



Density estimation with quadratic loss: a confidence intervals method

Pierre Alquier

► To cite this version:

Pierre Alquier. Density estimation with quadratic loss: a confidence intervals method. 2006. hal-00020740

HAL Id: hal-00020740

<https://hal.science/hal-00020740>

Preprint submitted on 14 Mar 2006

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.

DENSITY ESTIMATION WITH QUADRATIC LOSS: A CONFIDENCE INTERVALS METHOD

PIERRE ALQUIER

ABSTRACT. In [1], a least square regression estimation procedure was proposed: first, we consider a family of functions f_k and study the properties of an estimator in every unidimensionnal model $\{\alpha f_k, \alpha \in \mathbb{R}\}$; we then show how to aggregate these estimators. The purpose of this paper is to extend this method to the case of density estimation. We first give a general overview of the method, adapted to the density estimation problem. We then show that this leads to adaptative estimators, that means that the estimator reaches the best possible rate of convergence (up to a log factor). Finally we show some ways to improve and generalize the method.

1. INTRODUCTION: THE DENSITY ESTIMATION SETTING

Let us assume that we are given a measure space $(\mathcal{X}, \mathcal{B}, \lambda)$ where λ is positive and σ -finite, and a probability measure P on $(\mathcal{X}, \mathcal{B})$ such that P has a density with respect to λ :

$$P(dx) = f(x)\lambda(dx).$$

We assume that we observe a realisation of the canonical process $(X_1, \dots, X_N)$ on $(\mathcal{X}^N, \mathcal{B}^{\otimes N}, P^{\otimes N})$. Our objective here is to estimate f on the basis of the observations $X_1, \dots, X_N$.

More precisely, let $\mathcal{L}^2(\mathcal{X}, \lambda)$ denote the set of all measurable functions from $(\mathcal{X}, \mathcal{B})$ to $(\mathbb{R}, \mathcal{B}_R)$ where $\mathcal{B}_R$ is the Borel σ -algebra on $\mathbb{R}$. We will write $\mathcal{L}^2(\mathcal{X}, \lambda) = \mathcal{L}^2$ for short. Remark that $f \in \mathcal{L}^2$. Let us put, for any $(g, h) \in (\mathcal{L}^2)^2$:

$$d^2(g, h) = \int_{\mathcal{X}} (g(x) - h(x))^2 \lambda(dx),$$

and let $\|\cdot\|$ and $\langle \cdot, \cdot \rangle$ denote the corresponding norm and scalar product. We are here looking for an estimator $\hat{f}$ that tries to minimize our objective:

$$d^2(\hat{f}, f).$$

Let us choose an integer $m \in \mathbb{N}$ and a family of functions $(f_1, \dots, f_m) \in (\mathcal{L}^2)^m$. There is no particular assumptions about this family: it is not necessarily linearly independant for example.

In a first time, we are going to study estimators of f in every unidimensionnal model $\{\alpha f_k(\cdot), \alpha \in \mathbb{R}\}$ (as done in [1]). Usually these models are too small and the obtained estimators do not have good properties. We then propose an iterative method that selects and aggregate such estimators in order to build a suitable estimator of f (section 2).

Date: March 14, 2006.

2000 Mathematics Subject Classification. Primary 62G07; Secondary 62G15, 68T05.

Key words and phrases. Density estimation, statistical learning, confidence regions, thresholding methods, support vector machines.

I would like to thank my PhD advisor, Professor Olivier Catoni, for his kind and constant help, and Professors Patrice Bertail, Emmanuelle Gautherat and Hugo Harari-Kermadec for their remark that Panchenko's lemma could improve theorem 2.1.

In section 3 we study the rate of convergence of the obtained estimator in a particular case.

In section 4 we propose several improvements and generalizations of the method.

Finally, in section 5 we make some simulations in order to compare the practical performances of our estimator with other ones.

2. ESTIMATION METHOD

2.1. Hypothesis. In this section we will use a particular hypothesis about f and/or the basis functions $f_k, k \in \{1, \dots, m\}$.

Definition 2.1. We will say that f and $(f_1, \dots, f_m)$ satisfies the conditions $\mathcal{H}(p)$ for $1 < p < +\infty$ if, for:

$$\frac{1}{p} + \frac{1}{q} = 1,$$

there exists some $(c, c_1, \dots, c_m) \in (\mathbb{R}_+^*)^{m+1}$ (known to the statistician) such that:

$$\forall k \in \{1, \dots, m\}, \quad \left(\int_{\mathcal{X}} |f_k|^{2p} \lambda(dx) \right)^{\frac{1}{p}} \leq c_k \int_{\mathcal{X}} |f_k|^2 \lambda(dx)$$

$$\text{and} \quad \left(\int_{\mathcal{X}} |f|^q \lambda(dx) \right)^{\frac{1}{p}} \leq c \int_{\mathcal{X}} |f| \lambda(dx) \quad (= c).$$

For $p = 1$ the condition $\mathcal{H}(1)$ is: f is bounded by a (known) constant c and we put $c_1 = \dots = c_m = 1$. For $p = +\infty$ the condition $\mathcal{H}(+\infty)$ is just that every $|f_k|$ is bounded by

$$\sqrt{c_k \int_{\mathcal{X}} f_k(x)^2 \lambda(dx)}$$

where c_k is known, and we put $c = 1$. In any case, we put, for any k :

$$C_k = c_k c.$$

Definition 2.2. We put, for any $k \in \{1, \dots, m\}$:

$$D_k = \int_{\mathcal{X}} |f_k|^2 \lambda(dx) = d^2(f_k, 0) = \|f_k\|^2.$$

2.2. Unidimensionnal models. Let us choose $k \in \{1, \dots, m\}$ and consider the unidimensionnal model $\mathcal{M}_k = \{\alpha f_k(\cdot), \alpha \in \mathbb{R}\}$. Remark that the orthogonal projection (denoted by $\Pi_{\mathcal{M}_k}$) of f on $\mathcal{M}_k$ is known, it is namely:

$$\Pi_{\mathcal{M}_k} f(\cdot) = \bar{\alpha}_k f_k(\cdot)$$

where:

$$\bar{\alpha}_k = \arg \min_{\alpha \in \mathbb{R}} d^2(\alpha f_k, f) = \frac{\int_{\mathcal{X}} f_k(x) f(x) \lambda(dx)}{\int_{\mathcal{X}} f_k(x)^2 \lambda(dx)} = \frac{\int_{\mathcal{X}} f_k(x) f(x) \lambda(dx)}{D_k}.$$

A natural estimator of this coefficient is:

$$\hat{\alpha}_k = \frac{\frac{1}{N} \sum_{i=1}^N f_k(X_i)}{\int_{\mathcal{X}} f_k(x)^2 \lambda(dx)},$$

because we expect to have, by the law of large numbers:

$$\frac{1}{N} \sum_{i=1}^N f_k(X_i) \xrightarrow[N \rightarrow \infty]{a.s.} P[f_k(X)] = \int_{\mathcal{X}} f_k(x) f(x) \lambda(dx).$$

Actually, we can formulate a more precise result.

Theorem 2.1. *Let us assume that condition $\mathcal{H}(p)$ holds for some $p \in [1, +\infty]$. Then for any $\varepsilon > 0$ we have:*

$$P^{\otimes N} \left\{ \forall k \in \{1, \dots, m\}, d^2(\hat{\alpha}_k f_k, \bar{\alpha}_k f_k) \leq \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} \right\} \left[\frac{\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2}{D_k} + C_k \right] \right\} \geq 1 - \varepsilon.$$

The proof is given at the end of the section.

2.3. The selection algorithm. Until the end of this section we assume that $\mathcal{H}(p)$ is satisfied for some $1 \leq p \leq +\infty$.

Let $\beta(\varepsilon, k)$ denote the upper bound for the model k in theorem 2.1:

$$\forall \varepsilon > 0, \forall k \in \{1, \dots, m\} : \quad \beta(\varepsilon, k) = \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} \right\} \left[\frac{\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2}{D_k} + C_k \right].$$

Let us put:

$$\mathcal{CR}_{k,\varepsilon} = \left\{ g \in \mathcal{L}^2, d^2(\hat{\alpha}_k f_k, \Pi_{\mathcal{M}_k} g) \leq \beta(\varepsilon, k) \right\}.$$

Then theorem 2.1 implies the following result.

Corollary 2.2. *For any $\varepsilon > 0$ we have:*

$$P^{\otimes N} \left\{ \forall k \in \{1, \dots, m\}, f \in \mathcal{CR}_{k,\varepsilon} \right\} \geq 1 - \varepsilon.$$

So for any k , $\mathcal{CR}_{k,\varepsilon}$ is a confidence region at level k for f . Moreover, $\mathcal{CR}_{k,\varepsilon}$ being convex we have the following corollary.

Corollary 2.3. *For any $\varepsilon > 0$ we have:*

$$P^{\otimes N} \left\{ \forall k \in \{1, \dots, m\}, \forall g \in \mathcal{L}^2, d^2(\Pi_{\mathcal{CR}_{k,\varepsilon}} g, f) \leq d^2(g, f) \right\} \geq 1 - \varepsilon.$$

It just means that for any g , $\Pi_{\mathcal{M}_k} g$ is a better estimator than g .

So we propose the following algorithm (generic form):

- we choose ε and start with $g_0 = 0$;
- at each step n , we choose a model $\mathcal{M}_{k(n)}$ where $k(n) \in \{1, \dots, m\}$ can be chosen on each way we want (it can of course depend on the data) and take:

$$g_{n+1} = \Pi_{\mathcal{CR}_{k(n),\varepsilon}} g_n;$$

- we choose a stopping time n_s on each way we want and take:

$$\hat{f} = g_{n_s}.$$

So corollary 2.3 implies that:

$$P^{\otimes N} \left\{ d^2(\hat{f}, f) = d^2(g_{n_s}, f) \leq \dots \leq d^2(g_0, f) = d^2(0, f) \right\} \geq 1 - \varepsilon.$$

Actually, a more accurate version of corollary 2.3 can give an idea of the way to choose $k(n)$ in the algorithm. Let us use corollary 2.2 and remember the fact that each $\mathcal{CR}_{k,\varepsilon}$ is convex.

