needed to reach convergence—we let the chains run 5000 steps in this context. Nevertheless, as a general rule, we encourage the users to inspect the convergence of the chains by checking if $R_n(\theta^{(t)}, f^{(t)})$ is stabilized, and to run several chains starting from different points to avoid their attraction into local minima.

3.2 Simulation study

In this subsection, we illustrate the finite sample performance of the presented estimation method on three synthetic data sets and compare its predictive capabilities with those of three standard statistical procedures. In all our experiments, we took as dictionary the (non-normalized) trigonometric system $\{\varphi_j\}_{j=1}^\infty$ and denote accordingly the resulting regression function estimate defined in Section 2 by $\hat{F}_{\text{Fourier}}$. In accordance with the order of magnitude indicated by the theoretical results, we set $\lambda = 4n$. This choice can undoubtedly be improved a bit but, as the numerical results show, it seems sufficient for our procedure to be fairly competitive.

The tested competing methods are the Lasso (Tibshirani [41]), the standard regression kernel estimate (Nadaraya [34, 35] and Watson [45], see also Tsybakov [42]), and the estimation strategy discussed in Härdle, Hall and Ichimura [25]. While the procedure of Härdle, Hall and Ichimura is specifically tailored for single-index models, the Lasso is designed to deal with the estimation of sparse linear models. On the other hand, the nonparametric kernel method is one of the best options when no obvious assumption (such as the single-index one) can be made on the shape of the targeted regression function.

We briefly recall that, for a linear model of the form $Y = \theta^*{}^T \mathbf{X} + W$, the Lasso estimate takes the form $\hat{F}_{\text{Lasso}}(\mathbf{x}) = \hat{\theta}_{\text{Lasso}}^T \mathbf{x}$, where

$$\hat{\theta}_{\text{Lasso}} \in \arg \min_{\theta \in \mathbb{R}^p} \left\{ \frac{1}{n} \sum_{i=1}^n (Y_i - \theta^T \mathbf{X}_i)^2 + \xi \sum_{j=1}^p |\theta_j| \right\}$$

and $\xi > 0$ is a regularization parameter. Theoretical results (see e.g. Bunea, Tsybakov and Wegkamp [10]) indicate that ξ should be of the order $\xi^* = \sigma \sqrt{\log(p)/n}$. Throughout, σ is assumed to be known, and we let $\xi = \xi^*/3$, since this choice is known to give good practical results. The Nadaraya-Watson kernel estimate will be denoted by $\hat{F}_{\text{NW}}$. It is defined by

$$\hat{F}_{\text{NW}}(\mathbf{x}) = \frac{\sum_{i=1}^n Y_i K_h(\mathbf{x} - \mathbf{X}_i)}{\sum_{i=1}^n K_h(\mathbf{x} - \mathbf{X}_i)}$$

for some nonnegative kernel K on $\mathbb{R}^p$ and $K_h(\mathbf{z}) = K(\mathbf{z}/h)/h$. In the experiments, we let K be the Gaussian kernel $K(\mathbf{z}) = \exp(-\mathbf{z}^T \mathbf{z})$ and chose the smoothing parameter h via a classical leave-one-out procedure on the grid $\mathcal{G} = \{0.75^k, k = 0, \dots, \lfloor \log(n) \rfloor\}$, see, e.g., Györfi, Kohler, Krzyżak and Walk [24] (notation $\lfloor \cdot \rfloor$ stands for the floor function). Finally, the estimation

procedure advocated in Härdle, Hall and Ichimura [25] takes the form

$$\hat{F}_{\text{HHI}}(\mathbf{x}) = \frac{\sum_{i=1}^n Y_i G_{\hat{h}} \left(\hat{\theta}_{\text{HHI}}^T (\mathbf{x} - \mathbf{X}_i) \right)}{\sum_{i=1}^n G_{\hat{h}} \left(\hat{\theta}_{\text{HHI}}^T (\mathbf{x} - \mathbf{X}_i) \right)}$$

for some kernel G on $\mathbb{R}$, with $G_h(\mathbf{z}) = G(\mathbf{z}/h)/h$ and

$$\left(\hat{h}, \hat{\theta}_{\text{HHI}} \right) \in \arg \min_{h>0, \theta \in \mathbb{R}^p} \sum_{i=1}^n \left[Y_i - \frac{\sum_{j \neq i} Y_j G_h \left(\theta^T (\mathbf{X}_j - \mathbf{X}_i) \right)}{\sum_{j \neq i} G_h \left(\theta^T (\mathbf{X}_j - \mathbf{X}_i) \right)} \right]^2.$$

All calculations were performed with the Gaussian kernel. We used the grid $\mathcal{G}$ for the optimization with respect to h , whereas the best search for θ was implemented via a pathwise coordinate optimization.

The various methods were tested for the general regression model

$$Y_i = F(\mathbf{X}_i) + W_i, \quad i = 1, \dots, n,$$

for three different choices of F (single-index or not) and two values of n , namely $n = 50$ and $n = 100$. In each of these models, the observations $\mathbf{X}_i$ take values in $\mathbb{R}^p$, with $p = 10$ and $p = 50$, and have independent components uniformly distributed on $[-1, 1]$. The noise variables $W_1, \dots, W_n$ are independently distributed according to a Gaussian $\mathcal{N}(0, \sigma^2)$, with $\sigma = 0.2$. It is worth pointing out that for $n = 50$ and $p = 50$, p and n are of the same order, which means that the setting is nonasymptotic. It is essentially in this case that the use of estimates tailored to sparsity, which reduce the variance, is expected to improve the performance over generalist methods. On the other hand, the situation $n = 100$ and $p = 10$ is less difficult and mimics the asymptotic setting.

The three examined functions $F(\mathbf{x})$, for $\mathbf{x} = (x_1, \dots, x_p)$, were the following ones:

[Model 1] A linear model $F_{\text{Linear}}(\mathbf{x}) = 2\theta^{*T} \mathbf{x}$.

[Model 2] A single-index function $F_{\text{SI}}(\mathbf{x}) = 2(\theta^{*T} \mathbf{x})^2 + \theta^{*T} \mathbf{x}$.

[Model 3] A purely nonparametric model $F_{\text{NP}}(\mathbf{x}) = 2|x_2|\sqrt{|x_1|} - x_3^3$,

where, in the first and second model, $\theta^* = (0.5, 0.5, 0, \dots, 0)^T$. Thus, in **[Model 1]** and **[Model 2]**, even if the ambient dimension is large, the intrinsic dimension of the model is in fact equal to 2.

$n = 50$	$p = 10$	$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
F_{Linear}	median	0.061	0.063	0.046	0.293
	mean	0.061	0.063	0.047	0.290
	s.d.	0.016	0.014	0.011	0.063
F_{SI}	median	0.050	0.067	0.307	0.198
	mean	0.069	0.080	0.338	0.208
	s.d.	0.081	0.057	0.082	0.072
F_{NP}	median	0.375	0.405	0.830	0.354
	mean	0.402	0.407	0.890	0.336
	s.d.	0.166	0.110	0.176	0.006
$n = 100$	$p = 10$	$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
F_{Linear}	median	0.053	0.051	0.042	0.227
	mean	0.056	0.050	0.043	0.237
	s.d.	0.011	0.006	0.004	0.044
F_{SI}	median	0.047	0.052	0.332	0.209
	mean	0.049	0.053	0.337	0.218
	s.d.	0.009	0.012	0.063	0.045
F_{NP}	median	0.305	0.343	0.793	0.333
	mean	0.321	0.338	0.833	0.324
	s.d.	0.092	0.042	0.145	0.041

Table 1: Numerical results for the simulated data, with $n = 50$ and $n = 100$, $p = 10$. The characters in bold indicate the best performance.

For each experiment, a learning set of size n was generated to compute the estimates and their performance, in terms of mean square prevision error, was evaluated on a separate test set of the same size. The results are shown in Table 1 ($p = 10$) and Table 2 ($p = 50$). As each experiment was repeated 20 times, these tables report the median, the mean and the standard deviation (s.d.) of the prevision error of each procedure.

Some comments are in order. First, we note without surprise that:

1. The Lasso performs well in the linear setting [**Model 1**].
2. The single-index methods $\hat{F}_{\text{Fourier}}$ and $\hat{F}_{\text{HHI}}$ are the best ones when the targeted regression function really involves a single-index model [**Model 2**].
3. The kernel method gives good results in the purely nonparametric setting [**Model 3**].

$n = 50$	$p = 50$	$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
F_{Linear}	median	0.057	1.156	0.060	0.507
	mean	0.095	1.124	0.066	0.533
	s.d.	0.143	0.241	0.026	0.081
F_{SI}	median	0.050	0.502	0.795	0.308
	mean	0.051	0.539	0.776	0.326
	s.d.	0.011	0.200	0.208	0.109
F_{NP}	median	0.358	0.788	1.910	0.374
	mean	0.504	0.771	1.931	0.391
	s.d.	0.320	0.168	0.468	0.101
$n = 100$	$p = 50$	$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
F_{Linear}	median	0.053	0.092	0.050	0.519
	mean	0.054	0.100	0.050	0.508
	s.d.	0.007	0.026	0.006	0.026
F_{SI}	median	0.047	0.242	0.503	0.329
	mean	0.070	0.267	0.502	0.339
	s.d.	0.099	0.111	0.106	0.073
F_{NP}	median	0.361	0.736	1.968	0.418
	mean	0.557	0.765	2.045	0.406
	s.d.	0.519	0.226	0.546	0.076

Table 2: Numerical results for the simulated data, with $n = 50$ and $n = 100$, $p = 50$. The characters in bold indicate the best performance.

Interestingly, $\hat{F}_{\text{Fourier}}$ provides slightly better results than the single-index-tailored estimate $\hat{F}_{\text{HHI}}$, especially for $p = 50$. This observation can be easily explained by the fact that $\hat{F}_{\text{HHI}}$ does not integrate any sparsity information regarding the parameter θ^* , whereas $\hat{F}_{\text{Fourier}}$ tries to focus on the dimension of the active coordinates, which is equal to 2 in this simulation. As a general finding, we retain that $\hat{F}_{\text{Fourier}}$ is the most robust of all the tested procedures.

