



On the layered nearest neighbour estimate, the bagged nearest neighbour estimate and the random forest method in regression and classification

G rard Biau^{a,b,*}, Luc Devroye^c

^a LSTA & LPMA, Universit  Pierre et Marie Curie—Paris VI, Bo te 158, Tour 15–25, 2 me  tage, 4 place Jussieu, 75252 Paris Cedex 05, France

^b DMA, Ecole Normale Sup rieure, 45 rue d'Ulm, 75230 Paris Cedex 05, France

^c School of Computer Science, McGill University, Montreal, Canada H3A 2K6

ARTICLE INFO

Article history:

Received 23 September 2008

Available online 7 July 2010

AMS 2000 subject classifications:

62G05

62G20

Keywords:

Regression estimation

Layered nearest neighbours

One nearest neighbour estimate

Bagging

Random forests

ABSTRACT

Let $\mathbf{X}_1, \dots, \mathbf{X}_n$ be identically distributed random vectors in $\mathbb{R}^d$, independently drawn according to some probability density. An observation $\mathbf{X}_i$ is said to be a layered nearest neighbour (LNN) of a point $\mathbf{x}$ if the hyperrectangle defined by $\mathbf{x}$ and $\mathbf{X}_i$ contains no other data points. We first establish consistency results on $L_n(\mathbf{x})$, the number of LNN of $\mathbf{x}$. Then, given a sample $(\mathbf{X}, Y), (\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)$ of independent identically distributed random vectors from $\mathbb{R}^d \times \mathbb{R}$, one may estimate the regression function $r(\mathbf{x}) = \mathbb{E}[Y | \mathbf{X} = \mathbf{x}]$ by the LNN estimate $r_n(\mathbf{x})$, defined as an average over the Y_i 's corresponding to those $\mathbf{X}_i$ which are LNN of $\mathbf{x}$. Under mild conditions on r , we establish the consistency of $\mathbb{E}|r_n(\mathbf{x}) - r(\mathbf{x})|^p$ towards 0 as $n \rightarrow \infty$, for almost all $\mathbf{x}$ and all $p \geq 1$, and discuss the links between r_n and the random forest estimates of Breiman (2001) [8]. We finally show the universal consistency of the bagged (bootstrap-aggregated) nearest neighbour method for regression and classification.

  2010 Elsevier Inc. All rights reserved.

1. Introduction

Let $\mathcal{D}_n = \{\mathbf{X}_1, \dots, \mathbf{X}_n\}$ be a sample of independent and identically distributed (i.i.d.) random vectors in $\mathbb{R}^d$, $d \geq 2$. An observation $\mathbf{X}_i$ is said to be a *layered nearest neighbour* (LNN) of a target point $\mathbf{x}$ if the hyperrectangle defined by $\mathbf{x}$ and $\mathbf{X}_i$ contains no other data points. As illustrated in Fig. 1, the number of LNNs of $\mathbf{x}$ is typically larger than one and depends on the number and configuration of the sample points.

The LNN concept, which is briefly discussed in the monograph [13, Chapter 11, Problem 11.6], has strong connections with the notions of *dominance* and *maxima* in random vectors. Recall that a point $\mathbf{X}_i = (X_{i1}, \dots, X_{id})$ is said to be dominated by $\mathbf{X}_j$ if $X_{ik} \leq X_{jk}$ for all $k = 1, \dots, d$, and a point $\mathbf{X}_i$ is called a maximum of $\mathcal{D}_n$ if none of the other points dominates. One can distinguish between *dominance* ($\mathbf{X}_{ik} \leq \mathbf{X}_{jk}$ for all k), *strong dominance* (at least one inequality is strict) and *strict dominance* (all inequalities are strict). The actual kind of dominance will not matter in this paper because we will assume throughout that the common distribution of the data has a density, so that equality of coordinates happens with zero probability. The study of the maxima of a set of points was initiated by Barndorff-Nielsen and Sobel [4] as an attempt to describe the boundary of a set of random points in $\mathbb{R}^d$. Dominance deals with the natural order relations for multivariate observations. Due to its close relationship with the convex hull, dominance is important in computational geometry, pattern classification, graphics, economics and data analysis. The reader is referred to Bai et al. [2] for more information and references.

* Corresponding author at: LSTA & LPMA, Universit  Pierre et Marie Curie—Paris VI, Bo te 158, Tour 15–25, 2 me  tage, 4 place Jussieu, 75252 Paris Cedex 05, France.

E-mail addresses: gerard.biau@upmc.fr (G. Biau), lucdevroye@gmail.com (L. Devroye).

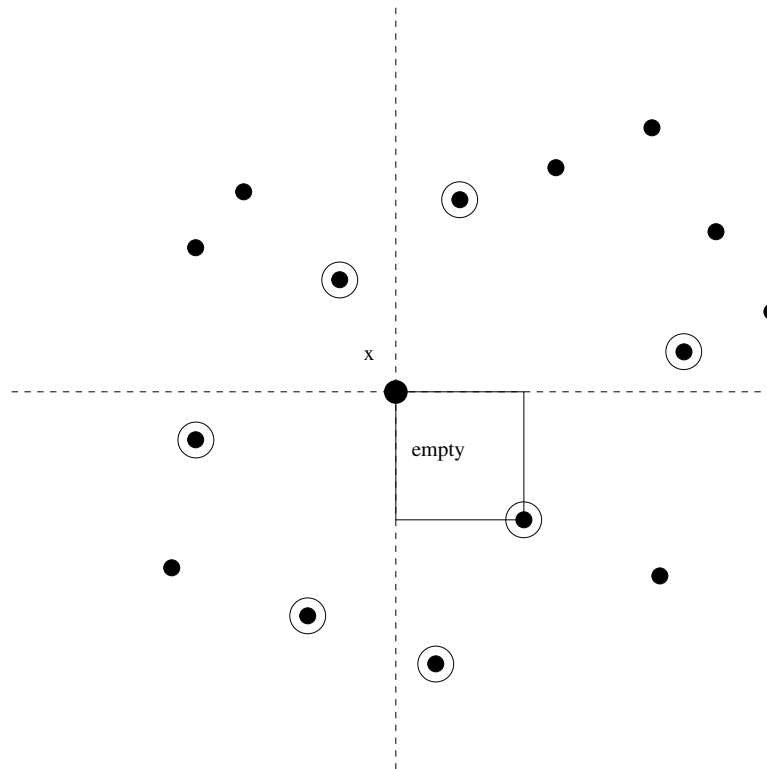


Fig. 1. The layered nearest neighbours (LNN) of a point $\mathbf{x}$.

Denote by L_n the number of maxima in the set $\mathcal{D}_n$. Under the assumption that the observations are independently and uniformly distributed over $(0, 1)^d$, a lot is known about the statistical properties of L_n [4,3,2]. For example, it can be shown that

$$\mathbb{E}L_n = \frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}((\log n)^{d-2}),$$

and

$$\frac{(d-1)!L_n}{(\log n)^{d-1}} \rightarrow 1 \quad \text{in probability,}$$

as $n \rightarrow \infty$ and $d \geq 2$ is fixed. From this, one deduces that when the random vectors $\mathbf{X}_1, \dots, \mathbf{X}_n$ are independently and uniformly distributed over $(0, 1)^d$, then the number of LNNs of any point $\mathbf{x}$ in $(0, 1)^d$, denoted hereafter by $L_n(\mathbf{x})$, satisfies

$$\mathbb{E}L_n(\mathbf{x}) = \frac{2^d (\log n)^{d-1}}{(d-1)!} + \mathcal{O}((\log n)^{d-2})$$

and

$$\frac{(d-1)!L_n(\mathbf{x})}{2^d (\log n)^{d-1}} \rightarrow 1 \quad \text{in probability as } n \rightarrow \infty.$$

Here, the extra factor represents the contribution of the 2^d quadrants surrounding the point $\mathbf{x}$.