Corollary 2.4. *For any $\varepsilon > 0$ we have:*

$$P^{\otimes N} \left\{ \forall k \in \{1, \dots, m\}, \forall g \in \mathcal{L}^2, d^2(\Pi_{\mathcal{CR}_{k,\varepsilon}} g, f) \leq d^2(g, f) - d^2(\Pi_{\mathcal{CR}_{k,\varepsilon}} g, g) \right\} \geq 1 - \varepsilon.$$

So we propose the following version of our previous algorithm (this is not necessarily the better choice!):

- we choose ε and $0 < \kappa \leq 1/N$ and start with $g_0 = 0$;
- at each step n , we take:

$$k(n) = \arg \max_{k \in \{1, \dots, m\}} d^2(\Pi_{\mathcal{RC}_{k,\varepsilon}} g_n, g_n)$$

and:

$$g_{n+1} = \Pi_{\mathcal{RC}_{k(n),\varepsilon}} g_n;$$

- we take:

$$n_s = \inf \{n \in \mathbb{N} : d^2(g_n, g_{n-1}) \leq \kappa\}$$

and:

$$\hat{f} = g_{n_s}.$$

So corollary 2.4 implies that:

$$P^{\otimes N} \left\{ d^2(\hat{f}, f) \leq d^2(0, f) - \sum_{n=0}^{n_s-1} d^2(g_n, g_{n+1}) \right\} \geq 1 - \varepsilon.$$

2.4. Remarks on the intersection of the confidence regions. Actually, corollary 2.2 could motivate another method. Note that:

$$\forall k \in \{1, \dots, m\}, f \in \mathcal{RC}_{k,\varepsilon} \Leftrightarrow f \in \bigcap_{k=1}^m \mathcal{RC}_{k,\varepsilon}.$$

Let us put, for any $I \subset \{1, \dots, m\}$:

$$\mathcal{RC}_{I,\varepsilon} = \bigcap_{k \in I} \mathcal{RC}_{k,\varepsilon},$$

and:

$$\hat{f}_I = \Pi_{\mathcal{RC}_{I,\varepsilon}} 0.$$

Then $\mathcal{RC}_{I,\varepsilon}$ is still a convex region that contains f and is a subset of every $\mathcal{RC}_{k,\varepsilon}$ for $k \in I$. So we have the following result.

Corollary 2.5. *For any $\varepsilon > 0$:*

$$P^{\otimes N} \left\{ \forall I \subset \{1, \dots, m\}, \forall k \in I, d(\hat{f}_{\{1, \dots, m\}}, f) \leq d(\hat{f}_I, f) \leq d(\Pi_{\mathcal{RC}_{k,\varepsilon}} 0, f) \right\} \geq 1 - \varepsilon.$$

In the case where we are interested in "model selection type aggregation" of estimators, note that, with probability at least $1 - \varepsilon$:

$$d(\Pi_{\mathcal{RC}_{k,\varepsilon}} 0, f) \leq d(\Pi_{\mathcal{RC}_{k,\varepsilon}} \bar{\alpha}_k f_k) + d(\bar{\alpha}_k f_k, f) \leq \beta(\varepsilon, k) + d(f_k, f).$$

So we have the following result.

Corollary 2.6. *For any $\varepsilon > 0$:*

$$P^{\otimes N} \left\{ d(\hat{f}_{\{1, \dots, m\}}, f) \leq \inf_{k \in \{1, \dots, m\}} [d(f_k, f) + \beta(\varepsilon, k)] \right\} \geq 1 - \varepsilon.$$

The estimator $\hat{f}_{1, \dots, m}$ can be reached by solving the following optimization problem:

$$\begin{aligned} & \min_{g \in \mathcal{L}^2} \|g\|^2, \\ & s.t. \quad \forall k \in \{1, \dots, m\} : \\ & \quad \begin{cases} \langle g - \hat{\alpha}_k f_k, f_k \rangle - \sqrt{D_k \beta(\varepsilon, k)} \leq 0, \\ -\langle g - \hat{\alpha}_k f_k, f_k \rangle - \sqrt{D_k \beta(\varepsilon, k)} \leq 0. \end{cases} \end{aligned}$$

The problem can be solved in dual form:

$$\max_{\gamma \in \mathbb{R}^m} \left[- \sum_{i=1}^m \sum_{k=1}^m \gamma_i \gamma_k \langle f_i, f_k \rangle + 2 \sum_{k=1}^m \gamma_k \hat{\alpha}_k \|f_k\|^2 - 2 \sum_{k=1}^m |\gamma_k| \sqrt{D_k \beta(\varepsilon, k)} \right].$$

with solution $\gamma^* = (\gamma_1^*, \dots, \gamma_m^*)$ and:

$$\hat{f}_{\{1, \dots, m\}} = \sum_{k=1}^m \gamma_k^* f_k.$$

As:

$$- \sum_{i=1}^m \sum_{k=1}^m \gamma_i^* \gamma_k^* \langle f_i, f_k \rangle = \|f^*\|^2$$

and:

$$2 \sum_{k=1}^m \gamma_k^* \hat{\alpha}_k \|f_k\|^2 = 2 \sum_{k=1}^m \gamma_k^* \frac{\frac{1}{N} \sum_{i=1}^N f_k(X_i)}{\|f_k\|^2} \|f_k\|^2 = \frac{2}{N} \sum_{i=1}^N f^*(X_i)$$

we can see this as a penalized maximization of the likelihood.

We can note that it is easier and more computationnaly efficient to project successively on every region $\mathcal{RC}(k, \varepsilon)$ than to project once on $\mathcal{RC}(\{1, \dots, m\}, \varepsilon)$.

2.5. An example: the histogram. Let us assume that λ is a finite measure and let $A_1, \dots, A_m$ be a partition of $\mathcal{X}$. We put, for any $k \in \{1, \dots, m\}$:

$$f_k(\cdot) = \mathbb{1}_{A_k}(\cdot).$$

Remark that:

$$D_k = \int_{\mathcal{X}} f_k(x)^2 \lambda(dx) = \lambda(A_k),$$

and that condition $\mathcal{H}(+\infty)$ is satisfied with constants:

$$c_k = \frac{1}{\lambda(A_k)}$$

and (as we have the convention $c = 1$ in this case) $C_k = c_k c = c_k$.

In this context we have:

$$\begin{aligned} \bar{\alpha}_k &= \frac{P(X \in A_k)}{\lambda(A_k)}, \\ \hat{\alpha}_k &= \frac{\frac{1}{N} \sum_{i=1}^N \mathbb{1}_{A_k}(X_i)}{\lambda(A_k)}, \\ \beta(\varepsilon, k) &= \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N \lambda(A_k)} \right\} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + 1 \right]. \end{aligned}$$

Finally, note that all the confidence regions $\mathcal{CR}_{k, \varepsilon}$ are all orthogonal in this case. So the order of projection does not affect the obtained estimator here, and we can take:

$$\hat{f} = \Pi_{\mathcal{CR}_{m, \varepsilon}} \dots \Pi_{\mathcal{CR}_{1, \varepsilon}} 0$$

(and note that $\hat{f} = \hat{f}_{\{1, \dots, m\}}$ here, following the notations of subsection 2.4). We have:

$$\hat{f}(x) = \sum_{k=1}^m \left(\hat{\alpha}_k - \sqrt{\lambda(A_k) \beta(\varepsilon, k)} \right)_+ f_k(x)$$

where, for any $y \in \mathbb{R}$:

$$(y)_+ = \begin{cases} y & \text{if } y \geq 0 \\ 0 & \text{otherwise.} \end{cases}$$

In this case corollary 2.4 becomes:

$$P^{\otimes N} \left\{ d^2(\hat{f}, f) \leq d^2(0, f) - \sum_{k=1}^m \left(\hat{\alpha}_k - \sqrt{\lambda(A_k) \beta(\varepsilon, k)} \right)_+^2 \lambda(A_k) \right\} \geq 1 - \varepsilon.$$

2.6. Proof of the theorem. Before giving the proof, let us state two lemmas that we will use in the proof. The first one is a variant of a lemma by Catoni [6], the second one is due to Panchenko [11].

Lemma 2.7. *Let $(T_1, \dots, T_{2N})$ be a random vector taking values in $\mathbb{R}^{2N}$ distributed according to a distribution $\mathcal{P}^{\otimes 2N}$. For any $\eta \in \mathbb{R}$, for any measurable function $\lambda : \mathbb{R}^{2N} \rightarrow \mathbb{R}_+^*$ that is exchangeable with respect to its $2 \times 2N$ arguments:*

$$\mathcal{P}^{\otimes 2N} \exp \left(\frac{\lambda}{N} \sum_{i=1}^N \{T_{i+N} - T_i\} - \frac{\lambda^2}{N^2} \sum_{i=1}^{2N} T_i^2 - \eta \right) \leq \exp(-\eta)$$

and the reverse inequality:

$$\mathcal{P}^{\otimes 2N} \exp \left(\frac{\lambda}{N} \sum_{i=1}^N \{T_i - T_{i+1}\} - \frac{\lambda^2}{N^2} \sum_{i=1}^{2N} T_i^2 - \eta \right) \leq \exp(-\eta),$$

where we write:

$$\begin{aligned} \eta &= \eta(T_1, \dots, T_{2N}) \\ \lambda &= \lambda(T_1, \dots, T_{2N}) \end{aligned}$$

for short.