3.3 Real data

The real-life data sets used in this second series of experiments are from two different sources. The first one, called **AIR-QUALITY** data ($n = 111$, $p = 3$), has been first used by Chambers, Cleveland, Kleiner and Tukey [14] and has been later considered as a benchmark in the study and comparison of single-index models (see, for example, Antoniadis, Grégoire and McKeague [3] and Wang [44], among others). This data set originated from

an environmental study relating $n = 111$ ozone concentration measures at $p = 3$ meteorological variables, namely wind speed, temperature and radiation. The data is available as a package in the software **R** [37], which we employed in all the numerical experiments. The programs are available upon request from the authors.

The second category of data arises from the UC Irvine Machine Learning Repository <http://archive.ics.uci.edu/ml>, where the following packages have been downloaded from:

- **AUTO-MPG** (Quinlan [36], $n = 392$, $p = 7$).
- **CONCRETE** (Yeh [46], $n = 1030$, $p = 8$).
- **HOUSING** (Harrison and Rubinfeld [26], $n = 508$, $p = 13$).
- **SLUMP-1**, **SLUMP-2** and **SLUMP-3**, which correspond to the concrete slump test data introduced by Yeh [47] ($n = 51$, $p = 7$). Since there are 3 different output variables Y in the original data set, we created a single experiment for each of these variables (1 refers to the output “slump”, 2 to the output “flow” and 3 to the output “28-day Compressive Strength”).
- **WINE-RED** and **WINE-WHITE** (Cortez, Cerdeira, Almeida, Matos and Reis [17], $n = 1599$, $n = 4898$, $p = 11$).

We refer to the above-mentioned references for a precise description of the meaning of the variables involved in these data sets. For homogeneity reasons, all data were normalized to force the input variables to lie in $[-1, 1]$ —in accordance with the setting of our method—and to ensure that all output variables have standard deviation 0.5. In two data sets (**AIR-QUALITY** and **AUTO-MPG**) there were some missing values and the corresponding observations were simply removed.

For each method and each of the nine data sets, we randomly split the observations in a learning and a test set of equal sizes, computed the estimate on the learning set, evaluated the prediction error on the test set, and repeated this protocol 20 times. The results are summarized in Table 3.

We see that all the tested methods provide reasonable results on most data sets. The Lasso is very competitive, especially in the nonasymptotic framework. The estimation procedure $\hat{F}_{\text{Fourier}}$ offers outcomes which are similar to the ones of $\hat{F}_{\text{HHI}}$, with a slight advantage for the latter method however. Altogether, $\hat{F}_{\text{Fourier}}$ and $\hat{F}_{\text{HHI}}$ provide the best performance in terms of prediction error in 6 out of 9 experiments. Besides, when it is not the best, the

Data set		$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
AIR QUALITY $n = 111$ $p = 3$	median	0.117	0.099	0.107	0.129
	mean	0.128	0.096	0.113	0.130
	s.d.	0.044	0.029	0.029	0.035
AUTO-MPG $n = 392$ $p = 7$	median	0.044	0.049	0.070	0.068
	mean	0.051	0.050	0.072	0.069
	s.d.	0.017	0.006	0.011	0.009
CONCRETE $n = 1030$ $p = 8$	median	0.089	0.087	0.106	0.094
	mean	0.091	0.087	0.107	0.094
	s.d.	0.008	0.003	0.005	0.004
HOUSING $n = 508$ $p = 11$	median	0.074	0.059	0.086	0.086
	mean	0.076	0.061	0.085	0.088
	s.d.	0.015	0.013	0.012	0.016
SLUMP-1 $n = 51$ $p = 7$	median	0.289	0.171	0.201	0.208
	mean	0.244	0.187	0.213	0.226
	s.d.	0.062	0.050	0.049	0.047
SLUMP-2 $n = 51$ $p = 7$	median	0.219	0.196	0.172	0.215
	mean	0.216	0.194	0.171	0.213
	s.d.	0.053	0.025	0.019	0.022
SLUMP-3 $n = 51$ $p = 7$	median	0.065	0.070	0.053	0.116
	mean	0.073	0.079	0.052	0.126
	s.d.	0.033	0.027	0.010	0.026
WINE-RED $n = 1599$ $p = 11$	median	0.173	0.171	0.183	0.171
	mean	0.174	0.170	0.174	0.183
	s.d.	0.009	0.008	0.007	0.010
WINE-WHITE $n = 4898$ $p = 11$	median	0.191	0.187	0.185	0.184
	mean	0.202	0.188	0.186	0.185
	s.d.	0.045	0.003	0.004	0.004

Table 3: Numerical results for the real-life data sets. The characters in bold indicate the best performance.

method $\hat{F}_{\text{Fourier}}$ is close to the best one, as for example in **SLUMP-3** and **WINE-RED**. As an illustrative example, the plot of the resulting fit of our procedure to the data set **AUTO-MPG** is shown in Figure 1.

Clearly, all data sets under study have a dimension p which is small compared to n . To correct this situation, we ran the same series of experiments by

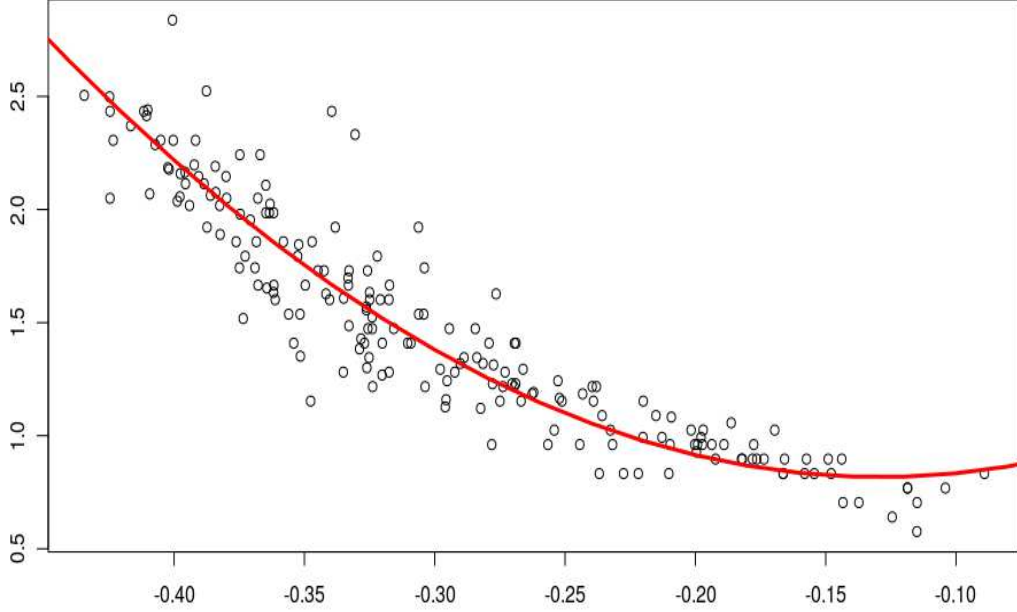


Figure 1: **AUTO-MPG** example: Estimated link function by the method $\hat{F}_{\text{Fourier}}$.

adding some additional irrelevant dimensions to the data. Specifically, the observations were embedded into a space of dimension $p \times 4$ by letting the new fake coordinates follow independent uniform $[0, 1]$ random variables. The results are shown in Table 4. In this nonasymptotic framework, the method $\hat{F}_{\text{HHI}}$ —which is not designed for sparsity—collapses, whereas $\hat{F}_{\text{Fourier}}$ takes a clear advantage over its competitors. In fact, it provides the best results in 3 out of 9 experiments (**AUTO-MPG**, **CONCRETE** and **HOUSING**). Besides, when it is not the best, the method $\hat{F}_{\text{Fourier}}$ is very close to the best one, as for example in **SLUMP-3** and **WINE-RED**.

Thus, as a general conclusion to this experimental section, we may say that our PAC-Bayesian oriented procedure has an excellent predictive ability, even in nonasymptotic/high-dimensional situations. It is fast, robust, and exhibits performance at the level of the gold standard Lasso.

Augmented data set		$\hat{F}_{\text{Fourier}}$	$\hat{F}_{\text{HHI}}$	$\hat{F}_{\text{Lasso}}$	$\hat{F}_{\text{NW}}$
AIR QUALITY $n = 111$ $p = 12$	median	0.172	0.272	0.164	0.281
	mean	0.244	0.291	0.163	0.291
	s.d.	0.163	0.116	0.038	0.046
AUTO-MPG $n = 392$ $p = 28$	median	0.043	0.062	0.085	0.202
	mean	0.044	0.072	0.086	0.203
	s.d.	0.009	0.018	0.008	0.014
CONCRETE $n = 1030$ $p = 32$	median	0.087	0.093	0.113	0.245
	mean	0.087	0.094	0.112	0.094
	s.d.	0.007	0.008	0.005	0.009
HOUSING $n = 508$ $p = 44$	median	0.071	0.199	0.092	0.226
	mean	0.075	0.181	0.095	0.227
	s.d.	0.023	0.084	0.013	0.018
SLUMP-1 $n = 51$ $p = 44$	median	0.270	0.426	0.276	0.271
	mean	0.290	0.409	0.274	0.262
	s.d.	0.101	0.079	0.055	0.042
SLUMP-2 $n = 51$ $p = 44$	median	0.276	0.332	0.195	0.253
	mean	0.285	0.349	0.198	0.254
	s.d.	0.075	0.063	0.043	0.034
SLUMP-3 $n = 51$ $p = 28$	median	0.079	0.371	0.061	0.372
	mean	0.082	0.361	0.058	0.279
	s.d.	0.025	0.079	0.013	0.031
WINE-RED $n = 1599$ $p = 44$	median	0.178	0.222	0.172	0.245
	mean	0.176	0.226	0.174	0.246
	s.d.	0.085	0.033	0.006	0.029
WINE-WHITE $n = 4898$ $p = 11$	median	0.199	0.239	0.187	0.252
	mean	0.204	0.256	0.188	0.260
	s.d.	0.091	0.041	0.005	0.019

Table 4: Numerical results for the real-life data sets augmented with noise variables. The characters in bold indicate the best performance.

4 Proofs

4.1 Preliminary results

Throughout this section, we let π be the prior probability measure on $\mathbb{R}^p \times \mathcal{F}_n(C+1)$ equipped with its canonical Borel σ -field. Recall that $\mathcal{F}_n(C+1) \subset \mathcal{F}$ and that, for each $f \in \mathcal{F}_n(C+1)$, we have $\|f\|_\infty \leq C+1$.