On the other hand, to the best of our knowledge, little if nothing is known about the behaviour of $L_n(\mathbf{x})$ under the much more general assumption that the sample points are distributed according to some arbitrary (i.e., non-necessarily uniform) density. A quick inspection reveals that the analysis of this important case is non-trivial and that it may not be readily deduced from the above-mentioned results. Thus, the first objective of this paper is to fill the gap and to offer consistency results about $L_n(\mathbf{x})$ when $\mathbf{X}_1, \dots, \mathbf{X}_n$ are independently drawn according to some arbitrary density f . This will be the topic of Section 2.

Next, we wish to emphasise that the LNN concept also has important statistical consequences. To formalise this idea, assume that we are given a sequence $(\mathbf{X}, Y), (\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)$ of i.i.d. $\mathbb{R}^d \times \mathbb{R}$ -valued random variables with $\mathbb{E}|Y| < \infty$. Then, denoting by $\mathcal{L}_n(\mathbf{x})$ the set of LNNs of $\mathbf{x} \in \mathbb{R}^d$, the regression function $r(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$ may be estimated by

$$r_n(\mathbf{x}) = \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n Y_i \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]}.$$

(Note that $L_n(\mathbf{x}) \geq 1$, so that the division makes sense). In other words, the estimate $r_n(\mathbf{x})$ just outputs the average of the Y_i 's associated with the LNN of the target point $\mathbf{x}$.

The interest of studying the LNN regression estimate r_n , which was first mentioned in [13], is threefold. First, we observe that this estimate has no smoothing parameter, a somewhat unusual situation in nonparametric estimation. Second, it is scale-invariant, which is clearly a desirable feature when the components of the vector represent physically different quantities. And third, it turns out that r_n is intimately related to the *random forests* classification and regression estimates, which were defined by Breiman in [8].

Breiman [8] takes data $(\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)$ and partitions $\mathbb{R}^d$ randomly into “pure” rectangles, i.e., rectangles that each contain one data point. If $A(\mathbf{X})$ is the rectangle to which $\mathbf{X}$ belongs, then $\mathbf{X}$ votes “ Y_i ”, where $\mathbf{X}_i$ is the unique data point in $A(\mathbf{X})$. Breiman repeats such voting and call the principle “random forests”. Classification is done by a majority vote. Regression is done by averaging all Y_i 's. Observe that each voting $\mathbf{X}_i$ is a LNN of $\mathbf{X}$, so that random forests lead to a weighted LNN estimate. In contrast, the estimate r_n above assigns uniform weights. Biau et al. [5] point out that the random forest method is not universally consistent, but the question of consistency remains open when $\mathbf{X}$ is assumed to have a density.

This paper is indeed concerned with minimal conditions of convergence. We say that a regression function estimate r_n is L_p -consistent ($p > 0$) if $\mathbb{E}|r_n(\mathbf{X}) - r(\mathbf{X})|^p \rightarrow 0$, as $n \rightarrow \infty$. It is universally L_p -consistent if this property is true for all distributions of $(\mathbf{X}, Y)$ with $\mathbb{E}|Y|^p < \infty$. Universal consistency was the driving theme of the monograph [12], and we try as much as possible to adhere to the style and notation of that text.

In classification, we have $Y \in \{0, 1\}$, and construct a $\{0, 1\}$ -valued estimate $g_n(\mathbf{x})$ of Y . This is related to regression function estimation since one could use a regression function estimate $r_n(\mathbf{x})$ of $r(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$, and set

$$g_n(\mathbf{x}) = \mathbf{1}_{[r_n(\mathbf{x}) \geq 1/2]}. \quad (1)$$

That estimate has the property that if $\mathbb{E}|r_n(\mathbf{X}) - r(\mathbf{X})| \rightarrow 0$ as $n \rightarrow \infty$, then

$$\mathbb{P}(g_n(\mathbf{X}) \neq Y) \rightarrow \inf_{g: \mathbb{R}^d \rightarrow \{0, 1\}} \mathbb{P}(g(\mathbf{X}) \neq Y),$$

a property which is called Bayes risk consistency (see [13]). It is universally Bayes risk consistent if this property is true for all distributions of $(\mathbf{X}, Y)$.

Random forests are some of the most successful ensemble methods that exhibit performance on the level of boosting and support vector machines. Fast and robust to noise, random forests do not overfit and offer possibilities for explanation and visualisation of the input, such as variable selection. Moreover, random forests have been shown to give excellent performance on a number of practical problems and are among the most accurate general-purpose regression methods available.

An important attempt to investigate the driving forces behind the consistency of random forests is due to Lin and Jeon [18], who show that a forest can be seen as a weighted LNN regression estimate and argue that the LNN approach provides an interesting data-dependent way of measuring proximities between observations.

As a new step towards understanding random forests, we study the consistency of the (uniformly weighted) LNN regression estimate r_n and thoroughly discuss the links between r_n and the random forest estimates of Breiman [8] (Section 3). We finally show in Section 4 the universal consistency of the bagged (bootstrap-aggregated) nearest neighbour method for regression and classification. Proofs of the most technical results are gathered in Section 5.

2. Some consistency properties of the LNN

Throughout this section, we let $\mathcal{D}_n = \{\mathbf{X}_1, \dots, \mathbf{X}_n\}$ be $\mathbb{R}^d$ -valued ($d \geq 2$) independent random variables, identically distributed according to some density f with respect to the Lebesgue measure λ . For any $\mathbf{x} \in \mathbb{R}^d$, we denote by $\mathcal{L}_n(\mathbf{x})$ the set of layered nearest neighbours (LNN) of $\mathbf{x}$ in $\mathcal{D}_n$, and we let $L_n(\mathbf{x})$ be the cardinality of $\mathcal{L}_n(\mathbf{x})$ (i.e., $L_n(\mathbf{x}) = |\mathcal{L}_n(\mathbf{x})|$). Finally, we denote the probability measure associated to f by μ .

The following two basic consistency theorems are proved in Section 5:

Theorem 1. Assume that $\mathbf{X}_1$ has a density with respect to the Lebesgue measure on $\mathbb{R}^d$. Then, for μ -almost all $\mathbf{x} \in \mathbb{R}^d$, one has

$$L_n(\mathbf{x}) \rightarrow \infty \quad \text{in probability as } n \rightarrow \infty.$$

Theorem 2. Assume that $\mathbf{X}_1$ has a density f with respect to the Lebesgue measure on $\mathbb{R}^d$ and that f is λ -almost everywhere continuous. Then

$$\frac{(d-1)! \mathbb{E} L_n(\mathbf{x})}{2^d (\log n)^{d-1}} \rightarrow 1 \quad \text{as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x} \in \mathbb{R}^d$.