Proof of lemma 2.7. In order to prove the first inequality, we write:

$$\begin{aligned} &\mathcal{P}^{\otimes 2N} \exp \left(\frac{\lambda}{N} \sum_{i=1}^N \{T_{i+N} - T_i\} - \frac{\lambda^2}{N^2} \sum_{i=1}^{2N} T_i^2 - \eta \right) \\ &= \mathcal{P}^{\otimes 2N} \exp \left(\sum_{i=1}^N \log \cosh \left\{ \frac{\lambda}{N} (T_{i+1} - T_i) \right\} - \frac{\lambda^2}{N^2} \sum_{i=1}^{2N} T_i^2 - \eta \right). \end{aligned}$$

We now use the inequality:

$$\forall x \in \mathbb{R}, \log \cosh x \leq \frac{x^2}{2}.$$

We obtain:

$$\log \cosh \left\{ \frac{\lambda}{N} (T_{i+1} - T_i) \right\} \leq \frac{\lambda^2}{2N^2} (T_{i+1} - T_i)^2 \leq \frac{\lambda^2}{N^2} (T_{i+1}^2 + T_i^2).$$

The proof for the reverse inequality is exactly the same. $\square$

Lemma 2.8 (Panchenko [11], corollary 1). *Let us assume that we have i.i.d. variables $T_1, \dots, T_N$ (with distribution $\mathcal{P}$ and values in $\mathbb{R}$) and an independant copy $T' = (T_{N+1}, \dots, T_{2N})$ of $T = (T_1, \dots, T_N)$. Let $\xi_j(T, T')$ for $j \in \{1, 2, 3\}$ be three measurables functions taking values in $\mathbb{R}$, and $\xi_3 \geq 0$. Let us assume that we know two constants $A \geq 1$ and $a > 0$ such that, for any $u > 0$:*

$$P^{\otimes 2N} \left[\xi_1(T, T') \geq \xi_2(T, T') + \sqrt{\xi_3(T, T')u} \right] \leq A \exp(-au).$$

Then, for any $u > 0$:

$$\begin{aligned} P^{\otimes 2N} \left\{ P^{\otimes 2N} [\xi_1(T, T') | T] \right. \\ \left. \geq P^{\otimes 2N} [\xi_2(T, T') | T] + \sqrt{P^{\otimes 2N} [\xi_3(T, T') | T] u} \right\} \leq A \exp(1 - au). \end{aligned}$$

The proof of this lemma can be found in [11]. We can now give the proof of theorem 2.1.

Proof of theorem 2.1. Let $(X_{N+1}, \dots, X_{2N})$ be an independant copy of our sample $(X_1, \dots, X_N)$. Let us choose $k \in \{1, \dots, m\}$. Let us apply lemma 2.7 with $\mathcal{P} = P$ and, for any $i \in \{1, \dots, 2N\}$:

$$T_i = f_k(X_i).$$

We obtain, for any measurable function $\eta_k \in \mathbb{R}$, for any measurable function $\lambda_k : \mathbb{R}^{2N} \rightarrow \mathbb{R}_+^*$ that is exchangeable with respect to its $2 \times 2N$ arguments:

$$P^{\otimes 2N} \exp \left(\frac{\lambda_k}{N} \sum_{i=1}^N \left\{ f_k(X_{i+N}) - f_k(X_i) \right\} - \frac{\lambda_k^2}{N^2} \sum_{i=1}^{2N} f_k(X_i)^2 - \eta_k \right) \leq \exp(-\eta_k)$$

and the reverse inequality:

$$P^{\otimes 2N} \exp \left(\frac{\lambda_k}{N} \sum_{i=1}^N \left\{ f_k(X_i) - f_k(X_{i+N}) \right\} - \frac{\lambda_k^2}{N^2} \sum_{i=1}^{2N} f_k(X_i)^2 - \eta_k \right) \leq \exp(-\eta_k)$$

as well. This implies that:

$$P^{\otimes 2N} \left[\frac{1}{N} \sum_{i=1}^N \left\{ f_k(X_i) - f_k(X_{i+N}) \right\} \leq \frac{\lambda_k}{N^2} \sum_{i=1}^{2N} f_k(X_i)^2 + \frac{\eta_k}{\lambda_k} \right] \leq \exp(-\eta_k)$$

and:

$$P^{\otimes 2N} \left[\frac{1}{N} \sum_{i=1}^N \left\{ f_k(X_{i+N}) - f_k(X_i) \right\} \leq \frac{\lambda_k}{N^2} \sum_{i=1}^{2N} f_k(X_i)^2 + \frac{\eta_k}{\lambda_k} \right] \leq \exp(-\eta_k).$$

Let us choose:

$$\lambda_k = \sqrt{\frac{N\eta_k}{\frac{1}{N} \sum_{i=1}^{2N} f_k(X_i)^2}}$$

in both inequalities, we obtain for the first one:

$$P^{\otimes 2N} \left[\frac{1}{N} \sum_{i=1}^N \left\{ f_k(X_i) - f_k(X_{i+N}) \right\} \geq 2\sqrt{\frac{\eta_k \frac{1}{N} \sum_{i=1}^{2N} f_k(X_i)^2}{N}} \right] \leq \exp(-\eta_k).$$

We now apply lemma 2.8 with the same $T_i = f_k(X_i)$, $\eta_k = u$, $A = 1$, $a = 1$, $\xi_2 = 0$,

$$\xi_1 = \frac{1}{N} \sum_{i=1}^N \left\{ f_k(X_i) - f_k(X_{i+N}) \right\} \quad \text{and}$$

$$\xi_3 = \frac{4 \frac{1}{N} \sum_{i=1}^{2N} f_k(X_i)^2}{N}.$$

We obtain:

$$\begin{aligned} P^{\otimes N} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i) - P[f_k(X)] \geq 2 \sqrt{\frac{\eta_k \left\{ \frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + P[f_k(X)^2] \right\}}{N}} \right] \\ = P^{\otimes 2N} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i) - P[f_k(X)] \geq 2 \sqrt{\frac{\eta_k \left\{ \frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + P[f_k(X)^2] \right\}}{N}} \right] \\ \leq \exp(1 - \eta_k). \end{aligned}$$

Remark that:

$$P[f_k(X)^2] = \int_{\mathcal{X}} f_k(x)^2 f(x) \lambda(dx).$$

So, using condition $\mathcal{H}(p)$ and Hlder's inequality we have:

$$\begin{aligned} P[f_k(X)^2] &\leq \left(\int_{\mathcal{X}} |f_k(x)|^{2p} \lambda(dx) \right)^{\frac{1}{p}} \left(\int_{\mathcal{X}} f(x)^q \lambda(dx) \right)^{\frac{1}{q}} \\ &\leq \left(c_k \int_{\mathcal{X}} f_k(x)^2 \lambda(dx) \right) \left(c \int_{\mathcal{X}} f(x) \lambda(dx) \right) \\ &= (c_k c) \int_{\mathcal{X}} f_k(x)^2 \lambda(dx) = C_k D_k. \end{aligned}$$

Now, let us combine this inequality with the reverse one by a union bound argument, we have:

$$\begin{aligned} P^{\otimes N} \left[\left| \frac{1}{N} \sum_{i=1}^N f_k(X_i) - P[f_k(X)] \right| \geq 2 \sqrt{\frac{\eta_k \left\{ \frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + C_k D_k \right\}}{N}} \right] \\ \leq 2 \exp(1 - \eta_k). \end{aligned}$$

We now make a union bound on $k \in \{1, \dots, m\}$ and put:

$$\eta_k = 1 + \log \frac{2m}{\varepsilon}.$$

We obtain:

$$\begin{aligned} P^{\otimes N} \left[\forall k \in \{1, \dots, m\}, \left| \frac{1}{N} \sum_{i=1}^N f_k(X_i) - P[f_k(X)] \right| \geq 2 \sqrt{\frac{(1 + \log \frac{2m}{\varepsilon}) \left\{ \frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + C_k D_k \right\}}{N}} \right] \\ \geq 1 - \varepsilon. \end{aligned}$$

We end the proof by noting that:

$$d^2(\hat{\alpha}_k f_k, \bar{\alpha}_k f_k) = \frac{\left[\frac{1}{N} \sum_{i=1}^N f_k(X_i) - P[f_k(X)] \right]^2}{\int_{\mathcal{X}} f_k(x)^2 \lambda(dx)}.$$

□

3. SOME EXAMPLES WITH RATES OF CONVERGENCE

3.1. General remarks when $(f_k)_k$ is an orthonormal family and condition $\mathcal{H}(1)$ is satisfied. In subsections 3.1, 3.2 and 3.3, we study the rate of convergence of our estimator in the special case where $(f_k)_{k \in \mathbb{N}^*}$ is an orthonormal basis of $\mathcal{L}^2$, so we have:

$$D_k = \int_{\mathcal{X}} f_k(x)^2 \lambda(dx) = 1$$

and:

$$\int_{\mathcal{X}} f_k(x) f_{k'}(x) \lambda(dx) = 0$$

if $k \neq k'$.

We also assume that condition $\mathcal{H}(1)$ is satisfied: $\forall x \in \mathcal{X}, f(x) \leq c$, remember that in this case we have taken $c_k = 1$ and so $C_k = c$, so:

$$\beta(\varepsilon, k) = \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} \right\} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + c \right].$$

Note that in this case all the order of application of the projections $\Pi_{\mathcal{RC}_{k,\varepsilon}}$ does not matter because these projections works on orthogonal directions. So we can define, once m is chosen:

$$\hat{f} = \Pi_{\mathcal{RC}_{m,\varepsilon}} \dots \Pi_{\mathcal{RC}_{1,\varepsilon}} 0 = \Pi_{\mathcal{RC}_{\{1,\dots,m\},\varepsilon}} 0 = \hat{f}_{\{1,\dots,m\}}$$

(following the notations of subsection 2.4). Note that:

$$\hat{f}(x) = \sum_{k=1}^m \text{sign}(\hat{\alpha}_k) \left(|\hat{\alpha}_k| - \sqrt{\beta(\varepsilon, k)} \right)_+ f_k(x)$$

where $\text{sign}(x)$ is the sign of x (namely $+1$ if $x > 0$ and -1 otherwise), and so $\hat{f}$ is a soft-thresholded estimator. Let us also make the following remark. As for any x , $f(x) \leq c$, we have:

$$d^2(f, 0) \leq c.$$

So the region:

$$\mathcal{B} = \left\{ g \in \mathcal{L}^2 : \forall k \in \mathbb{N}^*, \int_{\mathcal{X}} g(x) f_k(x) \lambda(dx) \leq \sqrt{c} \right\}$$

is convex, and contains f . So the projection on $\mathcal{B}$, $\Pi_{\mathcal{B}}$ can only improve $\hat{f}$. We put:

$$\tilde{f} = \Pi_{\mathcal{B}} \hat{f}.$$

Note that this transformation is needed to obtain the following theorem, but does not have practical incidence in general. Actually:

$$\tilde{f}(x) = \sum_{k=1}^m \text{sign}(\hat{\alpha}_k) \left\{ \left(|\hat{\alpha}_k| - \sqrt{\beta(\varepsilon, k)} \right)_+ \wedge \sqrt{c} \right\} f_k(x).$$

3.2. Rate of convergence in Sobolev spaces. It is well known that if f has regularity β (known by the statistician) then we have the choice

$$m = N^{\frac{1}{2\beta+1}}$$

and a standard estimation of coefficients leads to the optimal rate of convergence:

$$N^{\frac{-2\beta}{2\beta+1}}.$$

Here, we assume that we don't know β , and we show that taking $m = N$ leads to the rate of convergence:

$$N^{\frac{-2\beta}{2\beta+1}} \log N$$

namely the optimal rate of convergence up to a $\log N$ factor.