Besides, since $\mathbb{E}[Y|\mathbf{X}] = f^*(\theta^{*T}\mathbf{X})$ almost surely, we note once and for all that for all $(\theta, f) \in \mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$,

$$\begin{aligned} R(\theta, f) - R(\theta^*, f^*) &= \mathbb{E} [Y - f(\theta^T \mathbf{X})]^2 - \mathbb{E} [Y - f^*(\theta^{*T} \mathbf{X})]^2 \\ &= \mathbb{E} [f(\theta^T \mathbf{X}) - f^*(\theta^{*T} \mathbf{X})]^2 \end{aligned}$$

(Pythagora's theorem). We start with four technical lemmas. Lemma 4.1 is a version of Bernstein's inequality, whose proof can be found in Massart [31, Chapter 2, inequality (2.21)]. Lemma 4.2 is a classical result, whose proof can be found, for example, in Catoni [13, page 4]. For a random variable Z , the notation $(Z)_+$ means the positive part of Z .

Lemma 4.1 *Let $T_1, \dots, T_n$ be independent real-valued random variables. Assume that there exist two positive constants v and w such that, for all integers $k \geq 2$,*

$$\sum_{i=1}^n \mathbb{E} [(T_i)_+^k] \leq \frac{k!}{2} v w^{k-2}.$$

Then, for any $\zeta \in]0, 1/w[$,

$$\mathbb{E} \left[\exp \left(\zeta \sum_{i=1}^n [T_i - \mathbb{E} T_i] \right) \right] \leq \exp \left(\frac{v \zeta^2}{2(1 - w \zeta)} \right).$$

Given a measurable space $(E, \mathcal{E})$ and two probability measures μ_1 and μ_2 on $(E, \mathcal{E})$, we denote by $\mathcal{K}(\mu_1, \mu_2)$ the Kullback-Leibler divergence of μ_1 with respect to μ_2 , defined by

$$\mathcal{K}(\mu_1, \mu_2) = \begin{cases} \int \log \left(\frac{d\mu_1}{d\mu_2} \right) d\mu_1 & \text{if } \mu_1 \ll \mu_2, \\ \infty & \text{otherwise.} \end{cases}$$

(Notation $\mu_1 \ll \mu_2$ means “ μ_1 is absolutely continuous with respect to μ_2 ”.)

Lemma 4.2 *Let $(E, \mathcal{E})$ be a measurable space. For any probability measure μ on $(E, \mathcal{E})$ and any measurable function $h : E \rightarrow \mathbb{R}$ such that $\int (\exp \circ h) d\mu < \infty$, we have*

$$\log \int (\exp \circ h) d\mu = \sup_m \left(\int h dm - \mathcal{K}(m, \mu) \right), \quad (4.1)$$

where the supremum is taken over all probability measures on $(E, \mathcal{E})$ and, by convention, $\infty - \infty = -\infty$. Moreover, as soon as h is bounded from above

on the support of μ , the supremum with respect to m on the right-hand side of (4.1) is reached for the Gibbs distribution g given by

$$\frac{dg}{d\mu}(e) = \frac{\exp[h(e)]}{\int (\exp \circ h) d\mu}, \quad e \in E.$$

Lemma 4.3 Assume that Assumption **N** holds. Set $w = 8(2C + 1) \max[L, 2C + 1]$ and take

$$\lambda \in \left] 0, \frac{n}{w + [(2C + 1)^2 + 4\sigma^2]} \right[.$$

Then, for all $\delta \in]0, 1[$ and any data-dependent probability measure $\hat{\rho}$ absolutely continuous with respect to π we have, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}, \hat{f}) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \left(R_n(\hat{\theta}, \hat{f}) - R_n(\theta^*, f^*) + \frac{\log\left(\frac{d\hat{\rho}}{d\pi}(\hat{\theta}, \hat{f})\right) + \log\left(\frac{1}{\delta}\right)}{\lambda} \right), \end{aligned}$$

where the pair $(\hat{\theta}, \hat{f})$ is distributed according to $\hat{\rho}$.

Proof of Lemma 4.3. Fix $\theta \in \mathcal{S}_{1,+}^p$ and $f \in \mathcal{F}_n(C + 1)$. The proof starts with an application of Lemma 4.1 to the random variables

$$T_i = - (Y_i - f(\theta^T \mathbf{X}_i))^2 + (Y_i - f^*(\theta^{*T} \mathbf{X}_i))^2, \quad i = 1, \dots, n.$$

Note that these random variables are independent, identically distributed, and that

$$\begin{aligned} & \sum_{i=1}^n \mathbb{E} T_i^2 \\ & = \sum_{i=1}^n \mathbb{E} \left\{ [2Y_i - f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)]^2 [f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)]^2 \right\} \\ & = \sum_{i=1}^n \mathbb{E} \left\{ [2W_i + f^*(\theta^{*T} \mathbf{X}_i) - f(\theta^T \mathbf{X}_i)]^2 [f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)]^2 \right\} \\ & \leq \sum_{i=1}^n \mathbb{E} \left\{ [4W_i^2 + (2C + 1)^2] [f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)]^2 \right\} \\ & \quad (\text{since } \mathbb{E}[W_i | \mathbf{X}_i] = 0). \end{aligned}$$

Thus, by Assumption **N**,

$$\begin{aligned} \sum_{i=1}^n \mathbb{E} T_i^2 &\leq [(2C+1)^2 + 4\sigma^2] \sum_{i=1}^n \mathbb{E} [f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)]^2 \\ &\leq v, \end{aligned}$$

where we set

$$v = 2n[(2C+1)^2 + 4\sigma^2] [R(\theta, f) - R(\theta^*, f^*)].$$

More generally, for all integers $k \geq 3$,

$$\begin{aligned} \sum_{i=1}^n \mathbb{E} [(T_i)_+^k] &\leq \sum_{i=1}^n \mathbb{E} \left\{ |2Y_i - f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^k |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^k \right\} \\ &= \sum_{i=1}^n \mathbb{E} \left\{ |2W_i + f^*(\theta^{*T} \mathbf{X}_i) - f(\theta^T \mathbf{X}_i)|^k |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^k \right\} \\ &\leq 2^{k-1} \sum_{i=1}^n \mathbb{E} \left\{ [2^k |W_i|^k + (2C+1)^k] (2C+1)^{k-2} |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^2 \right\}. \end{aligned}$$

In the last inequality, we used the fact that $|a+b|^k \leq 2^{k-1}(|a|^k + |b|^k)$ together with

$$\begin{aligned} &|f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^k \\ &= |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^{k-2} \times |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^2 \\ &\leq (2C+1)^{k-2} |f(\theta^T \mathbf{X}_i) - f^*(\theta^{*T} \mathbf{X}_i)|^2. \end{aligned}$$

Therefore, by Assumption **N**,

$$\begin{aligned} \sum_{i=1}^n \mathbb{E} [(T_i)_+^k] &\leq \sum_{i=1}^n [2^{2k-2} k! \sigma^2 L^{k-2} + 2^{k-1} (2C+1)^k] (2C+1)^{k-2} [R(\theta, f) - R(\theta^*, f^*)] \\ &= v \times \frac{[2^{2k-2} k! \sigma^2 L^{k-2} + 2^{k-1} (2C+1)^k] (2C+1)^{k-2}}{[(2C+1)^2 + 4\sigma^2]} \\ &\leq v \times \frac{8^{k-2} k! \max[L^{k-2}, (2C+1)^{k-2}] (2C+1)^{k-2}}{2} \\ &= \frac{k!}{2} v w^{k-2}, \end{aligned}$$

with $w = 8(2C + 1) \max[L, 2C + 1]$.

Thus, for any inverse temperature parameter $\lambda \in]0, n/w[$, taking $\zeta = \lambda/n$, we may write by Lemma 4.1

$$\begin{aligned} & \mathbb{E} \left\{ \exp [\lambda (R(\theta, f) - R(\theta^*, f^*) - R_n(\theta, f) + R_n(\theta^*, f^*))] \right\} \\ & \leq \exp \left(\frac{v\lambda^2}{2n^2(1 - \frac{w\lambda}{n})} \right). \end{aligned}$$

Therefore, using the definition of v , we obtain

$$\begin{aligned} & \mathbb{E} \left\{ \exp \left[\left(\lambda - \frac{\lambda^2 [(2C + 1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta, f) - R(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. + \lambda (-R_n(\theta, f) + R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] \right\} \leq \delta. \end{aligned}$$

Next, we use a standard PAC-Bayesian approach (Catoni [12, 13], Audibert [4] and Alquier [1]). Let us remind the reader that π is a prior probability measure on the set $\mathcal{S}_{1,+}^p \times \mathcal{F}_n(C + 1)$. We have

$$\begin{aligned} & \int \mathbb{E} \left\{ \exp \left[\left(\lambda - \frac{\lambda^2 [(2C + 1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta, f) - R(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. + \lambda (-R_n(\theta, f) + R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] \right\} d\pi(\theta, f) \leq \delta \end{aligned}$$

and consequently, using Fubini's theorem,

$$\begin{aligned} & \mathbb{E} \left\{ \int \exp \left[\left(\lambda - \frac{\lambda^2 [(2C + 1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta, f) - R(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. + \lambda (-R_n(\theta, f) + R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] d\pi(\theta, f) \right\} \leq \delta. \end{aligned}$$

Therefore, for any data-dependent posterior probability measure $\hat{\rho}$ absolutely

continuous with respect to π , adopting the convention $\infty \times 0 = 0$,

$$\begin{aligned} & \mathbb{E} \left\{ \int \exp \left[\left(\lambda - \frac{\lambda^2 [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta, f) - R(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. + \lambda (-R_n(\theta, f) + R_n(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. - \log \left(\frac{d\hat{\rho}}{d\pi}(\theta, f) \right) - \log \left(\frac{1}{\delta} \right) \right] d\hat{\rho}(\theta, f) \right\} \\ & \leq \delta. \end{aligned}$$

Recalling that $\mathbf{P}^{\otimes n}$ stands for the distribution of the sample $\mathcal{D}_n$, the latter inequality can be more conveniently written as

$$\begin{aligned} & \mathbb{E}_{\mathcal{D}_n \sim \mathbf{P}^{\otimes n}} \mathbb{E}_{(\hat{\theta}, \hat{f}) \sim \hat{\rho}} \left\{ \exp \left[\left(\lambda - \frac{\lambda^2 [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\hat{\theta}, \hat{f}) - R(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. + \lambda (-R_n(\hat{\theta}, \hat{f}) + R_n(\theta^*, f^*)) - \log \left(\frac{d\hat{\rho}}{d\pi}(\hat{\theta}, \hat{f}) \right) - \log \left(\frac{1}{\delta} \right) \right] \right\} \\ & \leq \delta. \end{aligned}$$