The two theorems above are not universal, since we assume that $\mathbf{X}$ has a density. In fact, no universal consistency result is possible. To see this, let $\mathbf{X}$ be $\mathbb{R}^2$ -valued uniformly distributed on the diagonal $D = \{\mathbf{x} = (x_1, x_2) : 0 \leq x_1 \leq 1, x_2 = x_1\}$. Then each $\mathbf{x}$ on the diagonal has almost surely at most 2 LNN and $L_n(\mathbf{x})$ does not converge to infinity.

Finally, the question can be asked what happens when the dimension depends upon n . In this case, for each n , one has in fact another data distribution—statisticians call this a triangular situation. Negative minimax style results imply that for any sequence of estimates, there exists a sequence of distributions such that the error (measured in some sense) does not tend to zero, even under restrictive assumptions on the smoothness of the distribution. Devroye et al. [13] have some theorems like this for discrimination. One is thus forced to make assumptions on the sequence of distributions, e.g., by only considering extensions of distributions of smaller dimensions, i.e., by nesting. For nested sequences, the convergence theorems should be looked at again in general—indeed, there are virtually no universal consistency results on this topic available today.

3. LNN regression estimation

3.1. Consistency

Denote by $(\mathbf{X}, Y), (\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)$ i.i.d. random vectors of $\mathbb{R}^d \times \mathbb{R}$, and let $\mathcal{D}_n$ be the set of data defined by

$$\mathcal{D}_n = \{(\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)\}.$$

In this section we will assume that $|Y| \leq \gamma < \infty$ and that $\mathbf{X}$ has a density. We consider the general regression function estimation problem, where one wants to use the data $\mathcal{D}_n$ in order to construct an estimate $r_n : \mathbb{R}^d \rightarrow \mathbb{R}$ of the regression function $r(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$. Here $r_n(\mathbf{x}) = r_n(\mathbf{x}, \mathcal{D}_n)$ is a measurable function of $\mathbf{x}$ and the data. For simplicity, we will omit $\mathcal{D}_n$ in the notation and write $r_n(\mathbf{x})$ instead of $r_n(\mathbf{x}, \mathcal{D}_n)$.

As in Section 2, for fixed $\mathbf{x} \in \mathbb{R}^d$, we denote by $\mathcal{L}_n(\mathbf{x})$ the LNN of $\mathbf{x}$ in the sample $\{\mathbf{X}_1, \dots, \mathbf{X}_n\}$ and let $L_n(\mathbf{x})$ be the cardinality of $\mathcal{L}_n(\mathbf{x})$ (i.e., $L_n(\mathbf{x}) = |\mathcal{L}_n(\mathbf{x})|$ —note that $L_n(\mathbf{x}) \geq 1$). As stated in the introduction, we will be concerned in this section with the consistency properties of the LNN regression estimate, which is defined by

$$r_n(\mathbf{x}) = \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n Y_i \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]}.$$

Our main result is the following theorem:

Theorem 3 (Pointwise L_p -Consistency). Assume that $\mathbf{X}$ has a density with respect to the Lebesgue measure on $\mathbb{R}^d$, that Y is bounded and that the regression function r is λ -almost everywhere continuous. Then, for μ -almost all $\mathbf{x} \in \mathbb{R}^d$ and all $p \geq 1$,

$$\mathbb{E} |r_n(\mathbf{x}) - r(\mathbf{x})|^p \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

The following corollary is a consequence of Theorem 3 and the dominated convergence theorem.

Theorem 4 (Global L_p -Consistency). Assume that $\mathbf{X}$ has a density with respect to the Lebesgue measure on $\mathbb{R}^d$, that Y is bounded and that the regression function r is λ -almost everywhere continuous. Then, for all $p \geq 1$,

$$\mathbb{E} |r_n(\mathbf{X}) - r(\mathbf{X})|^p \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

The theorems above are not universal—indeed, we assume that r is λ -almost everywhere continuous and that $\mathbf{X}$ has a density. It is noteworthy that no universal consistency result is possible and that the condition that a density exists, while light, cannot be omitted. Indeed, the LNN estimate is not convergent in general without it. For $d = 2$, consider a distribution that puts all its mass smoothly on the diagonal $x_2 = x_1$, let $r(x_1, x_2) = 0$ and let Y be independent of $\mathbf{X}$ taking values $+1$ and -1 with equal probability. Then the LNN estimate averages over its two nearest neighbours on the diagonal, and thus at almost all $\mathbf{x}$ between the extremes of the support, the LNN estimate is distributed as Z where $Z \in \{-1, 0, 1\}$ and the probabilities of the atoms of Z are $1/4$, $1/2$ and $1/4$. Convergence to zero is not possible.

In a second, even more striking example for $d = 2$, let X_1 be uniform on $[-1, 1]$, let $X_2 = 1/X_1$, let $r(x_1, x_2) = \text{sign}(x_1)$, and let $Y = r(\mathbf{X})$ with probability one. Then at almost all points (with respect to the distribution of $\mathbf{X}$), the LNN estimate averages over about half the points in the opposite quadrant, so for $\mathbf{x}$ in quadrant 1, the estimate converges to -1 almost surely, and for $\mathbf{x}$ in quadrant 3, it converges to $+1$ almost surely. The error thus converges to 2 almost surely at almost all $\mathbf{x}$. In other words, singular continuity of the measure of $\mathbf{X}$ causes havoc.

On the positive side, the results do not impose any condition on the density. They are also scale-free, i.e., the estimate does not change when all coordinates of $\mathbf{X}$ are transformed in a strictly monotone manner. In particular, one can without loss of generality assume that $\mathbf{X}$ is supported on $[0, 1]^d$.

Thus, in particular, since (1) is equivalent to taking a majority vote over LNN, we have Bayes risk consistency whenever r is λ -almost everywhere continuous and $\mathbf{X}$ has a density. This partially solves an exercise in [13].

In view of Theorem 2, averaging in the LNN is never over more than $\mathcal{O}((\log n)^{d-1})$ elements. One cannot expect a great rate of convergence for these estimates. The same is true, mutatis mutandis, for Breiman's random forests because averaging is over a subset of size $\mathcal{O}((\log n)^{d-1})$. However, one can hope to improve the averaging rate by the judicious use of subsampling in bagging (bootstrap-aggregating). Bagging, which was suggested by Breiman in [6], is a simple way of randomising and averaging predictors in order to improve their performance. In bagging, randomisation is achieved by generating many bootstrap samples from the original data set. This is illustrated in the next section on 1-nearest neighbour bagging.

3.2. Random forests and LNN

As stated in the introduction, a random forest is a tree-ensemble learning algorithm, where each tree depends on the values of a random vector sampled independently and with the same distribution for all trees. Thus, a random forest consists of many decision trees and outputs the average of the decisions provided by individual trees. Random forests have been shown to give excellent performance on a number of practical problems. They work fast, generally exhibit a substantial performance improvement over single tree algorithms such as CART, and yield generalisation error rates that compare favourably to traditional statistical methods. In fact, random forests are among the most accurate general-purpose learning algorithms available [Breiman 8]).

Algorithms for inducing a random forest were first developed by Breiman and Cutler, and “Random Forests” is their trademark. The web page <http://www.stat.berkeley.edu/users/breiman/RandomForests> provides a collection of downloadable technical reports, and gives an overview of random forests as well as comments on the features of the method.