Theorem 3.1. *Let us assume that $(f_k)_{k \in \mathbb{N}^*}$ is an orthonormal basis of $\mathcal{L}^2$. Let us put:*

$$\bar{f}_m = \arg \min_{g \in \text{Span}(f_1, \dots, f_m)} d^2(g, f),$$

and let us assume that $f \in \mathcal{L}^2$ satisfies condition $\mathcal{H}(1)$ and is such that there are unknown constants $D > 0$ and $\beta \geq 1$ such that:

$$d^2(\bar{f}_m, f) \leq Dm^{-2\beta}.$$

Let us choose $m = N$ and $\varepsilon = N^{-2}$ in the definition of $\tilde{f}$. Then we have, for any $N \geq 2$:

$$P^{\otimes N} d^2(\tilde{f}, f) \leq D'(c, D) \left(\frac{\log N}{N} \right)^{\frac{2\beta}{2\beta+1}}.$$

Here again, the proof of the theorems are given at the end of the section. Let us just remark that, in the case where $\mathcal{X} = [0, 1]$, λ is the Lebesgue measure, and $(f_k)_{k \in \mathbb{N}^*}$ is the trigonometric basis, the condition:

$$d^2(\bar{f}_m, f) \leq Dm^{-2\beta}$$

is satisfied for $D = D(\beta, L) = L^2 \pi^{-2\beta}$ as soon as $f \in W(\beta, L)$ where $W(\beta, L)$ is the Sobolev class:

$$\left\{ f \in \mathcal{L}^2 : f^{(\beta-1)} \text{ is absolutely continuous and } \int_0^1 f^{(\beta)}(x)^2 \lambda(dx) \leq L^2 \right\},$$

see Tsybakov [14] for example. The minimax rate of convergence in $W(\beta, L)$ is $N^{-\frac{2\beta}{2\beta+1}}$, so we can see that our estimator reaches the best rate of convergence up to a $\log N$ factor with an unknown β .

3.3. Rate of convergence in Besov spaces. We here extend the previous result to the case of a Besov space $B_{s,p,q}$. Note that we have, for any $L \geq 0$ and $\beta \geq 0$:

$$W(\beta, L) \subset B_{\beta,2,2}$$

so this result is really an extension of the previous one (see Hrdle, Kerkycharian, Picard and Tsybakov [10], or Donoho, Johnstone, Kerkycharian and Picard [9]). We define the Besov space:

$$B_{s,p,q} = \left\{ g : [0, 1] \rightarrow \mathbb{R}, \quad g(\cdot) = \alpha \phi(\cdot) + \sum_{j=0}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k} \psi_{j,k}(\cdot), \right. \\ \left. \sum_{j=0}^{\infty} 2^{jq(s-\frac{1}{2}-\frac{1}{p})} \left[\sum_{k=1}^{2^j} |\beta_{j,k}|^p \right]^{\frac{q}{p}} = \|g\|_{s,p,q}^q < +\infty \right\},$$

with obvious changes for $p = +\infty$ or $q = +\infty$. We also define the weak Besov space:

$$\begin{aligned} W_{\rho,\pi} = & \left\{ g : [0, 1] \rightarrow \mathbb{R}, \quad g(\cdot) = \alpha\phi(\cdot) + \sum_{j=0}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k} \psi_{j,k}(\cdot), \right. \\ & \left. \sup_{\lambda>0} \lambda^{\rho} \sum_{j=0}^{\infty} 2^{j(\frac{\pi}{2}-1)} \sum_{k=1}^{2^j} \mathbb{1}_{\{|\beta_{j,k}|>\lambda\}} < +\infty \right\} \\ = & \left\{ g : [0, 1] \rightarrow \mathbb{R}, \quad g(\cdot) = \alpha\phi(\cdot) + \sum_{j=0}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k} \psi_{j,k}(\cdot), \right. \\ & \left. \sup_{\lambda>0} \lambda^{\pi-\rho} \sum_{j=0}^{\infty} 2^{j(\frac{\pi}{2}-1)} \sum_{k=1}^{2^j} |\beta_{j,k}|^{\pi} \mathbb{1}_{\{|\beta_{j,k}|\leq\lambda\}} < +\infty \right\}, \end{aligned}$$

see Cohen [7] for the equivalence of both definitions. Let us remark that $B_{s,p,q}$ is a set of functions with regularity s while $W_{\rho,\pi}$ is a set of functions with regularity:

$$s' = \frac{1}{2} \left(\frac{\pi}{\rho} - 1 \right).$$

Theorem 3.2. *Let us assume that $\mathcal{X} = [0, 1]$, and that $(\psi_{j,k})_{j=0,\dots,+\infty, k \in \{1,\dots,2^j\}}$ is a wavelet basis, together with a function ϕ , satisfying the conditions given in [9] and having regularity R (for example Daubechies' families), with ϕ and $\psi_{0,1}$ supported by $[-A, A]$. Let us assume that $f \in B_{s,p,q}$ with $R+1 \geq s > \frac{1}{p}$, $1 \leq q \leq \infty$, $2 \leq p \leq +\infty$, or that $f \in B_{s,p,q} \cap W_{\frac{2}{2s+1},2}$ with $R+1 \geq s > \frac{1}{p}$, $1 \leq p \leq +\infty$, with unknown constants s , p and q and that f satisfies condition $\mathcal{H}(1)$ with a known constant c . Let us choose:*

$$\{f_1, \dots, f_m\} = \{\phi\} \cup \{\psi_{j,k}, j = 1, \dots, 2^{\lfloor \frac{\log N}{\log 2} \rfloor}, k = 1, \dots, 2^j\}$$

(so $\frac{N}{2} \leq m \leq N$) and $\varepsilon = N^{-2}$ in the definition of $\tilde{f}$. Then we have:

$$P^{\otimes N} d^2(\tilde{f}, f) = \mathcal{O} \left(\left(\frac{\log N}{N} \right)^{\frac{2s}{2s+1}} \right).$$

Let us remark that we obtain nearly the same rate of convergence than in [9], namely the minimax rate of convergence up to a $\log N$ factor.

3.4. Kernel estimators. Here, we assume that $\mathcal{X} = \mathbb{R}$ and that f is compactly supported, say by $[0, 1]$. We put, for any $m \in \mathbb{N}$ and $k \in \{1, \dots, m\}$:

$$f_k(x) = K \left(\frac{k}{m}, x \right)$$

where K is some function $\mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ and we obtain some estimator that has the form of a kernel estimator:

$$\hat{f}_{\{1,\dots,m\}}(x) = \sum_{k=1}^m \tilde{\alpha}_k K \left(\frac{k}{m}, x \right).$$

Moreover, is is possible to use a multiple kernel estimator. Let us choose $n \in \mathbb{N}$, $h \in \mathbb{N}$, h kernels $K_1, \dots, K_h$ and put, for any $k = i + n * j \in \{1, \dots, m = hn\}$:

$$f_k(x) = K_j \left(\frac{i}{n}, x \right).$$

We obtain a multiple kernel estimator:

$$\hat{f}_{\{1, \dots, m\}}(x) = \sum_{i=1}^n \sum_{j=1}^h \tilde{\alpha}_{i+nj} K_j \left(\frac{i}{n}, x \right).$$

3.5. Proof of the theorems.

Proof of theorem 3.1. Let us begin the proof with a general m and ε , the reason of the choice $m = N$ and $\varepsilon = N^{-2}$ will become clear. Let us also write $\mathcal{E}(\varepsilon)$ the event satisfied with probability at least $1 - \varepsilon$ in theorem 2.1. We have:

$$P^{\otimes N} d^2(\tilde{f}, f) = P^{\otimes N} \left[\mathbf{1}_{\mathcal{E}(\varepsilon)} d^2(\tilde{f}, f) \right] + P^{\otimes N} \left[(1 - \mathbf{1}_{\mathcal{E}(\varepsilon)}) d^2(\tilde{f}, f) \right].$$

For the first term we have:

$$d^2(\tilde{f}, f) \leq 2 \int_{\mathcal{X}} f(x)^2 \lambda(dx) + 2 \int_{\mathcal{X}} \tilde{f}(x)^2 \lambda(dx) \leq 2c + 2mc = 2(m+1)c$$

and so:

$$P^{\otimes N} \left[(1 - \mathbf{1}_{\mathcal{E}(\varepsilon)}) d^2(\tilde{f}, f) \right] \leq 2\varepsilon(m+1)c.$$

For the other term, just remark that under $\mathcal{E}(\varepsilon)$:

$$\begin{aligned} d^2(\tilde{f}, f) &= d^2(\Pi_{\mathcal{B}} \Pi_{\mathcal{C}\mathcal{R}_{m,\varepsilon}} \dots \Pi_{\mathcal{C}\mathcal{R}_{1,\varepsilon}} 0, f) \\ &\leq d^2(\Pi_{\mathcal{C}\mathcal{R}_{m,\varepsilon}} \dots \Pi_{\mathcal{C}\mathcal{R}_{1,\varepsilon}} 0, f) \leq d^2(\Pi_{\mathcal{C}\mathcal{R}_{m',\varepsilon}} \dots \Pi_{\mathcal{C}\mathcal{R}_{1,\varepsilon}} 0, f) \end{aligned}$$

for any $m' \leq m$, because of theorem 2.1, more precisely of corollary 2.3. And we have:

$$\begin{aligned} d^2(\Pi_{\mathcal{M}_{m'}} \dots \Pi_{\mathcal{M}_1} 0, f) \\ \leq \sum_{k=1}^{m'} \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} \right\} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + c \right] + d^2(\bar{f}_m, f). \end{aligned}$$