Thus, using the elementary inequality $\exp(\lambda x) \geq \mathbf{1}_{\mathbb{R}_+}(x)$ we obtain, with probability at most δ ,

$$\begin{aligned} & \left(1 - \frac{\lambda [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\hat{\theta}, \hat{f}) - R(\theta^*, f^*)) \\ & \geq R_n(\hat{\theta}, \hat{f}) - R_n(\theta^*, f^*) + \frac{\log \left(\frac{d\hat{\rho}}{d\pi}(\hat{\theta}, \hat{f}) \right) + \log \left(\frac{1}{\delta} \right)}{\lambda}, \end{aligned}$$

where the probability is evaluated with respect to the distribution $\mathbf{P}^{\otimes n}$ of the data $\mathcal{D}_n$ and the conditional probability measure $\hat{\rho}$. Put differently, letting

$$\lambda \in \left] 0, \frac{n}{w + [(2C+1)^2 + 4\sigma^2]} \right[,$$

we have, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}, \hat{f}) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda [(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \left(R_n(\hat{\theta}, \hat{f}) - R_n(\theta^*, f^*) + \frac{\log \left(\frac{d\hat{\rho}}{d\pi}(\hat{\theta}, \hat{f}) \right) + \log \left(\frac{1}{\delta} \right)}{\lambda} \right). \end{aligned}$$

This concludes the proof of Lemma 4.3. ■

Lemma 4.4 *Under the conditions of Lemma 4.3 we have, with probability at least $1 - \delta$,*

$$\begin{aligned} & \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \\ & \leq \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}\right) \left(\int R(\theta, f) d\hat{\rho}(\theta, f) - R(\theta^*, f^*)\right) \\ & \quad + \frac{\mathcal{K}(\hat{\rho}, \pi) + \log\left(\frac{1}{\delta}\right)}{\lambda}. \end{aligned}$$

Proof of Lemma 4.4. The beginning of the proof is similar to the one of Lemma 4.3. More precisely, we apply Lemma 4.1 with $T_i = (Y_i - f(\theta^T \mathbf{X}_i))^2 - (Y_i - f^*(\theta^{*T} \mathbf{X}_i))^2$ and obtain, for any inverse temperature parameter $\lambda \in]0, n/w[$,

$$\begin{aligned} & \mathbb{E} \left\{ \exp \left[\lambda (R(\theta^*, f^*) - R(\theta, f) - R_n(\theta^*, f^*) + R_n(\theta, f)) \right] \right\} \\ & \leq \exp \left(\frac{v\lambda^2}{2n^2(1 - \frac{w\lambda}{n})} \right). \end{aligned}$$

Thus, using the definition of v ,

$$\begin{aligned} & \mathbb{E} \left\{ \exp \left[\left(\lambda + \frac{\lambda^2[(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta^*, f^*) - R(\theta, f)) \right. \right. \\ & \quad \left. \left. + \lambda (R_n(\theta, f) - R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] \right\} \leq \delta. \end{aligned}$$

Integrating with respect to π leads to

$$\begin{aligned} & \int \mathbb{E} \left\{ \exp \left[\left(\lambda + \frac{\lambda^2[(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta^*, f^*) - R(\theta, f)) \right. \right. \\ & \quad \left. \left. + \lambda (R_n(\theta, f) - R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] \right\} d\pi(\theta, f) \leq \delta \end{aligned}$$

whence, by Fubini's theorem,

$$\begin{aligned} & \mathbb{E} \left\{ \int \exp \left[\left(\lambda + \frac{\lambda^2[(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta^*, f^*) - R(\theta, f)) \right. \right. \\ & \quad \left. \left. + \lambda (R_n(\theta, f) - R_n(\theta^*, f^*)) - \log \left(\frac{1}{\delta} \right) \right] d\pi(\theta, f) \right\} \leq \delta. \end{aligned}$$

Thus, for any data-dependent posterior probability measure $\hat{\rho}$ absolutely continuous with respect to π ,

$$\begin{aligned} & \mathbb{E} \left\{ \int \exp \left[\left(\lambda + \frac{\lambda^2 [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta^*, f^*) - R(\theta, f)) \right. \right. \\ & \quad \left. \left. + \lambda (R_n(\theta, f) - R_n(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. - \log \left(\frac{d\hat{\rho}}{d\pi}(\theta, f) \right) - \log \left(\frac{1}{\delta} \right) \right] d\hat{\rho}(\theta, f) \right\} \\ & \leq \delta. \end{aligned}$$

Therefore, by Jensen's inequality,

$$\begin{aligned} & \mathbb{E} \left\{ \exp \int \left[\left(\lambda + \frac{\lambda^2 [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) (R(\theta^*, f^*) - R(\theta, f)) \right. \right. \\ & \quad \left. \left. + \lambda (R_n(\theta, f) - R_n(\theta^*, f^*)) \right. \right. \\ & \quad \left. \left. - \log \left(\frac{d\hat{\rho}}{d\pi}(\theta, f) \right) - \log \left(\frac{1}{\delta} \right) \right] d\hat{\rho}(\theta, f) \right\} \\ & = \mathbb{E} \left\{ \exp \left[\left(\lambda + \frac{\lambda^2 [(2C+1)^2 + 4\sigma^2]}{n(1 - \frac{w\lambda}{n})} \right) \left(R(\theta^*, f^*) - \int R(\theta, f) d\hat{\rho}(\theta, f) \right) \right. \right. \\ & \quad \left. \left. + \lambda \left(\int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \right) \right. \right. \\ & \quad \left. \left. - \mathcal{K}(\hat{\rho}, \pi) - \log \left(\frac{1}{\delta} \right) \right] \right\} \\ & \leq \delta. \end{aligned}$$

Consequently, by the elementary inequality $\exp(\lambda x) \geq \mathbf{1}_{\mathbb{R}_+}(x)$, we obtain, with probability at most δ ,

$$\begin{aligned} & \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \\ & \geq \left(1 + \frac{\lambda [(2C+1)^2 + 4\sigma^2]}{n - w\lambda} \right) \left(\int R(\theta, f) d\hat{\rho}(\theta, f) - R(\theta^*, f^*) \right) \\ & \quad + \frac{\mathcal{K}(\hat{\rho}, \pi) + \log \left(\frac{1}{\delta} \right)}{\lambda}. \end{aligned}$$

Equivalently, with probability at least $1 - \delta$,

$$\begin{aligned} & \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \\ & \leq \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}\right) \left(\int R(\theta, f) d\hat{\rho}(\theta, f) - R(\theta^*, f^*)\right) \\ & \quad + \frac{\mathcal{K}(\hat{\rho}, \pi) + \log\left(\frac{1}{\delta}\right)}{\lambda}. \end{aligned}$$

■

4.2 Proof of Theorem 2.1

The proof starts with an application of Lemma 4.3 with $\hat{\rho} = \hat{\rho}_\lambda$ (the Gibbs distribution) as posterior distribution. More precisely, we know that, with probability larger than $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \left(R_n(\hat{\theta}_\lambda, \hat{f}_\lambda) - R_n(\theta^*, f^*) \right. \\ & \quad \left. + \frac{\log\left(\frac{d\hat{\rho}_\lambda}{d\pi}(\hat{\theta}_\lambda, \hat{f}_\lambda)\right) + \log\left(\frac{1}{\delta}\right)}{\lambda} \right), \end{aligned}$$

where the probability is evaluated with respect to the distribution $\mathbf{P}^{\otimes n}$ of the data $\mathcal{D}_n$ and the conditional probability measure $\hat{\rho}_\lambda$. Observe that

$$\begin{aligned} \log\left(\frac{d\hat{\rho}_\lambda}{d\pi}(\hat{\theta}_\lambda, \hat{f}_\lambda)\right) &= \log\left(\frac{\exp[-\lambda R_n(\hat{\theta}_\lambda, \hat{f}_\lambda)]}{\int \exp[-\lambda R_n(\theta, f)] d\pi(\theta, f)}\right) \\ &= -\lambda R_n(\hat{\theta}_\lambda, \hat{f}_\lambda) - \log \int \exp[-\lambda R_n(\theta, f)] d\pi(\theta, f). \end{aligned}$$

Consequently, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{\lambda \left(1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}\right)} \left(-\log \int \exp[-\lambda R_n(\theta, f)] d\pi(\theta, f) \right. \\ & \quad \left. - \lambda R_n(\theta^*, f^*) + \log\left(\frac{1}{\delta}\right) \right). \end{aligned}$$

Next, using Lemma 4.2 we deduce that, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\hat{\rho}} \left\{ \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \right. \\ & \quad \left. + \frac{\mathcal{K}(\hat{\rho}, \pi) + \log\left(\frac{1}{\delta}\right)}{\lambda} \right\}, \end{aligned}$$

where the infimum is taken over all probability measures on $\mathcal{S}_{1,+}^p \times \mathcal{F}_n(C+1)$. In particular, letting $\mathcal{M}(I, M)$ be the set of all probability measures on $\mathcal{S}_{1,+}^p(I) \times \mathcal{F}_M(C+1)$, we have, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \inf_{\hat{\rho} \in \mathcal{M}(I, M)} \left\{ \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \right. \\ & \quad \left. + \frac{\mathcal{K}(\hat{\rho}, \pi) + \log\left(\frac{1}{\delta}\right)}{\lambda} \right\}. \end{aligned}$$

Next, observe that, for $\hat{\rho} \in \mathcal{M}(I, M)$,

$$\begin{aligned} \mathcal{K}(\hat{\rho}, \pi) &= \mathcal{K}(\hat{\rho}, \mu \otimes \nu) = \mathcal{K}(\hat{\rho}, \mu_I \otimes \nu_M) + \log \left[\frac{\left(1 - \left(\frac{1}{10}\right)^p\right) \left(1 - \left(\frac{1}{10}\right)^n\right) \binom{p}{|I|}}{10^{-|I|-M}} \right] \\ &\leq \mathcal{K}(\hat{\rho}, \mu_I \otimes \nu_M) + \log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right]. \end{aligned} \tag{4.2}$$