Following Biau et al. [5], who study consistency of various versions of random forests and other randomised ensemble classifiers, a regression forest may be modelled as follows. Assume that $\Theta_1, \dots, \Theta_m$ are i.i.d. draws of some randomising variable Θ , independent of the sample. Then, a random forest is a collection of m randomised trees $t_1(\mathbf{x}, \Theta_1, \mathcal{D}_n), \dots, t_m(\mathbf{x}, \Theta_m, \mathcal{D}_n)$, which are finally combined to form the aggregated regression estimate

$$r_n(\mathbf{x}) = \frac{1}{m} \sum_{j=1}^m t_j(\mathbf{x}, \Theta_j, \mathcal{D}_n).$$

The randomising variable Θ is used to determine how the successive cuts are performed when building the tree, such as selection of the coordinate to split and position of the split. In the model we have in mind, each individual randomised tree $t_j(\mathbf{x}, \Theta_j, \mathcal{D}_n)$ is typically constructed without pruning, that is, the tree building process continues until each terminal node contains no more than k data points, where k is some prespecified positive integer. Different random forests differ in how randomness is introduced in the tree building process, ranging from extreme random splitting strategies (Breiman [7], Cutler and Zhao [11]) to more involved data-dependent strategies (Amit and Geman [1], Breiman [8], Dietterich [14]). However, as pointed out by Lin and Jeon [18], no matter what splitting strategy is used, if the nodes of the individual trees define rectangular cells, then a random forest with $k = 1$ can be viewed as a *weighted* LNN regression estimate (see also Breiman [9]). Besides, if the randomised splitting scheme is independent of the responses $Y_1, \dots, Y_n$ —such a scheme is called *non-adaptive* in [18]—then so are the weights. One example of such a scheme is the purely random splitting where, for each internal node, we randomly choose a variable to split on, and the split point is chosen uniformly at random over all possible split points on that variable. Thus, for such non-adaptive strategies,

$$r_n(\mathbf{x}) = \sum_{i=1}^n Y_i W_{ni}(\mathbf{x}),$$

where the weights $(W_{n1}(\mathbf{x}), \dots, W_{nn}(\mathbf{x}))$ are nonnegative Borel measurable functions of the variables $\mathbf{x}, \mathbf{X}_1, \dots, \mathbf{X}_n, \Theta_1, \dots, \Theta_m$, and such that $W_{ni}(\mathbf{x}) = 0$ if $\mathbf{X}_i \notin \mathcal{L}_n(\mathbf{x})$ and

$$\sum_{i=1}^n W_{ni}(\mathbf{x}) = \sum_{i=1}^n W_{ni}(\mathbf{x}) \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} = 1.$$

The next proposition states a lower bound on the rate of convergence of the mean squared error of a random forest with non-adaptive splitting scheme. In this proposition, the symbol $\mathbb{V}$ denotes variance and $\mathbb{E}$ denotes expectation with respect to $\mathbf{X}_1, \dots, \mathbf{X}_n$ and $\Theta_1, \dots, \Theta_m$. Proof of the result is due to Lin and Jeon (see [18, Theorem 3, page 581]) and is given here for completeness.

Proposition 1. For any $\mathbf{x} \in \mathbb{R}^d$, assume that $\sigma^2 = \mathbb{V}[Y|\mathbf{X} = \mathbf{x}]$ does not depend upon $\mathbf{x}$. Then

$$\mathbb{E}[r_n(\mathbf{x}) - r(\mathbf{x})]^2 \geq \frac{\sigma^2}{\mathbb{E}L_n(\mathbf{x})}.$$

Proof of Proposition 1. We may write, using the independence of $\mathcal{D}_n$ and $\Theta_1, \dots, \Theta_m$,

$$\begin{aligned} \mathbb{E}[r_n(\mathbf{x}) - r(\mathbf{x})]^2 &\geq \mathbb{E}[\mathbb{V}[r_n(\mathbf{x})|\mathbf{X}_1, \dots, \mathbf{X}_n, \Theta_1, \dots, \Theta_m]] \\ &= \mathbb{E}\left[\sum_{i=1}^n W_{ni}^2(\mathbf{x}) \mathbb{V}[Y_i|\mathbf{X}_1, \dots, \mathbf{X}_n, \Theta_1, \dots, \Theta_m]\right] \\ &= \mathbb{E}\left[\sum_{i=1}^n W_{ni}^2(\mathbf{x}) \mathbb{V}[Y_i|\mathbf{X}_i]\right] \end{aligned}$$

$$\begin{aligned}
&= \sigma^2 \mathbb{E} \left[\sum_{i=1}^n W_{ni}^2(\mathbf{x}) \right] \\
&\geq \sigma^2 \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \left(\sum_{i=1}^n W_{ni}(\mathbf{x}) \right)^2 \right] \quad (\text{by the Cauchy-Schwarz inequality}) \\
&= \sigma^2 \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \right],
\end{aligned}$$

where, in the last equality, we used the fact that $\sum_{i=1}^n W_{ni}(\mathbf{x}) = 1$. The conclusion follows from Jensen's inequality. $\square$

Proposition 1 is thrown in here because we know that

$$\mathbb{E} L_n(\mathbf{x}) \sim \frac{2^d (\log n)^{d-1}}{(d-1)!}$$

at μ -almost all $\mathbf{x}$, when f is λ -almost everywhere continuous ([Theorem 2](#)). Thus, under this additional condition on f , at an $\mathbf{x}$ for which [Theorem 2](#) is valid, we have

$$\begin{aligned}
\mathbb{E} [r_n(\mathbf{x}) - r(\mathbf{x})]^2 &\geq \frac{\sigma^2}{\mathbb{E} L_n(\mathbf{x})} \\
&\sim \frac{\sigma^2 (d-1)!}{2^d (\log n)^{d-1}},
\end{aligned}$$

which is rather slow as a function of n .

As mentioned above, there are two related methods to possibly get a better rate of convergence:

- (i) One can modify the splitting method and stop as soon as a future rectangle split would cause a sub-rectangle to have fewer than k points. In this manner, if $k \rightarrow \infty$, $k/n \rightarrow 0$, one can obtain consistent regression function estimates and classifiers with variances of errors that are of the order $1/[k(\log n)^{d-1}]$. This approach has been thoroughly explored in the uniform case by Lin and Jeon [18], who call it k -potential nearest neighbours (k -PNN, see also Breiman [9]). In a sense, this generalises the classical k -nearest neighbour (k -NN) approach (Györfi et al. [17, Chapter 6]).
- (ii) As suggested by Breiman [6], one could resort to bagging and randomise using small random subsamples. In the next section, we illustrate how this can be done for the 1-NN rule of Fix and Hodges [16] (see also Cover and Hart [10]), thereby extending previous results of [5]. A random subsample of size k is drawn, and the method is repeated m times. The regression estimate takes the average over the mY -values corresponding to the nearest neighbours. In classification, a majority vote is taken. It is shown that for appropriate k and m , this 1-NN bagging is universally consistent, and indeed, that it corresponds to a weighted NN rule, roughly speaking, with geometrically decreasing weights (for more on weighted NN rules, see Stone [23], Devroye [12] or Györfi et al. [17]). Because of this equivalence, one can optimise using standard bias/variance trade-off methods, such as used, e.g., in [17].

4. The bagged 1-NN rule

Breiman's bagging principle has a simple application in the context of nearest neighbour methods. We proceed as follows, via a randomised basic regression estimate $r_{n,k}$ in which $1 \leq k \leq n$ is a parameter. The predictor $r_{n,k}$ is the 1-NN rule for a random sample $\mathcal{S}_n$ drawn with (without) replacement from $\{\mathbf{X}_1, \dots, \mathbf{X}_n\}$, with $|\mathcal{S}_n| = k$. Clearly, $r_{n,k}$ is not generally universally consistent.