So we have:

$$\begin{aligned} P^{\otimes N} \left[\mathbf{1}_{\mathcal{E}(\varepsilon)} d^2(\tilde{f}, f) \right] &\leq P^{\otimes N} \left[d^2(\tilde{f}, f) \right] \\ &\leq P^{\otimes N} \sum_{k=1}^{m'} \left\{ \frac{4 \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} \right\} \left[\frac{1}{N} \sum_{i=1}^N f_k(X_i)^2 + c \right] + (m')^{-2\beta} D \\ &\leq \frac{8m'c \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} + (m')^{-2\beta} D. \end{aligned}$$

So finally, we obtain, for any $m' \leq m$:

$$P^{\otimes N} d^2(\tilde{f}, f) \leq \frac{8m'c \left[1 + \log \frac{2m}{\varepsilon} \right]}{N} + (m')^{-2\beta} D + 2\varepsilon(m+1)c.$$

The choice of:

$$m' = \left(\frac{N}{\log N} \right)^{\frac{1}{2\beta+1}}$$

leads to a first term of order $N^{\frac{-2\beta}{2\beta+1}} \log \frac{m}{\varepsilon} (\log N)^{-\frac{1}{2\beta+1}}$ and a second term of order $N^{\frac{-2\beta}{2\beta+1}} (\log N)^{\frac{2\beta}{2\beta+1}}$. The choice of $m = N$ and $\varepsilon = N^{-2}$ gives a first and second term at order:

$$\left(\frac{\log N}{N} \right)^{\frac{2\beta}{2\beta+1}}$$

while keeping the third term at order N^{-1} . This proves the theorem. $\square$

Proof of theorem 3.2. Here again let us write $\mathcal{E}(\varepsilon)$ the event satisfied with probability at least $1 - \varepsilon$ in theorem 2.1. We have:

$$P^{\otimes N} d^2(\tilde{f}, f) = P^{\otimes N} \left[\mathbb{1}_{\mathcal{E}(\varepsilon)} d^2(\tilde{f}, f) \right] + P^{\otimes N} \left[(1 - \mathbb{1}_{\mathcal{E}(\varepsilon)}) d^2(\tilde{f}, f) \right].$$

For the first term we still have:

$$d^2(\tilde{f}, f) \leq 2(m+1)c.$$

For the second term, let us write the development of f into our wavelet basis:

$$f = \alpha\phi + \sum_{j=0}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k} \psi_{j,k},$$

and:

$$\hat{f}(x) = \tilde{\alpha}\phi + \sum_{j=0}^J \sum_{k=1}^{2^j} \tilde{\beta}_{j,k} \psi_{j,k}$$

the estimator $\hat{f}$. Let us put:

$$J = \left\lfloor \frac{\log N}{\log 2} \right\rfloor.$$

For any $J' \leq J$ we have:

$$\begin{aligned} d^2(\tilde{f}, f) &= d^2(\Pi_{\mathcal{B}} \Pi_{C\mathcal{R}_{m,\varepsilon}} \dots \Pi_{C\mathcal{R}_{1,\varepsilon}} 0, f) \leq d^2(\Pi_{C\mathcal{R}_{m,\varepsilon}} \dots \Pi_{C\mathcal{R}_{1,\varepsilon}} 0, f) \\ &= (\tilde{\alpha} - \alpha)^2 + \sum_{j=0}^J \sum_{k=1}^{2^j} (\tilde{\beta}_{j,k} - \beta_{j,k})^2 + \sum_{j=J+1}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k}^2 \\ &\leq (\tilde{\alpha} - \alpha)^2 + \sum_{j=0}^{J'} \sum_{k=1}^{2^j} (\tilde{\beta}_{j,k} - \beta_{j,k})^2 \mathbb{1}(|\beta_{j,k}| \geq \kappa) + \sum_{j=0}^{J'} \sum_{k=1}^{2^j} \beta_{j,k}^2 \mathbb{1}(|\beta_{j,k}| < \kappa) \\ &\quad + \sum_{j=J'+1}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k}^2 \end{aligned}$$

for any $\kappa \geq 0$, as soon as $\mathcal{E}(\varepsilon)$ is satisfied (here again we applied theorem 2.1). In the case where $p \geq 2$ we can take:

$$J' = \left\lfloor \frac{\log N^{\frac{1}{1+2s}}}{\log 2} \right\rfloor$$

and $\kappa = 0$ to obtain (let C be a generic constant in the whole proof):

$$\sum_{j=J'+1}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k}^2 \leq \sum_{j=J'+1}^{\infty} \left(\sum_{k=1}^{2^j} \beta_{j,k}^p \right)^{\frac{2}{p}} 2^{j(1-\frac{2}{p})}.$$

As $f \in B_{s,p,q} \subset B_{s,p,\infty}$ we have:

$$\left(\sum_{k=1}^{2^j} \beta_{j,k}^p \right)^{\frac{2}{p}} \leq C 2^{-2j(s+\frac{1}{2}-\frac{1}{p})}$$

and so:

$$\sum_{j=J+1}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k}^2 \leq C 2^{-2J's} \leq C N^{\frac{-2s}{1+2s}},$$

and:

$$\begin{aligned} \sum_{j=0}^J \sum_{k=1}^{2^j} (\tilde{\beta}_{j,k} - \beta_{j,k})^2 \mathbb{1}(|\beta_{j,k}| \geq \kappa) &\leq \frac{8c [1 + \log \frac{2m}{\varepsilon}]}{N} \sum_{j=0}^J \sum_{k=1}^{2^j} 1 \\ &\leq \frac{8c [1 + \log \frac{2m}{\varepsilon}]}{N} 2^{J'+1} \leq C \frac{8c [1 + \log \frac{2m}{\varepsilon}]}{N} N^{\frac{1}{1+2s}}. \end{aligned}$$

So we obtain the desired rate of convergence. In the case where $p < 2$ we let $J' = J$ and proceed as follows.

$$\begin{aligned} \sum_{j=0}^J \sum_{k=1}^{2^j} (\tilde{\beta}_{j,k} - \beta_{j,k})^2 \mathbb{1}(|\beta_{j,k}| \geq \kappa) &\leq \frac{8c [1 + \log \frac{2m}{\varepsilon}]}{N} \sum_{j=0}^J \sum_{k=1}^{2^j} \mathbb{1}(|\beta_{j,k}| \geq \kappa) \\ &\leq \frac{8c [1 + \log \frac{2m}{\varepsilon}]}{N} C \kappa^{-\frac{2}{2s+1}} \end{aligned}$$

because f is also assumed to be in the weak Besov space. We also have:

$$\sum_{j=0}^J \sum_{k=1}^{2^j} \beta_{j,k}^2 \mathbb{1}(|\beta_{j,k}| < \kappa) \leq C \kappa^{2 - \frac{2}{1+2s}}.$$

For the remainder term we use (see [10, 9]):

$$B_{s,p,q} \subset B_{s-\frac{1}{p}+\frac{1}{2},2,q}$$

to obtain:

$$\sum_{j=J+1}^{\infty} \sum_{k=1}^{2^j} \beta_{j,k}^2 \leq C 2^{-2J(s+\frac{1}{2}-\frac{1}{p})} \leq C 2^{-J}$$

as $s > \frac{1}{p}$. Let us remember that:

$$\frac{N}{2} \leq m = 2^J \leq N$$

and that $\varepsilon = N^{-2}$, and take:

$$\kappa = \sqrt{\frac{\log N}{N}}$$

to obtain the desired rate of convergence. $\square$

4. BETTER BOUNDS AND GENERALIZATIONS

Actually, as pointed out by Catoni [6], the symmetrization technique used in the proof of theorem 2.1 causes the loss of a factor 2 in the bound because we upper bound the variance of two samples instead of 1. In this section, we try to use this remark to improve our bound, using techniques already used by Catoni [4]. We also give a generalization of the obtained result that allows us to use a family $(f_1, \dots, f_m)$ of functions that is data-dependant. The technique used is due to Seeger [13], and it will allow us to use kernel estimators as Support Vector Machines.

Remark that the estimation technique described in section 2 does not necessarily require a bound on $d^2(\hat{\alpha}_k f_k, \bar{\alpha}_k f_k)$. Actually, a simple confidence interval on $\bar{\alpha}_k$ is sufficient.

4.1. An improvement of theorem 2.1 under condition $\mathcal{H}(+\infty)$. Let us remember that $\mathcal{H}(+\infty)$ just means that every f_k is bounded by $\sqrt{C_k D_k}$.

Theorem 4.1. *Under condition $\mathcal{H}(+\infty)$, for any $\varepsilon > 0$, for any $\beta_{k,1}, \beta_{k,2}$ such that:*

$$0 < \beta_{k,j} < \frac{N}{\sqrt{C_k D_k}}, \quad j \in \{1, 2\},$$

with $P^{\otimes N}$ -probability at least $1 - \varepsilon$, for any $k \in \{1, \dots, m\}$ we have:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) \leq \bar{\alpha}_k \leq \alpha_k^{\sup}(\varepsilon, \beta_{k,2})$$

with:

$$\alpha_k^{\sup}(\varepsilon, \beta_{k,2}) = \frac{N - N \exp \left[\frac{1}{N} \sum_{i=1}^N \log \left(1 - \frac{\beta_{k,2}}{N} f_k(X_i) \right) - \frac{\log \frac{2m}{\varepsilon}}{N} \right]}{D_k \beta_{k,2}}$$

and:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) = \frac{N \exp \left[\frac{1}{N} \sum_{i=1}^N \log \left(1 + \frac{\beta_{k,1}}{N} f_k(X_i) \right) - \frac{\log \frac{2m}{\varepsilon}}{N} \right] - N}{D_k \beta_{k,1}}.$$

Before we give the proof, let us see why this theorem really improves theorem 2.1. Let us choose put:

$$V_k = P \left\{ [f_k(X) - P(f_k(X))]^2 \right\}$$

and:

$$\beta_{k,1} = \beta_{k,2} = \sqrt{\frac{N \log \frac{2m}{\varepsilon}}{V_k}}.$$

Then we obtain:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) = \hat{\alpha}_k - \sqrt{\frac{2V_k \log \frac{2m}{\varepsilon}}{N}} + \mathcal{O}_P \left(\frac{\log \frac{2m}{\varepsilon}}{N} \right)$$

and:

$$\alpha_k^{\sup}(\varepsilon, \beta_{k,2}) = \hat{\alpha}_k + \sqrt{\frac{2V_k \log \frac{2m}{\varepsilon}}{N}} + \mathcal{O}_P \left(\frac{\log \frac{2m}{\varepsilon}}{N} \right).$$

So, the first order term for $d^2(\hat{\alpha}_k f_k, \bar{\alpha}_k f_k)$ is:

$$\frac{2V_k \log \frac{2m}{\varepsilon}}{N},$$

there is an improvement by a factor 4 when we compare this bound to theorem 2.1.