Therefore, with probability at least $1 - \delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \inf_{\hat{\rho} \in \mathcal{M}(I, M)} \left\{ \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \right. \\ & \quad \left. + \frac{\mathcal{K}(\hat{\rho}, \mu_I \otimes \nu_M) + \log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right] + \log\left(\frac{1}{\delta}\right)}{\lambda} \right\}. \end{aligned} \tag{4.3}$$

By Lemma 4.4 and inequality (4.2), for any data-dependent distribution $\hat{\rho} \in \mathcal{M}(I, M)$, with probability at least $1 - \delta$,

$$\begin{aligned}
& \int R_n(\theta, f) d\hat{\rho}(\theta, f) - R_n(\theta^*, f^*) \\
& \leq \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}\right) \left(\int R(\theta, f) d\hat{\rho}(\theta, f) - R(\theta^*, f^*)\right) \\
& \quad + \frac{\mathcal{K}(\hat{\rho}, \mu_I \otimes \nu_M) + \log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right] + \log \left(\frac{1}{\delta} \right)}{\lambda}.
\end{aligned} \tag{4.4}$$

Thus, combining inequalities (4.3) and (4.4), we may write, with probability at least $1 - 2\delta$,

$$\begin{aligned}
& R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\
& \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \inf_{\hat{\rho} \in \mathcal{M}(I, M)} \left\{ \right. \\
& \quad \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}\right) \left(\int R(\theta, f) d\hat{\rho}(\theta, f) - R(\theta^*, f^*)\right) \\
& \quad \left. + 2 \frac{\mathcal{K}(\hat{\rho}, \mu_I \otimes \nu_M) + \log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right] + \log \left(\frac{1}{\delta} \right)}{\lambda} \right\}.
\end{aligned} \tag{4.5}$$

For any subset I of $\{1, \dots, p\}$, any positive integer $M \leq n$ and any $\eta, \gamma \in]0, 1/n]$, let the probability measure $\rho_{I, M, \eta, \gamma}$ be defined by

$$d\rho_{I, M, \eta, \gamma}(\theta, f) = d\rho_{I, M, \eta}^1(\theta) d\rho_{I, M, \gamma}^2(f),$$

with

$$\frac{d\rho_{I, M, \eta}^1(\theta)}{d\mu_I}(\theta) \propto \mathbf{1}_{[\|\theta - \theta_{I, M}^*\|_1 \leq \eta]}$$

and

$$\frac{d\rho_{I, M, \gamma}^2(f)}{d\nu_M}(f) \propto \mathbf{1}_{[\|f - f_{I, M}^*\|_M \leq \gamma]}$$

where, for $f = \sum_{j=1}^M \beta_j \varphi_j \in \mathcal{F}_M(C+1)$, we put

$$\|f\|_M = \sum_{j=1}^M j|\beta_j|.$$

With this notation, inequality (4.5) leads to

$$\begin{aligned}
& R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\
& \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \inf_{\eta, \gamma > 0} \left\{ \right. \\
& \quad \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda} \right) \left(\int R(\theta, f) d\rho_{I, M, \eta, \gamma}(\theta, f) - R(\theta^*, f^*) \right) \\
& \quad \left. + 2 \frac{\mathcal{K}(\rho_{I, M, \eta, \gamma}, \mu_I \otimes \nu_M) + \log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right] + \log \left(\frac{1}{\delta} \right)}{\lambda} \right\}. \tag{4.6}
\end{aligned}$$

To finish the proof, we have to control the different terms in (4.6). Note first that

$$\log \left(\frac{p}{|I|} \right) \leq |I| \log \left(\frac{pe}{|I|} \right)$$

and, consequently,

$$\log \left[\frac{\binom{p}{|I|}}{10^{-|I|-M}} \right] \leq |I| \log \left(\frac{pe}{|I|} \right) + (|I| + M) \log 10. \tag{4.7}$$

Next,

$$\begin{aligned}
\mathcal{K}(\rho_{I, M, \eta, \gamma}, \mu_I \otimes \nu_M) &= \mathcal{K}(\rho_{I, M, \eta}^1 \otimes \rho_{I, M, \gamma}^2, \mu_I \otimes \nu_M) \\
&= \mathcal{K}(\rho_{I, M, \eta}^1, \mu_I) + \mathcal{K}(\rho_{I, M, \gamma}^2, \nu_M).
\end{aligned}$$

By technical Lemma 4.5, we know that

$$\mathcal{K}(\rho_{I, M, \eta}^1, \mu_I) \leq (|I| - 1) \log \left(\max \left[|I|, \frac{4}{\eta} \right] \right).$$

Similarly, by technical Lemma 4.6,

$$\mathcal{K}(\rho_{I, M, \gamma}^2, \nu_M) = M \log \left(\frac{C+1}{\gamma} \right).$$

Putting all the pieces together, we are led to

$$\mathcal{K}(\rho_{I, M, \eta, \gamma}, \mu_I \otimes \nu_M) \leq (|I| - 1) \log \left(\max \left[|I|, \frac{4}{\eta} \right] \right) + M \log \left(\frac{C+1}{\gamma} \right). \tag{4.8}$$

Finally, it remains to control the term

$$\int R(\theta, f) d\rho_{I, M, \eta, \gamma}(\theta, f).$$

To this aim, we write

$$\begin{aligned}
& \int R(\theta, f) d\rho_{I,M,\eta,\gamma}(\theta, f) \\
&= \int \mathbb{E} \left[(Y - f(\theta^T \mathbf{X}))^2 \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\
&= \int \mathbb{E} \left[(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) + f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X}) \right. \\
&\quad \left. + f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X}))^2 \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\
&= R(\theta_{I,M}^*, f_{I,M}^*) \\
&\quad + \int \mathbb{E} \left[(f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X}))^2 \right. \\
&\quad + (f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X}))^2 \\
&\quad + 2(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})) \\
&\quad + 2(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) (f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X})) \\
&\quad \left. + 2(f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})) (f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X})) \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\
&:= R(\theta_{I,M}^*, f_{I,M}^*) + \mathbf{A} + \mathbf{B} + \mathbf{C} + \mathbf{D} + \mathbf{E}.
\end{aligned}$$

Computation of C By Fubini's theorem,

$$\begin{aligned}
\mathbf{C} &= \mathbb{E} \left[\int 2(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})) d\rho_{I,M,\eta,\gamma}(\theta, f) \right] \\
&= \mathbb{E} \left\{ \int \left[2(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) \right. \right. \\
&\quad \left. \left. \times \int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})) d\rho_{I,M,\gamma}^2(f) \right] d\rho_{I,M,\eta}^1(\theta) \right\}.
\end{aligned}$$

By the triangle inequality, for $f = \sum_{j=1}^M \beta_j \varphi_j$ and $f_{I,M}^* = \sum_{j=1}^M (\beta_{I,M}^*)_j \varphi_j$, it holds

$$\sum_{j=1}^M j |\beta_j| \leq \sum_{j=1}^M j |\beta_j - (\beta_{I,M}^*)_j| + \sum_{j=1}^M j |(\beta_{I,M}^*)_j|.$$

Since $f_{I,M}^* \in \mathcal{F}_M(C)$, we have $\sum_{j=1}^M j |(\beta_{I,M}^*)_j| \leq C$, so that $\sum_{j=1}^M j |\beta_j| \leq C + 1$ as soon as $\|f - f_{I,M}^*\|_M \leq 1$. This shows that the set

$$\left\{ f = \sum_{j=1}^M \beta_j \varphi_j : \|f - f_{I,M}^*\|_M \leq \gamma \right\}$$

is contained in the support of ν_M . In particular, this implies that $\rho_{I,M,\gamma}^2$ is centered at $f_{I,M}^*$ and, consequently,

$$\int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})) \, d\rho_{I,M,\gamma}^2(f) = 0.$$

This proves that $\mathbf{C} = 0$.

Control of A Clearly,

$$\mathbf{A} \leq \int \sup_{y \in \mathbb{R}} ((f_{I,M}^*(y) - f(y))^2 \, d\rho_{I,M,\gamma}^2(f) \leq \gamma^2.$$

Control of B We have

$$\begin{aligned} \mathbf{B} &= \int \mathbb{E} \left[(f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X}))^2 \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\ &\leq \int \mathbb{E} \left[(\ell(C+1)(\theta_{I,M}^{*T} - \theta^T) \mathbf{X})^2 \right] d\rho_{I,M,\eta}^1(\theta) \\ &\quad (\text{by the mean value theorem}) \\ &\leq \ell^2(C+1)^2 \mathbb{E} [\|\mathbf{X}\|_\infty^2] \int \|\theta_{I,M}^* - \theta\|_1^2 d\rho_{I,M,\eta}^1(\theta) \\ &\leq \ell^2(C+1)^2 \eta^2 \\ &\quad (\text{by Assumption D}). \end{aligned}$$

Control of E Write

$$\begin{aligned} |\mathbf{E}| &\leq 2 \int \mathbb{E} \left[|f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})| \right. \\ &\quad \left. \times |f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X})| \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\ &\leq 2 \int \mathbb{E} \left[|f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X})| \right. \\ &\quad \left. \times \ell(C+1) |(\theta_{I,M}^{*T} - \theta^T) \mathbf{X}| \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \\ &\leq 2 \left(\int \mathbb{E} \left[(f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta_{I,M}^{*T} \mathbf{X}))^2 \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \right)^{\frac{1}{2}} \\ &\quad \left(\int \mathbb{E} \left[(\ell(C+1)(\theta_{I,M}^{*T} - \theta^T) \mathbf{X})^2 \right] d\rho_{I,M,\eta,\gamma}(\theta, f) \right)^{\frac{1}{2}} \\ &\quad (\text{by the Cauchy-Schwarz inequality}) \\ &\leq 2 (\gamma^2)^{\frac{1}{2}} (\ell^2(C+1)^2 \eta^2)^{\frac{1}{2}} \\ &= 2\ell(C+1)\gamma\eta. \end{aligned}$$

Control of D Finally,

$$\begin{aligned}
\mathbf{D} &= 2 \int \mathbb{E} [(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) (f(\theta_{I,M}^{*T} \mathbf{X}) - f(\theta^T \mathbf{X}))] d\rho_{I,M,\eta,\gamma}(\theta, f) \\
&= 2 \int \mathbb{E} [(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X}))] d\rho_{I,M,\eta}^1(\theta) \\
&\quad (\text{since } \int f d\rho_{I,M,\gamma}^2(f) = f_{I,M}^*) \\
&= 2 \mathbb{E} \left[(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X})) \int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})) d\rho_{I,M,\eta}^1(\theta) \right] \\
&\leq 2 \sqrt{\mathbb{E} [(Y - f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}))^2]} \\
&\quad \times \sqrt{\mathbb{E} \left[\int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})) d\rho_{I,M,\eta}^1(\theta) \right]^2} \\
&\quad (\text{by the Cauchy-Schwarz inequality}) \\
&= 2 \sqrt{R(\theta_{I,M}^*, f_{I,M}^*)} \sqrt{\mathbb{E} \left[\int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})) d\rho_{I,M,\eta}^1(\theta) \right]^2}.
\end{aligned}$$