We apply bagging, that is, we repeat the random sampling m times, and take the average of the individual outcomes. Formally, if $Z_j = r_{n,k}(\mathbf{x})$ is the prediction in the j -th round of bagging, we let the bagged regression estimate r_n^* be defined as

$$r_n^*(\mathbf{x}) = \frac{1}{m} \sum_{j=1}^m Z_j,$$

where $Z_1, \dots, Z_m$ are the outcomes in the individual rounds. In the context of classification, $Z_j \in \{0, 1\}$, and we classify $\mathbf{x}$ as being in class 1 if $r_n^*(\mathbf{x}) \geq 1/2$, that is

$$\sum_{j=1}^m \mathbf{1}_{[Z_j=1]} \geq \sum_{j=1}^m \mathbf{1}_{[Z_j=0]}.$$

The corresponding bagged classifier is denoted by g_n^* .

Theorem 5. If $m \rightarrow \infty$ (or $m = \infty$), $k \rightarrow \infty$ and $k/n \rightarrow 0$, then r_n^* is universally L_p -consistent for all $p \geq 1$.

Corollary 1. If $m \rightarrow \infty$ (or $m = \infty$), $k \rightarrow \infty$ and $k/n \rightarrow 0$, then g_n^* is universally Bayes risk consistent.

Remark. In the theorem, the fact that sampling was done with/without replacement is irrelevant.

Before proving Theorem 5, recall that if we let $V_{n1} \geq V_{n2} \geq \dots \geq V_{nn} \geq 0$ denote weights that sum to one, and $V_{n1} \rightarrow 0$, $\sum_{i>\varepsilon n} V_{ni} \rightarrow 0$ for all $\varepsilon > 0$ as $n \rightarrow \infty$, then the regression estimate

$$\sum_{i=1}^n V_{ni} Y_{(i)}(\mathbf{x}),$$

with $(\mathbf{X}_{(1)}(\mathbf{x}), Y_{(1)}(\mathbf{x})), \dots, (\mathbf{X}_{(n)}(\mathbf{x}), Y_{(n)}(\mathbf{x}))$ the reordering of the data such that

$$\|\mathbf{x} - \mathbf{X}_{(1)}(\mathbf{x})\| \leq \dots \leq \|\mathbf{x} - \mathbf{X}_{(n)}(\mathbf{x})\|$$

is called the *weighted nearest neighbour regression estimate*. It is universally L_p -consistent for all $p \geq 1$ (Stone [23], and Problems 11.7, 11.8 of Devroye et al. [13]). In the sequel, to shorten notation, we omit the index n in the weights and write, for instance, V_1 instead of V_{n1} .

Proof of Theorem 5. We first observe that if $m = \infty$, r_n^* is in fact a weighted nearest neighbour estimate with

$$V_i = \mathbb{P}(i\text{-th nearest neighbour of } \mathbf{x} \text{ is chosen in a random selection}).$$

To avoid trouble, we have a unique way of breaking distance ties, that is, any tie is broken by using indices to declare a winner. Then, a moment's thought shows that for the "without replacement" sampling, V_i is hypergeometric:

$$V_i = \begin{cases} \frac{\binom{n-i}{k-1}}{\binom{n}{k}}, & i \leq n-k+1 \\ 0, & i > n-k+1. \end{cases}$$

We have

$$\begin{aligned} V_i &= \frac{k}{n-k+1} \cdot \frac{n-i}{n} \cdot \frac{n-i-1}{n-1} \dots \frac{n-i-k+2}{n-k+2} \\ &= \frac{k}{n-k+1} \prod_{j=0}^{k-2} \left(1 - \frac{i}{n-j}\right) \\ &\in \left[\frac{k}{n-k+1} \exp\left(\frac{-i(k-1)}{n-k-i+2}\right), \frac{k}{n-k+1} \exp\left(\frac{-i(k-1)}{n}\right) \right], \end{aligned}$$

where we used $\exp(-u/(1-u)) \leq 1-u \leq \exp(-u)$, $0 \leq u < 1$. Clearly, V_i is nonincreasing, with

$$V_1 = \frac{k}{n} \rightarrow 0.$$

Also,

$$\begin{aligned} \sum_{i>\varepsilon n} V_i &\leq \frac{k}{n-k+1} \sum_{i>\varepsilon n} e^{-i(k-1)/n} \\ &\leq \frac{k}{n-k+1} \cdot \frac{e^{-\varepsilon(k-1)}}{(1 - e^{-(k-1)/n})} \\ &\sim e^{-\varepsilon(k-1)} \rightarrow 0 \quad \text{as } k \rightarrow \infty. \end{aligned}$$

For sampling with replacement,

$$\begin{aligned} V_i &= \left(1 - \frac{i-1}{n}\right)^k - \left(1 - \frac{i}{n}\right)^k \\ &= \left(1 - \frac{i-1}{n}\right)^k \left[1 - \left(1 - \frac{1}{n-i+1}\right)^k\right] \\ &\in \left[e^{-(i-1)k/(n-i+1)} \left[\frac{k}{n-i+1} - \frac{k(k-1)}{2} \left(\frac{1}{n-i+1}\right)^2 \right], e^{-(i-1)k/n} \cdot \frac{k}{n-i+1} \right], \end{aligned}$$

where we used $1 - \alpha u \leq (1 - u)^\alpha \leq 1 - \alpha u + \alpha(\alpha - 1)u^2/2$ for integer $\alpha \geq 1$, $0 \leq u \leq 1$. Again, V_i is nonincreasing, and

$$V_1 = 1 - \left(1 - \frac{1}{n}\right)^k \leq \frac{k}{n} \rightarrow 0.$$

Also

$$\sum_{i > \varepsilon n} V_i = \left(1 - \frac{\lfloor \varepsilon n \rfloor}{n}\right)^k \rightarrow 0$$

since $\varepsilon > 0$ is fixed and $k \rightarrow \infty$.

For $m < \infty$, $m \rightarrow \infty$, the weights of the neighbours are random variables $(W_1, \dots, W_n)$, with $\sum_{i=1}^n W_i = 1$, and, in fact,

$$(W_1, \dots, W_n) \stackrel{\mathcal{L}}{=} \frac{\text{Multinomial}(m; V_1, \dots, V_n)}{m}.$$

We note that this random vector is *independent* of the data!