Remark that this particular choice for $\beta_{k,1}$ and $\beta_{k,2}$ is valid as soon as:

$$\sqrt{\frac{N \log \frac{2m}{\varepsilon}}{V_k}} < \frac{N}{\sqrt{C_k D_k}}$$

or equivalently as soon as N is greater than

$$\frac{C_k D_k \log \frac{2m}{\varepsilon}}{V_k}.$$

In practice, however, this particular $\beta_{k,1}$ and $\beta_{k,2}$ are unknown. We can use the following procedure (see Catoni [6]). We choose a value $a > 1$ and:

$$B = \left\{ a^l, 0 \leq l \leq \left\lfloor \frac{\log \frac{N}{\sqrt{C_k D_k}}}{\log a} \right\rfloor - 1 \right\}.$$

By taking a union bound over all possible values of B , with:

$$|B| \leq \frac{\log \frac{N}{\sqrt{C_k D_k}}}{\log a}$$

we obtain the following corollary.

Corollary 4.2. *Under condition $\mathcal{H}(+\infty)$, for any $a > 1$, for any $\varepsilon > 0$, with $P^{\otimes N}$ -probability at least $1 - \varepsilon$ we have:*

$$\sup_{\beta \in B} \alpha_k^{\inf} \left(\frac{\varepsilon \log a}{\log N - \frac{1}{2} \log C_k D_k}, \beta \right) \leq \bar{\alpha}_k \leq \inf_{\beta \in B} \alpha_k^{\sup} \left(\frac{\varepsilon \log a}{\log N - \frac{1}{2} \log C_k D_k}, \beta \right),$$

with:

$$B = \left\{ a^l, 0 \leq l \leq \left\lfloor \frac{\log \frac{N}{\sqrt{C_k D_k}}}{\log a} \right\rfloor - 1 \right\}.$$

Note that the price to pay for the optimization with respect to $\beta_{k,1}$ and $\beta_{k,2}$ was just a $\log \log N$ factor.

Proof of the theorem. The technique used in the proof is due to Catoni [5]. Let us choose $k \in \{1, \dots, m\}$, and:

$$\beta \in \left(0, \frac{N}{\sqrt{C_k D_k}} \right).$$

We have, for any $\eta \in \mathbb{R}$:

$$P^{\otimes N} \exp \left\{ \sum_{i=1}^N \log \left(1 - \frac{\beta}{N} f_k(X_i) \right) - \eta \right\} \leq \exp \left\{ N \log \left(1 - \frac{\beta}{N} P[f_k(X)] \right) - \eta \right\}.$$

Let us choose:

$$\eta = \log \frac{2m}{\varepsilon} + N \log \left(1 - \frac{\beta}{N} P[f_k(X)] \right).$$

We obtain:

$$P^{\otimes N} \exp \left\{ \sum_{i=1}^N \log \left(1 - \frac{\beta}{N} f_k(X_i) \right) - \log \frac{2m}{\varepsilon} - N \log \left(1 - \frac{\beta}{N} P[f_k(X)] \right) \right\} \leq \frac{\varepsilon}{2m},$$

and so:

$$P^{\otimes N} \left\{ \sum_{i=1}^N \log \left(1 - \frac{\beta}{N} f_k(X_i) \right) \geq \log \frac{2m}{\varepsilon} + N \log \left(1 - \frac{\beta}{N} P[f_k(X)] \right) \right\} \leq \frac{\varepsilon}{2m},$$

that becomes:

$$P^{\otimes N} \left\{ P[f_k(X)] \geq \frac{N}{\beta} \left[1 - \exp \left(\frac{1}{N} \sum_{i=1}^N \log \left(1 - \frac{\beta}{N} f_k(X_i) \right) - \frac{\log \frac{2m}{\varepsilon}}{N} \right) \right] \right\} \leq \frac{\varepsilon}{2m}.$$

We apply the same technique to:

$$P^{\otimes N} \exp \left\{ \sum_{i=1}^N \log \left(1 + \frac{\beta'}{N} f_k(X_i) \right) - \eta \right\} \leq \exp \left\{ N \log \left(1 + \frac{\beta'}{N} P[f_k(X)] \right) - \eta \right\}$$

to obtain the upper bound. We combine both result by a union bound argument. $\square$

4.2. A generalization to data-dependent basis functions. We now extend the previous method to the case where the family $(f_1, \dots, f_m)$ is allowed to be data-dependant, in a particular sense. This subsection requires some modifications of the notations of section 2.

Definition 4.1. For any $m' \in \mathbb{N}^*$ we define a function $\Theta_{m'} : \mathcal{X} \rightarrow (\mathcal{L}^2)^{m'}$. For any $i \in \{1, \dots, N\}$ we put:

$$\Theta_{m'}(X_i) = (f_{i,1}, \dots, f_{i,m'}).$$

Finally, consider the family of functions:

$$(f_1, \dots, f_m) = (f_{1,1}, \dots, f_{1,m'}, \dots, f_{N,1}, \dots, f_{N,m'}).$$

So we have $m = m'N$ (of course, m' is allowed to depend on N). Let us take, for any $i \in \{1, \dots, N\}$:

$$P_i(\cdot) = P^{\otimes N}(\cdot | X_i).$$

We put, for any $(i, k) \in \{1, \dots, N\} \times \{1, \dots, m'\}$:

$$D_{i,k} = \int_{\mathcal{X}} f_{i,k}(x)^2 \lambda(dx),$$

and we still assume that condition $\mathcal{H}(\infty)$ is satisfied, that means here that we have known constants $C_{i,k} = c_{i,k}$ such that:

$$\forall x \in \mathcal{X}, \quad |f_{i,k}(x)| \leq \sqrt{C_{i,k} D_{i,k}}.$$

Finally, we put:

$$\bar{\alpha}_{i,k} = \arg \min_{\alpha \in \mathbb{R}} d^2(\alpha f_{i,k}, f).$$

Let us choose $(i, k) \in \{1, \dots, N\} \times \{1, \dots, m'\}$. Using Seeger's idea, we follow the preceding proof, replacing $P^{\otimes N}$ by P_i , and using the $N - 1$ random variables:

$$\left(f_{i,k}(X_j) \right)_{\substack{j \in \{1, \dots, N\} \\ j \neq i}}$$

with

$$\eta = \log \frac{2m'N}{\varepsilon} + (N - 1) \log \left(1 - \frac{\beta}{N - 1} P[f_{i,k}(X)] \right)$$

and we obtain:

$$\begin{aligned} P_i \exp \left\{ \sum_{j \neq i} \log \left(1 - \frac{\beta}{N - 1} f_{i,k}(X_j) \right) - \log \frac{2m'N}{\varepsilon} \right. \\ \left. - (N - 1) \log \left(1 - \frac{\beta}{N - 1} P[f_{i,k}(X)] \right) \right\} \leq \frac{\varepsilon}{2m'N}. \end{aligned}$$

Note that for any random variable H that is a function of the X_i :

$$P^{\otimes N} P_i H = P^{\otimes N} H.$$

So we conclude exactly in the same way than for the previous theorem and we obtain the following result.

Theorem 4.3. For any $\varepsilon > 0$, for any $\beta_{i,k,1}, \beta_{i,k,2}$ such that:

$$0 < \beta_{i,k,j} < \frac{N - 1}{\sqrt{C_{i,k} D_{i,k}}}, \quad j \in \{1, 2\},$$

with $P^{\otimes N}$ -probability at least $1 - \varepsilon$, for any $i \in \{1, \dots, N\}$ and $k \in \{1, \dots, m\}$ we have:

$$\tilde{\alpha}_k^{\inf}(\varepsilon, \beta_{i,k,1}) \leq \bar{\alpha}_k \leq \tilde{\alpha}_k^{\sup}(\varepsilon, \beta_{i,k,2})$$

with:

$$\begin{aligned} \tilde{\alpha}_k^{\text{sup}}(\varepsilon, \beta_{i,k,2}) \\ = \frac{N-1 - (N-1) \exp \left[\frac{1}{N-1} \sum_{j \neq i} \log \left(1 - \frac{\beta_{i,k,2}}{N-1} f_{i,k}(X_j) \right) - \frac{\log \frac{2m'N}{\varepsilon}}{N-1} \right]}{D_{i,k} \beta_{i,k,2}} \end{aligned}$$

and:

$$\begin{aligned} \tilde{\alpha}_k^{\text{inf}}(\varepsilon, \beta_{i,k,1}) \\ = \frac{(N-1) \exp \left[\frac{1}{N-1} \sum_{j \neq i} \log \left(1 + \frac{\beta_{i,k,1}}{N-1} f_{i,k}(X_j) \right) - \frac{\log \frac{2m'N}{\varepsilon}}{N-1} \right] - N + 1}{D_{i,k} \beta_{i,k,1}}. \end{aligned}$$

Example 4.1 (Support Vector Machines). *Actually, SVM were firstly introduced by Guyon, Boser and Vapnik [3] in the context of classification, but the method was extended by Vapnik [15] to the context of least square regression estimation and of density estimation. The idea is to generalize the kernel estimator to the case where $\mathcal{X}$ is of large dimension, and so we cannot use a grid like we did in the $[0, 1]$ case. Let us choose a function:*

$$\begin{aligned} K : \mathcal{X}^2 &\rightarrow \mathbb{R} \\ (x, x') &\mapsto K(x, x'). \end{aligned}$$