The inequality

$$\begin{aligned}
|f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})| &\leq \ell(C+1) |(\theta_{I,M}^{*T} - \theta^T) \mathbf{X}| \\
&\leq \ell(C+1) \|\theta_{I,M}^* - \theta\|_1
\end{aligned}$$

leads to

$$\begin{aligned}
&\left[\int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})) d\rho_{I,M,\eta}^1(\theta) \right]^2 \\
&\leq \ell^2(C+1)^2 \left[\int \|\theta_{I,M}^* - \theta\|_1 d\rho_{I,M,\eta}^1(\theta) \right]^2.
\end{aligned}$$

Consequently,

$$\left[\int (f_{I,M}^*(\theta_{I,M}^{*T} \mathbf{X}) - f_{I,M}^*(\theta^T \mathbf{X})) d\rho_{I,M,\eta}^1(\theta) \right]^2 \leq \ell^2(C+1)^2 \eta^2,$$

and therefore

$$\begin{aligned}
\mathbf{D} &\leq 2\ell(C+1)\eta\sqrt{R(0,0)/2} \\
&\leq \sqrt{2}\ell(C+1)\eta\sqrt{C^2 + \sigma^2}.
\end{aligned}$$

Thus, taking $\eta = \gamma = 1/n$ and putting all the pieces together, we obtain

$$\mathbf{A} + \mathbf{B} + \mathbf{C} + \mathbf{D} + \mathbf{E} \leq \frac{\Xi_1}{n},$$

where Ξ_1 is a positive constant, function of C , σ and ℓ . Combining this inequality with (4.6)-(4.8) yields, with probability larger than $1 - 2\delta$,

$$\begin{aligned} & R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) \\ & \leq \frac{1}{1 - \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda}} \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \left\{ \left(1 + \frac{\lambda[(2C+1)^2 + 4\sigma^2]}{n - w\lambda} \right) \left(R(\theta_{I,M}^*, f_{I,M}^*) \right. \right. \\ & \quad \left. \left. - R(\theta^*, f^*) + \frac{\Xi_1}{n} \right) + 2 \frac{M \log(10(C+1)n) + |I| \log(40epn) + \log(\frac{1}{\delta})}{\lambda} \right\}. \end{aligned}$$

Choosing finally

$$\lambda = \frac{n}{w + 2[(2C+1)^2 + 4\sigma^2]},$$

we obtain that there exists a positive constant Ξ_2 , function of L , C , σ and ℓ such that, with probability at least $1 - 2\delta$,

$$\begin{aligned} R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) & \leq \Xi_2 \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \left\{ R(\theta_{I,M}^*, f_{I,M}^*) - R(\theta^*, f^*) \right. \\ & \quad \left. + \frac{M \log(10Cn) + |I| \log(40epn) + \log(\frac{1}{\delta})}{n} \right\}. \end{aligned}$$

This concludes the proof of Theorem 2.1.

4.3 Proof of Corollary 2.1

We already know, by Theorem 2.1, that with probability at least $1 - \delta$,

$$\begin{aligned} R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) & \leq \Xi \inf_{\substack{I \subset \{1, \dots, p\} \\ 1 \leq M \leq n}} \left\{ R(\theta_{I,M}^*, f_{I,M}^*) - R(\theta^*, f^*) \right. \\ & \quad \left. + \frac{M \log(Cn) + |I| \log(pn) + \log(\frac{2}{\delta})}{n} \right\}. \end{aligned}$$

By definition, for all $(\theta, f) \in \mathcal{S}_{1,+}^p(I) \times \mathcal{F}_M(C)$,

$$R(\theta_{I,M}^*, f_{I,M}^*) \leq R(\theta, f).$$

In particular, if $I^\star$ is such that $\theta^\star \in \mathcal{S}_{1,+}^p(I^\star)$, then

$$R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^\star, f^\star) \leq \Xi \inf_{\substack{1 \leq M \leq n \\ f \in \mathcal{F}_M(C)}} \left\{ R(\theta^\star, f) - R(\theta^\star, f^\star) + \frac{M \log(Cn) + |I^\star| \log(pn) + \log\left(\frac{2}{\delta}\right)}{n} \right\}. \quad (4.9)$$

Observe that, for any $f \in \mathcal{F}_M(C)$,

$$\begin{aligned} R(\theta^\star, f) - R(\theta^\star, f^\star) &= \int_{\mathbb{R}^p} [f(\theta^{\star T} \mathbf{x}) - f^\star(\theta^{\star T} \mathbf{x})]^2 d\mathbf{P}(\mathbf{x}, y) \\ &\leq B^2 \int_{-1}^1 [f(t) - f^\star(t)]^2 dt. \end{aligned}$$

Since $f^\star \in L_2([-1, 1])$, we may write

$$f^\star = \sum_{j=1}^{\infty} \beta_j^\star \varphi_j$$

and apply (4.9) with

$$f = \sum_{j=1}^M \beta_j^\star \varphi_j.$$

In order to do so, we just need to check that $f \in \mathcal{F}_M(C)$, that is

$$\sum_{j=1}^M j |\beta_j^\star| \leq C.$$

But, by the Cauchy-Schwarz inequality,

$$\begin{aligned} \sum_{j=1}^M j |\beta_j^\star| &= \sum_{j=1}^M j^k |\beta_j^\star| j^{1-k} \\ &\leq \sqrt{\sum_{j=1}^M j^{2k} (\beta_j^\star)^2} \sqrt{\sum_{j=1}^M j^{2-2k}}. \end{aligned}$$

Thus,

$$\begin{aligned} \sum_{j=1}^M j |\beta_j^*| &\leq \frac{\pi}{\sqrt{6}} \sqrt{\sum_{j=1}^M j^{2k} (\beta_j^*)^2} \\ &\quad (\text{since, by assumption, } k \geq 2) \\ &\leq C \\ &\quad (\text{since } f^* \in \mathcal{W}(k, 6C^2/\pi^2)). \end{aligned}$$

Next, with this choice of f ,

$$\int_{-1}^1 [f(t) - f^*(t)]^2 dt \leq \Lambda M^{-2k}$$

for some positive constant Λ depending only on k and C (see for instance Tsybakov [42]). Therefore, inequality (4.9) leads to

$$\begin{aligned} R(\hat{\theta}_\lambda, \hat{f}_\lambda) - R(\theta^*, f^*) &\leq \Xi \inf_{1 \leq M \leq n} \left\{ \Lambda M^{-2k} \right. \\ &\quad \left. + \frac{M \log(Cn) + |I^*| \log(pn) + \log\left(\frac{2}{\delta}\right)}{n} \right\}. \end{aligned} \tag{4.10}$$

Letting $\lceil \cdot \rceil$ be the ceiling function and choosing $M = \lceil (n/\log(Cn))^{\frac{1}{2\beta+1}} \rceil$ in (4.10) concludes the proof.

4.4 Some technical lemmas

Lemma 4.5 *For any subset I of $\{1, \dots, p\}$, any positive integer $M \leq n$ and any $\eta \in]0, 1/n]$, let the probability measure $\rho_{I,M,\eta}^1$ be defined by*

$$\frac{d\rho_{I,M,\eta}^1(\theta)}{d\mu_I} \propto \mathbf{1}_{\|\theta - \theta_{I,M}^*\|_1 \leq \eta}.$$

Then

$$\mathcal{K}(\rho_{I,M,\eta}^1, \mu_I) \leq (|I| - 1) \log \left(\max \left[|I|, \frac{4}{\eta} \right] \right).$$

Proof. For simplicity, we assume that $I = \{1, \dots, |I|\}$. Up to a permutation of the coordinates, the proof remains valid for any subset I of $\{1, \dots, p\}$. Still for simplicity, we let $\bar{\theta}$ denote $\theta_{I,M}^*$. By a symmetry argument, it can

be assumed that $\tilde{\theta}$ has nonnegative coordinates—this just means that $\tilde{\theta}$ is arbitrarily fixed in one of the $2^{|I|-1}$ faces of $\mathcal{S}_{1,+}^p(I)$. We denote by $\mathcal{FA}$ this face and note that

$$\mathcal{FA} = \left\{ \theta \in (\mathbb{R}_+)^{|I|} \times \{0\}^{p-|I|} : \sum_{j=1}^{|I|} \theta_j = 1 \right\}.$$

Finally, without loss of generality, we suppose that the largest coordinate in $\tilde{\theta}$ is $\tilde{\theta}_1$, and let ν be the uniform probability measure on $\mathcal{FA}$, defined by

$$\frac{d\nu}{d\mu_I}(\theta) = 2^{|I|-1} \mathbf{1}_{[\theta \in \mathcal{FA}]}.$$

Set $u = \min(1/|I|, \eta/2)$, and let

$$\begin{aligned} T_2 &= (\tilde{\theta}_1 - u, \tilde{\theta}_2 + u, \tilde{\theta}_3, \dots, \tilde{\theta}_{|I|}, 0, \dots, 0), \\ T_3 &= (\tilde{\theta}_1 - u, \tilde{\theta}_2, \tilde{\theta}_3 + u, \dots, \tilde{\theta}_{|I|}, 0, \dots, 0), \\ &\vdots \\ T_{|I|} &= (\tilde{\theta}_1 - u, \tilde{\theta}_2, \tilde{\theta}_3, \dots, \tilde{\theta}_{|I|} + u, 0, \dots, 0). \end{aligned}$$

Note that $u \leq 1/|I| \leq \tilde{\theta}_1$. Therefore, for each j , all the coordinates of T_j are nonnegative. Obviously $\|T_j\|_1 = 1$, so that, for all j , $T_j \in \mathcal{FA}$. Denoting by K the convex hull of the set $\{\tilde{\theta}, T_2, \dots, T_{|I|}\}$, we also have $K \subset \mathcal{FA}$. Next, observe that $\|T_j - \tilde{\theta}\|_1 = 2u \leq \eta$, which implies $K \subset \{\theta \in \mathbb{R}^p : \|\theta - \tilde{\theta}\|_1 \leq \eta\}$.