In the proof of the consistency result below, we use Stone's [23] general consistency theorem for locally weighted average estimates, see also [13, Theorem 6.3]. According to Stone's theorem, consistency holds if the following three conditions are satisfied:

(i)

$$\mathbb{E} \left[\max_{i=1, \dots, n} W_i \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

(ii) For all $\varepsilon > 0$,

$$\mathbb{E} \left[\sum_{i > \varepsilon n} W_i \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

(iii) There is a constant C such that, for every nonnegative measurable function f satisfying $\mathbb{E}f(\mathbf{X}) < \infty$,

$$\mathbb{E} \left[\sum_{i=1}^n W_i f(\mathbf{X}_{(i)}) \right] \leq C \mathbb{E}f(\mathbf{X}).$$

Checking Stone's conditions of convergence requires only minor work. To show (i), note that

$$\begin{aligned} \mathbb{P} \left(\max_{i=1, \dots, n} W_i \geq \varepsilon \right) &\leq \sum_{i=1}^n \mathbb{P}(W_i \geq \varepsilon) \\ &= \sum_{i=1}^n \mathbb{P}(\text{Bin}(m, V_i) \geq m\varepsilon) \\ &= \sum_{i=1}^n \mathbb{P}(\text{Bin}(m, V_i) \geq mV_i + m(\varepsilon - V_i)) \\ &\leq \sum_{i=1}^n \frac{\mathbb{V}[\text{Bin}(m, V_i)]}{(m(\varepsilon - V_i))^2} \quad (\text{by Chebyshev's inequality, for all } n \text{ large enough}) \\ &\leq \frac{\sum_{i=1}^n mV_i}{m^2(\varepsilon - V_1)^2} = \frac{1}{m(\varepsilon - V_1)^2} \rightarrow 0. \end{aligned}$$

Second, for (ii), we set $p = \sum_{i > \varepsilon n} V_i$, and need only show that

$$\mathbb{E}[\text{Bin}(m, p)/m] \rightarrow 0.$$

But this follows from $p \rightarrow 0$. Condition (iii) reduces to

$$\mathbb{E} \left[\sum_{i=1}^n V_i f(\mathbf{X}_{(i)}) \right],$$

which we know is bounded by a constant times $\mathbb{E}f(\mathbf{X})$ for any sequence of nonincreasing nonnegative weights V_i that sum to one (Stone [23], and [13, Chapter 11, Problems 11.7 and 11.8]).

Remark. The bagging weights for $m = \infty$ have been independently exhibited by Steele [22], who also shows on practical examples that substantial reductions in prediction error are possible by bagging the 1-NN estimate.

This concludes the proof. $\square$

5. Proofs

5.1. Proof of Theorem 1

In the sequel, for $\mathbf{x} = (x_1, \dots, x_d)$ and $\varepsilon > 0$, we let the hyperrectangle $\mathcal{R}_\varepsilon(\mathbf{x})$ be defined as $\mathcal{R}_\varepsilon(\mathbf{x}) = [x_1, x_1 + \varepsilon] \times \dots \times [x_d, x_d + \varepsilon]$. The crucial result from real analysis that is needed in the proof of Theorems 1 and 2 is the following (see for instance Wheeden and Zygmund [24]):

Lemma 1. Let g be locally integrable in $\mathbb{R}^d$. Then, for λ -almost all $\mathbf{x}$,

$$\frac{1}{\varepsilon^d} \int_{\mathcal{R}_\varepsilon(\mathbf{x})} |g(\mathbf{y}) - g(\mathbf{x})| d\mathbf{y} \rightarrow 0 \quad \text{as } \varepsilon \rightarrow 0. \quad (2)$$

Thus also, at λ -almost all $\mathbf{x}$,

$$\frac{1}{\varepsilon^d} \int_{\mathcal{R}_\varepsilon(\mathbf{x})} g(\mathbf{y}) d\mathbf{y} \rightarrow f(\mathbf{x}) \quad \text{as } \varepsilon \rightarrow 0. \quad (3)$$

The following useful corollary may be easily deduced from Lemma 1 and the fact that $f(\mathbf{x}) > 0$ for μ -almost all $\mathbf{x}$:

Corollary 2. Assume that $\mathbf{X}_1$ has a density with respect to the Lebesgue measure on $\mathbb{R}^d$. Let (ε_n) be a sequence of positive real numbers such that $\varepsilon_n \rightarrow 0$ and $n\varepsilon_n^d \rightarrow \infty$ as $n \rightarrow \infty$. Then, for μ -almost all $\mathbf{x} \in \mathbb{R}^d$, one has

$$n\mu(\mathcal{R}_{\varepsilon_n}(\mathbf{x})) \rightarrow \infty \quad \text{as } n \rightarrow \infty.$$

Remark. Lemma 1 only describes what happens if $\mathcal{R}_\varepsilon(\mathbf{x})$ is in the positive quadrant of $\mathbf{x}$. Trivially, it also holds for the $2^d - 1$ other quadrants centred at $\mathbf{x}$.

The proof of Theorem 1 uses a coupling argument. Roughly, the idea is to replace the sample $\mathbf{X}_1, \dots, \mathbf{X}_n$ by a sample $\mathbf{Z}_1, \dots, \mathbf{Z}_n$ which is (i) locally uniform around the point $\mathbf{x}$ and (ii) such that the probability of the event $[\mathbf{X}_i = \mathbf{Z}_i, i = 1, \dots, n]^c$ stays under control. This can be achieved via the following optimal coupling lemma (see for example Doeblin [15] or Rachev and Rüschendorf [21]):

Lemma 2. Let f and g be probability densities on $\mathbb{R}^d$. Then there exist random variables $\mathbf{W}$ and $\mathbf{Z}$ with density f and g , respectively, such that

$$\mathbb{P}(\mathbf{W} \neq \mathbf{Z}) = \frac{1}{2} \int_{\mathbb{R}^d} |f(\mathbf{y}) - g(\mathbf{y})| d\mathbf{y}.$$

We are now in a position to prove Theorem 1.

Proof of Theorem 1. Fix $\mathbf{x}$ for which (2) is true, and define the function g_ε as

$$g_\varepsilon(\mathbf{y}) = \begin{cases} \frac{\mu(\mathcal{R}_\varepsilon(\mathbf{x}))}{\varepsilon^d} & \text{if } \mathbf{y} \in \mathcal{R}_\varepsilon(\mathbf{x}) \\ f(\mathbf{y}) & \text{otherwise.} \end{cases}$$

Clearly, g_ε is a probability density on $\mathbb{R}^d$. Moreover,

$$\begin{aligned} \int_{\mathbb{R}^d} |f(\mathbf{y}) - g_\varepsilon(\mathbf{y})| d\mathbf{y} &= \int_{\mathcal{R}_\varepsilon(\mathbf{x})} \left| f(\mathbf{y}) - \frac{1}{\varepsilon^d} \int_{\mathcal{R}_\varepsilon(\mathbf{x})} f(\mathbf{z}) d\mathbf{z} \right| d\mathbf{y} \\ &\leq \int_{\mathcal{R}_\varepsilon(\mathbf{x})} |f(\mathbf{y}) - f(\mathbf{x})| d\mathbf{y} + \varepsilon^d \left| f(\mathbf{x}) - \frac{1}{\varepsilon^d} \int_{\mathcal{R}_\varepsilon(\mathbf{x})} f(\mathbf{z}) d\mathbf{z} \right| \\ &\leq 2 \int_{\mathcal{R}_\varepsilon(\mathbf{x})} |f(\mathbf{y}) - f(\mathbf{x})| d\mathbf{y} \\ &\leq 2\varepsilon^d \Phi(\varepsilon), \end{aligned} \quad (4)$$

where $\Phi(\varepsilon)$ is some nonnegative, nondecreasing function which has limit 0 as ε approaches 0.