We take $m' = 1$ and:

$$\Theta_1(x) = (K(x, \cdot))$$

then the obtained estimator has the form of a SVM:

$$\hat{f}(x) = \sum_{i=1}^N \tilde{\alpha}_i K(X_i, x)$$

where the set of i such that $\tilde{\alpha}_i \neq 0$ is expected to be small. Note that we do not need to assume that $K(\cdot, \cdot)$ is a Mercer's kernel as usual with SVM. Moreover, we can extend the method to the case where we have several kernels $K_1, \dots, K_{m'}$ by taking:

$$\Theta_{m'}(x) = (K_1(x, \cdot), \dots, K_{m'}(x, \cdot)).$$

The estimator becomes:

$$\hat{f}(x) = \sum_{j=1}^{m'} \sum_{i=1}^N \tilde{\alpha}_{i,j} K_j(X_i, x).$$

Note that a widely used kernel is the gaussian kernel; let $\delta(\cdot, \cdot)$ be a distance on $\mathcal{X}$ and $\gamma_1, \dots, \gamma_{m'} > 0$ then we put:

$$K_k(x, x') = \exp(-\gamma_k \delta^2(x, x')).$$

For example, if $\mathcal{X} = \mathbb{R}$ and λ is the Lebesgue measure then hypothesis $\mathcal{H}(\infty)$ is obviously satisfied with the gaussian kernel with

$$C_{i,k} = c_{i,k} = \sqrt{\frac{\gamma_k}{\pi}} = \frac{1}{D_{i,k}}.$$

4.3. Back to the histogram. In the case of the histogram, $f_k(\cdot) = \mathbb{1}_{A_k}(\cdot)$ can take only two values: 0 and 1. Remember that $D_k = \lambda(A_k)$. So:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) = \frac{N}{\lambda(A_k)\beta_{k,1}} \left\{ \left[\left(1 + \frac{\beta_{k,1}}{N} \right)^{|\{i: X_i \in A_k\}|} \frac{\varepsilon}{2m} \right]^{\frac{1}{N}} - 1 \right\}.$$

Remember that, for any $x \geq 0$:

$$(1+x)^\gamma \geq 1 + \gamma x + \frac{\gamma(\gamma-1)}{2} x^2$$

and so:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) \geq \hat{\alpha}_k \left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} \left[1 - \frac{\beta_{k,1}(1 - \hat{\alpha}_k D_k)}{2N} \right] - \frac{N}{D_k \beta_{k,1}} \left[1 - \left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} \right].$$

Now, we take the grid:

$$B = \left\{ 2^l, 0 \leq l \leq \left\lfloor \frac{\log \frac{N}{\sqrt{D_k}}}{\log 2} \right\rfloor - 1 \right\}.$$

Remark that, for any β in:

$$\left[1, \frac{N}{2\sqrt{D_k}} \right]$$

there is some $b \in B$ such that $\beta \leq b \leq 2\beta$, and so:

$$\alpha_k^{\inf}(\varepsilon, b) \geq \hat{\alpha}_k \left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} \left[1 - \frac{\beta_{k,1}(1 - \hat{\alpha}_k D_k)}{2N} \right] - \frac{N}{D_k 2\beta_{k,1}} \left[1 - \left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} \right].$$

This allows us to choose whatever value for $\beta_{k,1}$ in

$$\left[1, \frac{N}{2\sqrt{D_k}} \right].$$

Let us choose:

$$\beta_{k,1} = \sqrt{\frac{N^2 \left[\left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} - 1 \right]}{\hat{\alpha}_k D_k (1 - \hat{\alpha}_k D_k)}}$$

that is allowed for N large enough. So we have:

$$\alpha_k^{\inf}(\varepsilon, \beta_{k,1}) \geq \hat{\alpha}_k \left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} - \sqrt{\hat{\alpha}_k D_k (1 - \hat{\alpha}_k D_k) \left[\left(\frac{\varepsilon}{2m} \right)^{\frac{1}{N}} - 1 \right]}.$$

With the union bound term (over the grid B) we obtain:

$$\begin{aligned} & \alpha_k^{\inf} \left(\frac{\varepsilon \log 2}{\log \frac{N}{\sqrt{D_k}}}, \beta_{k,1} \right) \\ & \geq \hat{\alpha}_k \left(\frac{\varepsilon \log 2}{2m \log \frac{N}{\sqrt{D_k}}} \right)^{\frac{1}{N}} - \sqrt{\hat{\alpha}_k D_k (1 - \hat{\alpha}_k D_k) \left[\left(\frac{\varepsilon \log 2}{2m \log \frac{N}{\sqrt{D_k}}} \right)^{\frac{1}{N}} - 1 \right]} \\ & = \hat{\alpha}_k - \sqrt{\frac{\hat{\alpha}_k D_k (1 - \hat{\alpha}_k D_k) \log \frac{2m \log \frac{N}{\sqrt{D_k}}}{\varepsilon \log 2}}{N}} + \mathcal{O} \left(\frac{\log \frac{m \log N}{\varepsilon}}{N} \right), \end{aligned}$$

remark that we have this time the "real" variance term of $\mathbb{1}_{A_k}(X)$:

$$\hat{\alpha}_k D_k (1 - \hat{\alpha}_k D_k) = \frac{|\{i : X_i \in A_k\}|}{N} \left(1 - \frac{|\{i : X_i \in A_k\}|}{N} \right).$$

4.4. Another simple example: the Haar basis. Let us assume that $\mathcal{X} = [0, 1]$. Let (φ, ψ) be a father wavelet and the associated mother wavelet, and:

$$\psi_{j,k}(x) = \psi(2^j x + k)$$

for $k \in \{0, \dots, 2^j - 1\} = S_j$ (note that the wavelet basis is non-normalized here). Here, we use the Haar wavelets, with:

$$\begin{aligned}\varphi(x) &= \mathbb{1}_{[0,1]}(x) \\ \psi(x) &= \mathbb{1}_{[0, \frac{1}{2}]}(x) - \mathbb{1}_{[\frac{1}{2}, 1]}(x).\end{aligned}$$

For the sake of simplicity, let us write:

$$\psi_{-1,k}(x) = \varphi(x)$$

for $k \in \{0\} = S_{-1}$. By an obvious adaptation of our notations, let us put $\bar{\alpha}_{j,k}$ the coefficient associated to $\psi_{j,k}$:

$$\bar{\alpha}_{j,k} = \frac{P\psi_{j,k}(X)}{\int \psi_{j,k}^2} = P\psi_{j,k}(X),$$

remark that condition $\mathcal{H}(\infty)$ is satisfied with $D_{j,k} = 2^{-j}$ and $C_{j,k} = 1$. In this particular setting, note that $\bar{\alpha}_{-1,0} = 1$ is known, so the associated confidence interval is just $\{1\}$. Moreover, here $\psi_{j,k}(X)$ can take only three values: -1 , 0 and 1 . Let us put:

$$\bar{P} = \frac{1}{N} \sum_{i=1}^N \delta_{X_i}.$$

Remark that in this case we have:

$$\begin{aligned}\frac{1}{N} \sum_{i=1}^N \log \left(1 - \frac{\beta}{N} \psi_{j,k}(X_i) \right) \\ = \bar{P}(\psi_{j,k}(X) = 1) \log \left(1 - \frac{\beta}{N} \right) + \bar{P}(\psi_{j,k}(X) = -1) \log \left(1 + \frac{\beta}{N} \right) \\ = \frac{1}{2} \bar{P} [\psi_{j,k}(X)^2] \log \left(1 - \frac{\beta^2}{N^2} \right) + \frac{1}{2} \bar{P} [\psi_{j,k}(X)] \log \left(\frac{1 - \frac{\beta}{N}}{1 + \frac{\beta}{N}} \right).\end{aligned}$$

So we have:

$$\begin{aligned}\alpha_{j,k}^{\sup}(\varepsilon, \beta) \\ = \frac{N - N \exp \left[\frac{1}{2} \bar{P} [\psi_{j,k}(X)^2] \log \left(1 - \frac{\beta^2}{N^2} \right) - \frac{1}{2} \bar{P} [\psi_{j,k}(X)] \log \left(\frac{1 + \frac{\beta}{N}}{1 - \frac{\beta}{N}} \right) - \frac{\log \frac{2m}{\varepsilon}}{N} \right]}{D_k \beta_{k,2}}\end{aligned}$$

and:

$$\begin{aligned}\alpha_{j,k}^{\inf}(\varepsilon, \beta) \\ = \frac{N \exp \left[\frac{1}{2} \bar{P} [\psi_{j,k}(X)^2] \log \left(1 - \frac{\beta^2}{N^2} \right) + \frac{1}{2} \bar{P} [\psi_{j,k}(X)] \log \left(\frac{1 + \frac{\beta}{N}}{1 - \frac{\beta}{N}} \right) - \frac{\log \frac{2m}{\varepsilon}}{N} \right] - N}{D_k \beta_{k,1}}.\end{aligned}$$

5. SIMULATIONS

5.1. Description of the example. We assume that we observe X_i for $i \in \{1, \dots, N\}$ with $N = 2^{10} = 1024$, where the variables $X_i \in [0, 1] \subset \mathbb{R}$ are i.i.d. from a distribution with an unknown density f with respect to the Lebesgue measure. The goal is to estimate f .

Here, we will use three methods. The first estimation method will be a multiple kernel estimator obtained by the algorithm described previously, the second one a

thresholded wavelets estimate also obtained by this algorithm, and we will compare both estimators to a thresholded wavelet estimate as given by Donoho, Johnstone, Kerkycharian and Picard [8].

5.2. The estimators.

5.2.1. *Hard-thresholded wavelet estimator.* We first use a classical hard-thresholded wavelet estimator.

In the case of the Haar basis (see subsection 4.4), we take:

$$\hat{\alpha}_{j,k} = 2^j \frac{1}{N} \sum_{i=1}^N \psi_{j,k}(X_i).$$

For a given $\kappa \geq 0$ and $J \in \mathbb{N}$, we take:

$$\tilde{f}_J(\cdot) = \sum_{j=-1}^J \sum_{k \in S_j} \hat{\alpha}_{j,k} \mathbb{1}(|\hat{\alpha}_{j,k}| \geq \kappa t_{j,N}) \psi_{j,k}(\cdot)$$

where:

$$t_{j,N} = \sqrt{\frac{j}{N}}.$$

Actually, we must choose J in such a way that:

$$2^J \sim t_N^{-1}.$$

Here, we choose $\kappa = 0.7$ and $J = 7$.