Clearly,

$$\begin{aligned} \mathcal{K}(\rho_{I,M,\eta}^1, \mu_I) &= \log \left(\frac{1}{\int \mathbf{1}_{[\|\theta - \theta_{I,M}^*\|_1 \leq \eta]} d\mu_I(\theta)} \right) \\ &\leq \log \left(\frac{1}{\int \mathbf{1}_{[\theta \in \mathcal{FA}]} \mathbf{1}_{[\|\theta - \theta_{I,M}^*\|_1 \leq \eta]} d\mu_I(\theta)} \right). \end{aligned}$$

Thus,

$$\begin{aligned}\mathcal{K}(\rho_{I,M,\eta}^1, \mu_I) &\leq \log \left(\frac{2^{|I|-1}}{\int \mathbf{1}_{[\|\theta - \theta_{I,M}^*\|_1 \leq \eta]} d\nu(\theta)} \right) \\ &\leq \log \left(\frac{2^{|I|-1}}{\int \mathbf{1}_{[\theta \in K]} d\nu(\theta)} \right).\end{aligned}$$

Observe that K is homothetic to $\mathcal{FA}$, by a factor of u . This means that

$$\int \mathbf{1}_{[\theta \in K]} d\nu(\theta) = u^{|I|-1}.$$

Consequently, we obtain

$$\mathcal{K}(\rho_{I,M,\eta}^1, \mu_I) \leq \log \left(\left(\frac{2}{u} \right)^{|I|-1} \right) \leq (|I| - 1) \log \left(\max \left[|I|, \frac{4}{\eta} \right] \right).$$

■

Lemma 4.6 *For any subset I of $\{1, \dots, p\}$, any positive integer $M \leq n$ and any $\gamma \in]0, 1/n]$, let the probability measure $\rho_{I,M,\gamma}^2$ be defined by*

$$\frac{d\rho_{I,M,\gamma}^2}{d\nu_M}(f) \propto \mathbf{1}_{[\|f - f_{I,M}^*\|_M \leq \gamma]}$$

where, for $f = \sum_{j=1}^M \beta_j \varphi_j \in \mathcal{F}_M(C+1)$, we put

$$\|f\|_M = \sum_{j=1}^M j |\beta_j|.$$

Then

$$\mathcal{K}(\rho_{I,M,\gamma}^2, \nu_M) = M \log \left(\frac{C+1}{\gamma} \right).$$

Proof. Observe that

$$\mathcal{K}(\rho_{I,M,\gamma}^2, \nu_M) = \int \log \left(\frac{d\rho_{I,M,\gamma}^2}{d\nu_M}(f) \right) d\rho_{I,M,\gamma}^2(f).$$

Now,

$$\frac{d\rho_{I,M,\gamma}^2}{d\nu_M}(f) = \frac{\mathbf{1}_{[\|f-f_{I,M}^*\|_M \leq \gamma]}(f)}{\zeta},$$

where $\zeta = \int \mathbf{1}_{[\|f-f_{I,M}^*\|_M \leq \gamma]}(f) d\nu_M(f)$. It easily follows, using the fact that the support of $\rho_{I,M,\gamma}^2$ is included in the set $\{f \in \mathcal{F}_M(C+1) : \|f-f_{I,M}^*\| \leq \gamma\}$, that

$$\mathcal{K}(\rho_{I,M,\gamma}^2, \nu_M) = \log(1/\zeta).$$

Note that

$$\begin{aligned} \zeta &= \int \mathbf{1}_{[\|f-f_{I,M}^*\|_M \leq \gamma]}(f) d\nu_M(f) \\ &= \frac{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j - (\beta_{I,M}^*)_j| \leq \gamma]}(\beta) \mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]}(\beta) d\beta}{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]}(\beta) d\beta}, \end{aligned}$$

where the second equality is true since ν_M is (the image of) the uniform probability measure on $\{\beta \in \mathbb{R}^M : \sum_{j=1}^M j|\beta_j| \leq C+1\}$. This implies

$$\mathcal{K}(\rho_{I,M,\gamma}^2, \nu_M) = \log \left(\frac{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]}(\beta) d\beta}{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j - (\beta_{I,M}^*)_j| \leq \gamma]}(\beta) \mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]}(\beta) d\beta} \right).$$

By the triangle inequality,

$$\sum_{j=1}^M j|\beta_j| \leq \sum_{j=1}^M j|\beta_j - (\beta_{I,M}^*)_j| + \sum_{j=1}^M j|(\beta_{I,M}^*)_j|.$$

Since $f_{I,M}^* \in \mathcal{F}_M(C)$, we have $\sum_{j=1}^M j|(\beta_{I,M}^*)_j| \leq C$, so that

$$\mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]} \geq \mathbf{1}_{[\sum_{j=1}^M j|\beta_j - (\beta_{I,M}^*)_j| \leq \gamma]}$$

as soon as $\gamma \leq 1$. We conclude that

$$\mathcal{K}(\rho_{I,M,\gamma}^2, \nu_M) = \log \left(\frac{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j| \leq C+1]} d\beta}{\int \mathbf{1}_{[\sum_{j=1}^M j|\beta_j - (\beta_{I,M}^*)_j| \leq \gamma]} d\beta} \right) = M \log \left(\frac{C+1}{\gamma} \right).$$

■

5 Annex: Description of the MCMC algorithm

This annex is intended to make thoroughly clear the specification of the proposal conditional densities k_1 and k_2 introduced in Section 3.

5.1 Notation

To provide explicit formulas for the conditional densities $k_1((\tau, h)|(\theta, f))$ and $k_2((\tau, h)|(\theta, f))$, we first set

$$f = \sum_{j=1}^{m_f} \beta_{f,j} \varphi_j \quad \text{and} \quad h = \sum_{j=1}^{m_h} \beta_{h,j} \varphi_j,$$

where it is recalled that $\{\varphi_j\}_{j=1}^{\infty}$ denotes the (non-normalized) trigonometric system. We let I (respectively, J) be the set of nonzero coordinates of the vector θ (respectively, τ), and denote finally by θ_I (respectively, τ_J) the vector of dimension $|I|$ (respectively, $|J|$) which contains the nonzero coordinates of θ (respectively, τ). Recall that all densities are defined with respect to the prior π , which is made explicit in Subsection 2.2.

For a generic $h \in \mathcal{F}_{m_h}(C+1)$, given $\tau \in \mathcal{S}_{1,+}^p$ and $s > 0$, we let the density $\text{dens}_s(h|\tau, m_h)$ with respect to π be defined by

$$\begin{aligned} & \text{dens}_s(h|\tau, m_h) \\ & \propto \exp \left[-\frac{1}{2s^2} \sum_{j=1}^{m_h} \left(\beta_{h,j} - \tilde{\beta}_j(\tau, m_h) \right)^2 \right] \mathbf{1} \left[\sum_{j=1}^{m_h} j |\beta_{h,j}| \leq C+1 \right], \end{aligned}$$

where the $\tilde{\beta}_j(\tau, m_h)$ are the empirical least square coefficients given by

$$\left\{ \tilde{\beta}_j(\tau, m_h) \right\}_{j=1, \dots, m_h} \in \arg \min_{b \in \mathbb{R}^{m_h}} \sum_{i=1}^n \left(Y_i - \sum_{j=1}^{m_h} b_j \varphi_j(\tau^T \mathbf{X}_i) \right)^2.$$

In the experiments, we fixed $s = 0.1$. Note that simulating with respect to $\text{dens}_s(h|\tau, m_h)$ is an easy task, since one just needs to compute a least square estimate and then draw from a truncated Gaussian distribution.

5.2 Description of k_1

We take

$$\begin{aligned} k_1(\cdot | (\theta, f)) &= \frac{2k_{1,=}(\cdot | (\theta, f)) + k_{1,+}(\cdot | (\theta, f))}{3} \mathbf{1}_{[|I|=1]} \\ &+ \frac{k_{1,-}(\cdot | (\theta, f)) + 2k_{1,=}(\cdot | (\theta, f)) + k_{1,+}(\cdot | (\theta, f))}{4} \mathbf{1}_{[1 < |I| < p]} \\ &+ \frac{k_{1,-}(\cdot | (\theta, f)) + 2k_{1,=}(\cdot | (\theta, f))}{3} \mathbf{1}_{[|I|=p]}. \end{aligned}$$

Roughly, the idea is that $k_{1,-}$ tries to remove one component in θ , $k_{1,=}$ keeps the same number of components, whereas $k_{1,+}$ adds one component. The density $k_{1,=}$ takes the form

$$k_{1,=}((\tau, h) | (\theta, f)) = k_{1,=}(\tau | \theta) \text{dens}_s(h | \tau, m_f).$$

The density $k_{1,=}(\cdot | \theta)$ is the density of τ when $J = I$ and

$$\tau_I = \frac{\theta_I + E}{\|\theta_I + E\|_1} \text{sgn}((\theta_I + E)_{j(\theta_I + E)}),$$

where $E = (E_1, \dots, E_{|I|})$ and the E_i are independent random variables uniformly distributed in $[-\delta, \delta]$. Throughout, the value of δ was fixed at 0.5. It is noteworthy that when we change the parameter from θ to τ , then we also change the function from f to h . Thus, with this procedure, the link function h is more “adapted” to τ and the subsequent move is more likely to be accepted in the Hastings-Metropolis algorithm.

In the case where we are to remove one component, $k_{1,-}$ is given by

$$k_{1,-}((\tau, h) | (\theta, f)) = \sum_{j \in I} c_j \mathbf{1}_{[\tau = \theta_{-j}]} \text{dens}_s(h | \tau, m_f),$$

where θ_{-j} is just obtained from θ by setting the j -th component to 0 and by renormalizing the parameter in order to have $\|\theta_{-j}\|_1 = 1$. We set

$$c_j = \frac{\exp(-|\theta_j|) \mathbf{1}_{[|\theta_j| < \delta]}}{\sum_{\ell \in I} \exp(-|\theta_\ell|) \mathbf{1}_{[|\theta_\ell| < \delta]}}.$$

The idea is that smaller components are more likely to be removed than larger ones. Finally, the density $k_{1,+}$ takes the form

$$k_{1,+}((\tau, h) | (\theta, f)) = \sum_{j \notin I} c'_j \mathbf{1}_{[\tau_{-j} = \theta]} \frac{\mathbf{1}_{[|\tau_j| < \delta]}}{2\delta} \text{dens}_s(h | \tau, m_f).$$

We set

$$c'_j = \frac{\exp(|\sum_{i=1}^n (Y_i - f(\theta^T \mathbf{X}_i)) (\mathbf{X}_i)_j|)}{\sum_{\ell \notin I} \exp(|\sum_{i=1}^n (Y_i - f(\theta^T \mathbf{X}_i)) (\mathbf{X}_i)_\ell|)}$$

where $(\mathbf{X}_i)_j$ denotes the j -th component of $\mathbf{X}_i$. In words, the idea is that a new nonzero coordinate in θ is more likely to be interesting in the model if the corresponding feature is correlated with the current residual.