According to Lemma 2 and inequality (4), there exist random variables $\mathbf{W}$ and $\mathbf{Z}$ with density f and g_ε , respectively, such that

$$\mathbb{P}(\mathbf{W} \neq \mathbf{Z}) \leq \varepsilon^d \Phi(\varepsilon).$$

Define $\mathbf{W}$ and $\mathbf{Z}$ samples by creating n independent $(\mathbf{W}_1, \mathbf{Z}_1), \dots, (\mathbf{W}_n, \mathbf{Z}_n)$ pairs and assume, without loss of generality, that

$$(\mathbf{X}_1, \dots, \mathbf{X}_n) = (\mathbf{W}_1, \dots, \mathbf{W}_n).$$

Thus, denoting by $\mathcal{E}_n$ the event

$$[\mathbf{X}_i = \mathbf{Z}_i, i = 1, \dots, n],$$

we obtain, by construction of the optimal coupling,

$$\mathbb{P}(\mathcal{E}_n^c) \leq n\varepsilon^d \Phi(\varepsilon). \quad (5)$$

According to technical Lemma 4, there exists a sequence (ε_n) of positive real numbers such that $\varepsilon_n \rightarrow 0$, $n\varepsilon_n^d \rightarrow \infty$ and $n\varepsilon_n^d \Phi(\varepsilon_n) \rightarrow 0$ as $n \rightarrow \infty$. Thus, by choosing such a sequence, according to (5), the probability $\mathbb{P}(\mathcal{E}_n^c)$ can be made as small as desired for all n large enough.

To finish the proof of Theorem 1, denote by $L_{\varepsilon_n}(\mathbf{x})$ (respectively $L'_{\varepsilon_n}(\mathbf{x})$) the number of LNNs of $\mathbf{x}$ in the sample $\{\mathbf{X}_1, \dots, \mathbf{X}_n\}$ (respectively in the sample $\{\mathbf{Z}_1, \dots, \mathbf{Z}_n\}$) falling in $\mathcal{R}_{\varepsilon_n}(\mathbf{x})$. Clearly,

$$L_n(\mathbf{x}) \geq L_{\varepsilon_n}(\mathbf{x}), \quad (6)$$

and, on the event $\mathcal{E}_n$,

$$L_{\varepsilon_n}(\mathbf{x}) = L'_{\varepsilon_n}(\mathbf{x}). \quad (7)$$

By Lemma 5, since $n\varepsilon_n^d \rightarrow \infty$,

$$L'_{\varepsilon_n}(\mathbf{x}) \rightarrow \infty \quad \text{in probability as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x}$. This, together with (5)–(7) concludes the proof of the theorem. $\square$

5.2. Proof of Theorem 2

For $\mathbf{x} = (x_1, \dots, x_d)$ and $\varepsilon > 0$, let $\mathcal{C}_\varepsilon(\mathbf{x})$ be the hypercube $[x_1 - \varepsilon, x_1 + \varepsilon] \times \dots \times [x_d - \varepsilon, x_d + \varepsilon]$. Choose $\mathbf{x}$ in a set of μ -measure 1 such that $f(\mathbf{x}) > 0$, f is continuous at $\mathbf{x}$ and $\mu(\mathcal{C}_\varepsilon(\mathbf{x})) > 0$ for all $\varepsilon > 0$.

Fix $\delta \in (0, f(\mathbf{x}))$. Since f is continuous at $\mathbf{x}$, there exists $\varepsilon > 0$ such that $\mathbf{y} \in \mathcal{C}_\varepsilon(\mathbf{x})$ implies $|f(\mathbf{x}) - f(\mathbf{y})| < \delta$.

Let the hyperrectangle $\mathcal{R}_{(\mathbf{x}, \mathbf{y})}$ be defined by $\mathbf{x}$ and $\mathbf{y}$. We note that

$$\begin{aligned} \mathbb{E}L_n(\mathbf{x}) &= n \int_{\mathbb{R}^d} (1 - \mu(\mathcal{R}_{(\mathbf{x}, \mathbf{y})}))^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &= n \int_{\mathbb{R}^d} \left(1 - \int_{\mathcal{R}_{(\mathbf{x}, \mathbf{y})}} f(\mathbf{z}) d\mathbf{z}\right)^{n-1} f(\mathbf{y}) d\mathbf{y}. \end{aligned}$$

Thus, using the continuity of f at $\mathbf{x}$, we obtain

$$\begin{aligned} \mathbb{E}L_n(\mathbf{x}) &\geq n(f(\mathbf{x}) - \delta) \int_{\mathcal{C}_\varepsilon(\mathbf{x})} (1 - (f(\mathbf{x}) + \delta)\Pi|y_i - x_i|)^{n-1} d\mathbf{y} \\ &= n(f(\mathbf{x}) - \delta) \int_{\mathcal{C}_\varepsilon(\mathbf{0})} (1 - (f(\mathbf{x}) + \delta)\Pi|y_i|)^{n-1} d\mathbf{y} \\ &= n \frac{f(\mathbf{x}) - \delta}{f(\mathbf{x}) + \delta} \int_{\mathcal{C}_{\Delta\varepsilon}(\mathbf{0})} (1 - \Pi|y_i|)^{n-1} d\mathbf{y} \quad (\text{with } \Delta = (f(\mathbf{x}) + \delta)^{1/d}) \\ &= n2^d \frac{f(\mathbf{x}) - \delta}{f(\mathbf{x}) + \delta} \int_{\mathcal{R}_{\Delta\varepsilon}(\mathbf{0})} (1 - \Pi y_i)^{n-1} d\mathbf{y}, \end{aligned}$$

where the last equality follows from a symmetry argument. Thus, using technical Lemma 6, we conclude that

$$\mathbb{E}L_n(\mathbf{x}) \geq 2^d \frac{f(\mathbf{x}) - \delta}{f(\mathbf{x}) + \delta} \left[\frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}_{\Delta\varepsilon}(\log n)^{d-2} \right],$$

where the notation $\mathcal{O}_{\Delta\varepsilon}$ means that the constant in the $\mathcal{O}$ term depends on $\Delta\varepsilon$. Letting $\delta \rightarrow 0$ shows that

$$\liminf_{n \rightarrow \infty} \frac{(d-1)! \mathbb{E} L_n(\mathbf{x})}{2^d (\log n)^{d-1}} \geq 1.$$

To show the opposite inequality, we write, using the continuity of f at $\mathbf{x}$,

$$\begin{aligned} \mathbb{E} L_n(\mathbf{x}) &= n \int_{\mathcal{C}_\varepsilon(\mathbf{x})} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y} + n \int_{\mathbb{R}^d \setminus \mathcal{C}_\varepsilon(\mathbf{x})} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &\leq n 2^d \frac{f(\mathbf{x}) + \delta}{f(\mathbf{x}) - \delta} \int_{\mathcal{R}_{\Delta\varepsilon}(\mathbf{0})} (1 - \Pi y_i)^{n-1} d\mathbf{y} \quad (\text{with } \Delta = (f(\mathbf{x}) - \delta)^{1/d}) \\ &\quad + n \int_{\mathbb{R}^d \setminus \mathcal{C}_\varepsilon(\mathbf{x})} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y}. \end{aligned} \quad (8)$$

By technical Lemma 6, we have

$$n 2^d \frac{f(\mathbf{x}) + \delta}{f(\mathbf{x}) - \delta} \int_{\mathcal{R}_{\Delta\varepsilon}(\mathbf{0})} (1 - \Pi y_i)^{n-1} d\mathbf{y} = 2^d \frac{f(\mathbf{x}) + \delta}{f(\mathbf{x}) - \delta} \left[\frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}_{\Delta\varepsilon}(\log n)^{d-2} \right]. \quad (9)$$