5.2.2. *Wavelet estimators with our algorithm.* We also use the same family of functions, and we apply our thresholding method, with bounds given in subsection 4.4. So we take:

$$m = 2^J = 128.$$

We use an asymptotic version of our confidence intervals inspired by our theoretical confidence intervals:

$$\bar{\alpha}_{j,k} \in \left[\hat{\alpha}_{j,k} \pm \sqrt{\frac{\log \frac{2m}{\varepsilon} V_{j,k}}{N}} \right]$$

where $V_{j,k}$ is the estimated variance of $\psi_{j,k}(X)$:

$$V_{j,k} = \frac{1}{N} \sum_{i=1}^N \left[\psi_{j,k}(X_i) - \frac{1}{N} \sum_{h=1}^N \psi_{j,k}(X_h) \right]^2.$$

Let us remark that the union bound are always "pessimistic", and that we use a union bound argument over all the m models despite only a few of them are effectively used in the estimator. So, we propose to actually use the individual confidence interval for each model, replacing: the $\log \frac{2m}{\varepsilon}$ by $\log \frac{2}{\varepsilon}$.

5.2.3. *Multiple estimator.* Finally, we use the kernel estimator described in section 3, with function K :

$$K_j(u, v) = \exp \left[-2^{2j} (u - v)^2 \right]$$

with $n = N$ and $j \in \{1, \dots, h = 6\}$. We add the constant function 1 to the family.

Here again we use the individuals confidence intervals, and the asymptotic version of this intervals.

FIGURE 1. Values of t_i and c_i in the fonction $Blocks(\cdot)$.

i	1	2	3	4	5	6	7	8	9	10	11
c_i	4	-5	3	-4	5	4.2	-2.1	4.3	-3.1	2.1	-4.2
t_i	0.10	0.13	0.15	0.23	0.25	0.40	0.44	0.65	0.76	0.78	0.81

FIGURE 2. Results of the experiments. For each experiment, we give the mean distance of the estimator the the density ($d^2(\cdot, f)$).

Function $f(\cdot)$	standard thresh- olded wavelets	thresh. wav. with our method	multiple kernel
<i>Doppler</i>	0.104	0.127	0.083
<i>HeaviSine</i>	0.071	0.066	0.040
<i>Blocks</i>	0.110	0.142	0.121

5.3. Experiments and results. The simulations were realized with the R software [12].

For the experiments, we use the following functions f that are some variations of the functions used by Donoho and Johnstone for experiments on wavelets, for example in [8] (actually, these functions were used as regression functions, so the modification was to add them a constant in order to ensure they take nonnegative values):

$$Doppler(t) = 1 + 2\sqrt{t(1-t)} \sin \frac{2\pi(1+v)}{t+v} \quad \text{where } v = 0.05$$

$$HeaviSine(t) = 1.5 + \frac{1}{4} \left[4 \sin 4\pi t - \operatorname{sgn}(t - 0.3) - \operatorname{sgn}(0.72 - t) \right]$$

$$Blocks(t) = 1.05 + \frac{1}{4} \sum_{i=1}^{11} c_i \mathbb{1}_{(t_i, +\infty)}(t)$$

where $\operatorname{sgn}(t)$ is the sign of t (say -1 if $t \leq 0$ and $+1$ otherwise). The values of the c_i and t_i are given in figure 1.

We consider 3 experiments (for the three density functions), we choose $\varepsilon=10\%$, repeat each experiment 20 times; the results are reported in figure 2. We also give some illustrations (figure 3, 4 and 5).

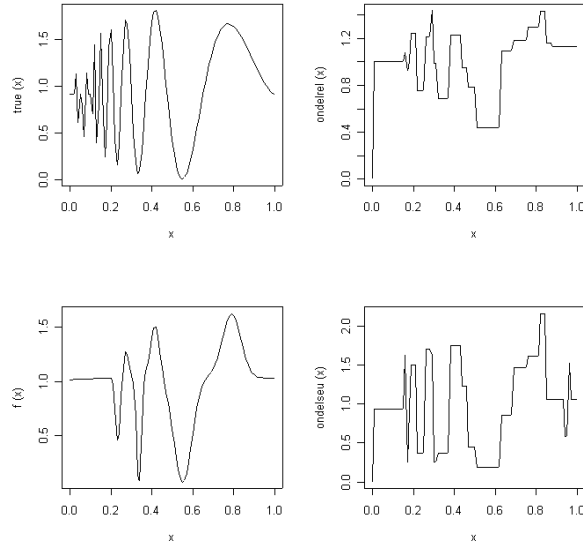


FIGURE 3. Experiment 1, $f = \text{Doppler}$. Up-left: true regression function (*true*). Down-left: SVM (f). Up-right: wavelet estimate with our algorithm (*ondelrel*). Down-right: "classical" wavelet estimate (*ondelseu*).

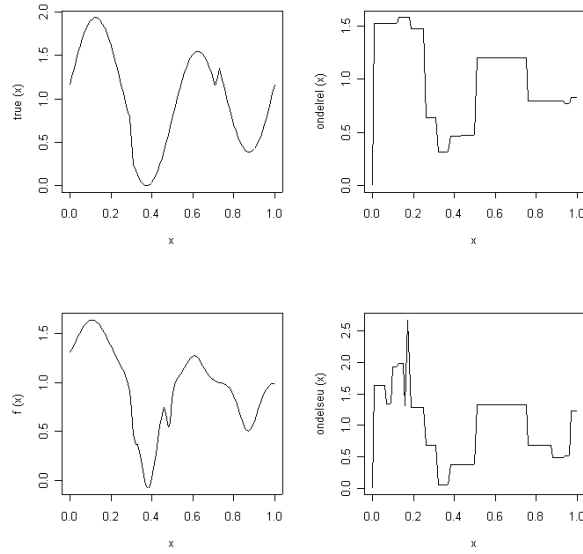


FIGURE 4. Experiment 2, $f = \text{HeaviSine}$ and $\sigma = 0.3$.

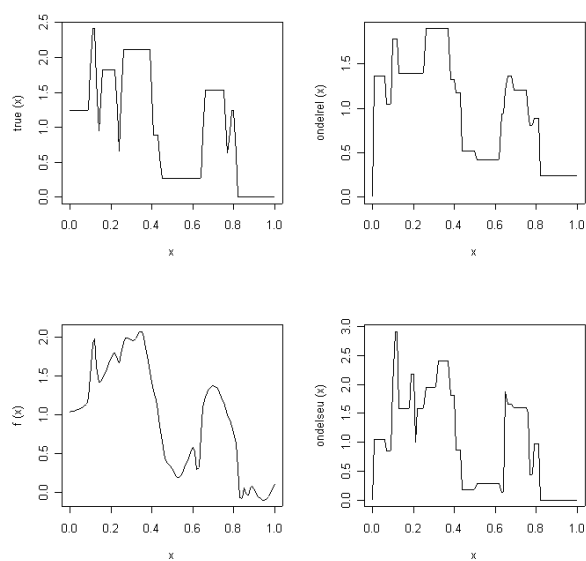


FIGURE 5. Experiment 3, $f = \text{Blocks}$ and $\sigma = 0.3$.

REFERENCES

- [1] P. Alquier, Iterative Feature Selection In Least Square Regression Estimation, *preprint Laboratoire de Probabilités et Modèles Aléatoires*, 2005.
- [2] L. Birgé and P. Massart, An adaptive compression algorithm in Besov spaces, *Constructive Approximation*, 2000, Vol. 16, No. 1.
- [3] B. E. Boser, I. M. Guyon and V. N. Vapnik, A training algorithm for optimal margin classifiers. In D. Haussler, editor, *Proceedings of the 5th Annual ACM Workshop on Computational Learning Theory*, pages 144-152. ACM Press, 1992.
- [4] O. Catoni, Statistical learning theory and stochastic optimization, *Lecture notes, Saint-Flour summer school on Probability Theory*, 2001, Springer.
- [5] O. Catoni, PAC-Bayesian Inductive and Transductive Learning, manuscript, 2006.
- [6] O. Catoni, A PAC-Bayesian approach to adaptive classification, *preprint Laboratoire de Probabilités et Modèles Aléatoires*, 2003.
- [7] A. Cohen, Wavelet methods in numerical analysis, in *Handbook of numerical analysis*, vol. VII, pages 417-711, North-Holland, Amsterdam, 2000.
- [8] D. L. Donoho and I. M. Johnstone, Ideal Spatial Adaptation by Wavelets, *Biometrika*, Vol. 81, No. 3 (Aug., 1994), 425-455.
- [9] D. L. Donoho, I. M. Johnstone, G. Kerkycharian and D. Picard, Density Estimation by Wavelet Thresholding, *Annals of Statistics*, 1996, 24: 508-539.
- [10] W. Hrdle, G. Kerkycharian, D. Picard and A. B. Tsybakov, *Wavelets, Approximations and Statistical Applications*, 1998, Lecture Notes in Statistics, Springer.
- [11] D. Panchenko, Symmetrization Approach to Concentration Inequalities for Empirical Processes, *The Annals Of Probability*, Vol. 31, No. 4 (2003), 2068-2081.
- [12] R Development Core Team, R: A Language And Environment For Statistical Computing, *R Foundation For Statistical Computing*, Vienna, Austria, 2004. URL <http://www.R-project.org>.
- [13] M. Seeger, PAC-Bayesian Generalization Error Bounds for Gaussian Process Classification, *Journal of Machine Learning Research* **3** (2002), 233-269.
- [14] A. B. Tsybakov, *Introduction à l'estimation non-paramétrique*, 2004, Mathématiques et Applications, Springer.
- [15] V. N. Vapnik, *The nature of statistical learning theory*, 1998, Springer Verlag.

LABORATOIRE DE PROBABILITS ET MODLES ALATOIRES, AND LABORATOIRE DE STATISTIQUE,
CREST, 3, AVENUE PIERRE LAROUSSE, 92240 MALAKOFF, FRANCE.

URL: <http://www.crest.fr/pageperso/alquier/alquier.htm>

E-mail address: alquier@ensae.fr