5.3 Description of k_2

In the same spirit, we let the conditional density k_2 be defined by

$$\begin{aligned} k_2(\cdot | (\theta, f)) &= \frac{2k_{2,=}(\cdot | (\theta, f)) + k_{2,+}(\cdot | (\theta, f))}{3} \mathbf{1}_{[m_f=1]} \\ &+ \frac{k_{2,-}(\cdot | (\theta, f)) + 2k_{2,=}(\cdot | (\theta, f)) + k_{2,+}(\cdot | (\theta, f))}{4} \mathbf{1}_{[1 < m_f < n]} \\ &+ \frac{k_{2,-}(\cdot | (\theta, f)) + 2k_{2,=}(\cdot | (\theta, f))}{3} \mathbf{1}_{[m_f=n]}. \end{aligned}$$

We choose

$$k_{2,=}((\tau, h) | (\theta, f)) = \mathbf{1}_{[\tau=\theta]} \text{dens}_s(h | \tau, m_f)$$

and

$$k_{2,+}((\tau, h) | (\theta, f)) = \mathbf{1}_{[\tau=\theta]} \text{dens}_s(h | \tau, m_f + 1).$$

With this choice, $m_h = m_f + 1$, which means that the proposal density tries to add one coefficient in the expansion of h , while leaving θ unchanged. Finally

$$k_{2,-}((\tau, h) | (\theta, f)) = \mathbf{1}_{[\tau=\theta]} \text{dens}_s(h | \tau, m_f - 1),$$

and the proposal tries to remove one coefficient in h .

Acknowledgments. The authors thank three referees for valuable comments and insightful suggestions, which lead to a substantial improvement of the paper. They also thank John O’Quigley for his careful reading of the article.

References

- [1] P. Alquier. PAC-Bayesian bounds for randomized empirical risk minimizers. *Mathematical Methods of Statistics*, 17:279–304, 2008.
- [2] P. Alquier and K. Lounici. PAC-Bayesian bounds for sparse regression estimation with exponential weights. *Electronic Journal of Statistics*, 5:127–145, 2011.

- [3] A. Antoniadis, G. Grégoire, and I.W. McKeague. Bayesian estimation in single-index models. *Statistica Sinica*, 14:1147–1164, 2004.
- [4] J.-Y. Audibert. Aggregated estimators and empirical complexity for least square regression. *Annales de l’Institut Henri Poincaré: Probability and Statistics*, 40:685–736, 2004.
- [5] J.-Y. Audibert and O. Catoni. Robust linear least squares regression. *The Annals of Statistics*, in press, 2011.
- [6] R.E. Bellman. *Adaptive Control Processes: A Guided Tour*. Princeton University Press, 1961.
- [7] P.J. Bickel, Y. Ritov, and A.B. Tsybakov. Simultaneous analysis of Lasso and Dantzig selector. *The Annals of Statistics*, 37:1705–1732, 2009.
- [8] A.M. Bruckstein, D.L. Donoho, and M. Elad. From sparse solutions of systems of equations to sparse modeling of signals and images. *SIAM Review*, 51:34–81, 2009.
- [9] P. Bühlmann and S. van de Geer. *Statistics for High-Dimensional Data*. Springer, New York, 2011.
- [10] F. Bunea, A. Tsybakov, and M. Wegkamp. Sparsity oracle inequalities for the Lasso. *Electronic Journal of Statistics*, 1:169–194, 2007.
- [11] E.J. Candès and T. Tao. The Dantzig selector: Statistical estimation when p is much larger than n . *The Annals of Statistics*, 35:2313–2351, 2005.
- [12] O. Catoni. *Statistical Learning Theory and Stochastic Optimization*. Springer, 2004.
- [13] O. Catoni. *PAC-Bayesian Supervised Classification: The Thermodynamics of Statistical Learning*, volume 56 of *Lecture Notes-Monograph Series*. IMS, 2007.
- [14] J.M. Chambers, W.S. Cleveland, B. Kleiner, and P.A. Tukey. *Graphical Methods for Data Analysis*. Wadsworth & Brooks, Belmont, 1983.
- [15] X. Chen, C. Zou, and R.D. Cook. Coordinate-independent sparse sufficient dimension reduction and variable selection. *The Annals of Statistics*, 38:3696–3723, 2010.

- [16] A. Cohen, I. Daubechies, R. DeVore, G. Kerkyacharian, and D. Picard. Capturing ridge functions in high dimension from point queries. *Constructive Approximation*, in press, 2011.
- [17] P. Cortez, A. Cerdeira, F. Almeida, T. Matos, and J. Reis. Modeling wine preferences by data mining from physicochemical properties. *Decision Support Systems*, 47:547–553, 2009.
- [18] A.S. Dalalyan, A. Juditsky, and V. Spokoiny. A new algorithm for estimating the effective dimension-reduction subspace. *Journal of Machine Learning Research*, 9:1647–1678, 2008.
- [19] A.S. Dalalyan and A.B. Tsybakov. Aggregation by exponential weighting, sharp PAC-Bayesian bounds and sparsity. *Machine Learning*, 72:39–61, 2008.
- [20] A.S. Dalalyan and A.B. Tsybakov. Sparse regression learning by aggregation and Langevin Monte-Carlo. *Journal of Computer and System Sciences*, in press, 2011.
- [21] M. Delecroix, M. Hristache, and V. Patilea. On semiparametric M -estimation in single-index regression. *Journal of Statistical Planning and Inference*, 136:730–769, 2006.
- [22] S. Gaïffas and G. Lecué. Optimal rates and adaptation in the single-index model using aggregation. *Electronic Journal of Statistics*, 1:538–573, 2007.
- [23] P.J. Green. Reversible jump Markov chain Monte Carlo computation and Bayesian model determination. *Biometrika*, 82:711–732, 1995.
- [24] L. Györfi, M. Kohler, A. Krzyżak, and H. Walk. *A Distribution-Free Theory of Nonparametric Regression*. Springer, New York, 2002.
- [25] W. Härdle, P. Hall, and H. Ichimura. Optimal smoothing in single-index models. *The Annals of Statistics*, 21:157–178, 1993.
- [26] D. Jr. Harrison and D.L. Rubinfeld. Hedonic housing prices and the demand for clean air. *Journal of Environmental Economics and Management*, 5:81–102, 1978.
- [27] J.L. Horowitz. *Semiparametric Methods in Econometrics*. Springer, 1998.

- [28] H. Ichimura. Semiparametric least squares (SLS) and weighted SLS estimation of single-index models. *Journal of Econometrics*, 58:71–120, 1993.
- [29] O. Lopez. Single-index regression models with right-censored responses. *Journal of Statistical Planning and Inference*, 139:1082–1097, 2009.
- [30] J.-M. Marin and C. Robert. *Bayesian Core: A Practical Approach to Computational Bayesian Analysis*. Springer, New York, 2007.
- [31] P. Massart. *Concentration Inequalities and Model Selection*. Springer, Berlin, 2007.
- [32] D.A. McAllester. Some PAC-Bayesian theorems. In *Proceedings of the Eleventh Annual Conference on Computational Learning Theory*, pages 230–234, New York, 1998. ACM.
- [33] P. McCullagh and J.A. Nelder. *Generalized Linear Models*. Chapman and Hall, 1983.
- [34] E.A. Nadaraya. On estimating regression. *Theory of Probability and its Applications*, 9:141–142, 1964.
- [35] E.A. Nadaraya. Remarks on nonparametric estimates for density functions and regression curves. *Theory of Probability and its Applications*, 15:134–137, 1970.
- [36] J.R. Quinlan. Combining instance-based and model-based learning. In *Proceedings of the Tenth International Conference on Machine Learning*, pages 236–243, Amherst, 1993. Morgan Kaufmann.
- [37] R Development Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, 2008.
- [38] Y. Seldin, N. Cesa-Bianchi, F. Laviolette, P. Auer, J. Shawe-Taylor, and J. Peters. *PAC-Bayesian analysis of the exploration-exploitation trade-off*. arXiv:1105.4585, 2011.
- [39] J. Shawe-Taylor and R. Williamson. A PAC analysis of a Bayes estimator. In *Proceedings of the Tenth Annual Conference on Computational Learning Theory*, pages 2–9, New York, 1997. ACM.
- [40] C.J. Stone. Optimal global rates of convergence for nonparametric regression. *The Annals of Statistics*, 10:1040–1053, 1982.

- [41] R. Tibshirani. Regression shrinkage and selection via the Lasso. *Journal of the Royal Statistical Society, Series B*, 58:267–288, 1996.
- [42] A.B. Tsybakov. *Introduction to Nonparametric Estimation*. Springer, 2009.
- [43] L.J. van’t Veer, H. Dai, M.J. van de Vijver, Y.D. He, A.A.M. Hart, M. Mao, H.L. Peterse, K. van der Kooy, M.J. Marton, A.T. Witteveen, G.J. Schreiber, R.M. Kerkhoven, C. Roberts, P.S. Linsley, R. Bernards, and S.H. Friend. Gene expression profiling predicts clinical outcome of breast cancer. *Nature*, 415:530–536, 2002.
- [44] H.B. Wang. Bayesian estimation and variable selection for single index models. *Computational Statistics and Data Analysis*, 53:2617–2627, 2009.
- [45] G.S. Watson. Smooth regression analysis. *Sankhyā Series A*, 26:359–372, 1964.
- [46] I.-C. Yeh. Modeling of strength of high-performance concrete using artificial neural networks. *Cement and Concrete Research*, 28:1797–1808, 1998.
- [47] I.-C. Yeh. Modeling slump flow of concrete using second-order regressions and artificial neural networks. *Cement and Concrete Composites*, 29:474–480, 2007.