Then, with respect to the second term in (8), we note that

$$\mathbb{R}^d \setminus \mathcal{C}_\varepsilon(\mathbf{x}) = \bigcup_{j=0}^{d-1} \mathcal{C}_j,$$

where, by definition, $\mathcal{C}_j$ denotes the collection of all $\mathbf{y}$ in $\mathbb{R}^d \setminus \mathcal{C}_\varepsilon(\mathbf{x})$ which have *exactly* j coordinates smaller than ε . Observe that, for each $j \in \{0, \dots, d-1\}$,

$$\mathcal{C}_j = \bigcup_{\underline{j}} \mathcal{C}_{\underline{j}},$$

where the index $\underline{j}$ runs over the $\binom{d}{j}$ possible j -uples coordinate choices smaller than ε . Associated to each of these choices is a marginal density of f , that we denote by $f_{\underline{j}}$. For $j \geq 1$, with a slight abuse of notation, we let $\mathcal{C}_\varepsilon(\mathbf{x}_{\underline{j}})$ be the j -dimensional hypercube with centre at the coordinates of $\mathbf{x}$ matching with $\underline{j}$ and side length 2ε . Finally, we choose ε small enough and $\mathbf{x}$ in a set of μ -measure 1 such that each marginal $f_{\underline{j}}$ is bounded over $\mathcal{C}_\varepsilon(\mathbf{x}_{\underline{j}})$ by, say, $\Lambda(\varepsilon)$.

Clearly, for $j = 0$,

$$\begin{aligned} n \int_{\mathcal{C}_0} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y} &= \int_{\mathcal{C}_0} \left(1 - \int_{\mathcal{R}(\mathbf{x}, \mathbf{y})} f(\mathbf{z}) d\mathbf{z} \right)^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &\leq n(1 - (f(\mathbf{x}) - \delta)\varepsilon^d)^{n-1} \int_{\mathcal{C}_0} f(\mathbf{y}) d\mathbf{y} \\ &\leq n(1 - (f(\mathbf{x}) - \delta)\varepsilon^d)^{n-1} \quad (\text{since } f \text{ is a density}) \\ &\leq 1 / [(f(\mathbf{x}) - \delta)\varepsilon^d], \end{aligned}$$

where, in the last inequality, we used the fact that $\sup_{x \in [0,1]} x(1-x)^{n-1} \leq 1/n$. Similarly, for $j \in \{1, \dots, d-1\}$, we may write

$$\begin{aligned} n \int_{\mathcal{C}_j} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y} &= n \sum_{\underline{j}} \int_{\mathcal{C}_{\underline{j}}} (1 - \mu(\mathcal{R}(\mathbf{x}, \mathbf{y})))^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &= n \sum_{\underline{j}} \int_{\mathcal{C}_{\underline{j}}} \left(1 - \int_{\mathcal{R}(\mathbf{x}, \mathbf{y})} f(\mathbf{z}) d\mathbf{z} \right)^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &\leq n \sum_{\underline{j}} \int_{\mathcal{C}_{\underline{j}}} (1 - (f(\mathbf{x}) - \delta)\varepsilon^{d-j} \Pi_{\ell} |y_{\ell} - x_{\ell}|)^{n-1} f(\mathbf{y}) d\mathbf{y}, \end{aligned}$$

where the notation Π_{ℓ} means the product over the j coordinates which are smaller than ε . Thus, integrating the density f over those coordinates which are larger than ε , we obtain

$$\begin{aligned} &\int_{\mathcal{C}_{\underline{j}}} (1 - (f(\mathbf{x}) - \delta)\varepsilon^{d-j} \Pi_{\ell} |y_{\ell} - x_{\ell}|)^{n-1} f(\mathbf{y}) d\mathbf{y} \\ &\leq \int_{\mathcal{C}_\varepsilon(\mathbf{x}_{\underline{j}})} (1 - (f(\mathbf{x}) - \delta)\varepsilon^{d-j} \Pi_{\ell} |y_{\ell} - x_{\ell}|)^{n-1} f_{\underline{j}}(\mathbf{y}_{\underline{j}}) d\mathbf{y}_{\underline{j}}. \end{aligned}$$

Finally using the fact that each marginal f_j is bounded by $\Lambda(\varepsilon)$ in the neighbourhood of $\mathbf{x}$, we obtain

$$\begin{aligned} & \int_{\mathcal{C}_j} (1 - (f(\mathbf{x}) - \delta)\varepsilon^{d-j} \Pi_\ell |y_\ell - x_\ell|)^{n-1} f(\mathbf{y}) d\mathbf{y} \\ & \leq \Lambda(\varepsilon) \int_{[0, \varepsilon]^j} (1 - (f(\mathbf{x}) - \delta)\varepsilon^{d-j} y_1 \dots y_j)^{n-1} dy_1 \dots dy_j \\ & = \frac{\Lambda(\varepsilon)}{(f(\mathbf{x}) - \delta)\varepsilon^{d-j}} \int_{[0, \Delta \varepsilon^{d-j/j}]^j} (1 - y_1 \dots y_j)^{n-1} dy_1 \dots dy_j \quad (\text{with } \Delta = (f(\mathbf{x}) - \delta)^{1/j}). \end{aligned}$$

Therefore, by Lemma 6, for $j \in \{2, \dots, d-1\}$,

$$n \int_{\mathcal{C}_j} (1 - \mu(\mathcal{R}_{(\mathbf{x}, \mathbf{y})}))^{n-1} f(\mathbf{y}) d\mathbf{y} \leq \frac{\binom{d}{j} \Lambda(\varepsilon)}{(f(\mathbf{x}) - \delta)\varepsilon^{d-j}} \left[\frac{(\log n)^{j-1}}{(j-1)!} + \mathcal{O}_{\Delta \varepsilon^{d-j/j}}(\log n)^{j-2} \right],$$

and clearly, for $j = 1$,

$$n \int_{\mathcal{C}_1} (1 - \mu(\mathcal{R}_{(\mathbf{x}, \mathbf{y})}))^{n-1} f(\mathbf{y}) d\mathbf{y} \leq \frac{\Lambda(\varepsilon)}{(f(\mathbf{x}) - \delta)\varepsilon^{d-1}}.$$

Putting all pieces together, we conclude that, for all $j \in \{0, \dots, d-1\}$,

$$\limsup_{n \rightarrow \infty} \frac{n}{(\log n)^{d-1}} \int_{\mathcal{C}_j} (1 - \mu(\mathcal{R}_{(\mathbf{x}, \mathbf{y})}))^{n-1} f(\mathbf{y}) d\mathbf{y} = 0,$$

and, consequently,

$$\limsup_{n \rightarrow \infty} \frac{n}{(\log n)^{d-1}} \int_{\mathbb{R}^d \setminus \mathcal{C}_\varepsilon(\mathbf{x})} (1 - \mu(\mathcal{R}_{(\mathbf{x}, \mathbf{y})}))^{n-1} f(\mathbf{y}) d\mathbf{y} = 0.$$

This, together with inequalities (8)–(9) and letting $\delta \rightarrow 0$ leads to

$$\limsup_{n \rightarrow \infty} \frac{(d-1)! \mathbb{E} L_n(\mathbf{x})}{2^d (\log n)^{d-1}} \leq 1.$$

5.3. Proof of Theorem 3

The elementary result needed to prove Theorem 3 is:

Lemma 3. Assume that $\mathbf{X}$ has a density with respect to the Lebesgue measure on $\mathbb{R}^d$, that Y is bounded and that the regression function r is λ -almost everywhere continuous. Then, for fixed $p \geq 1$,

$$\mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x} \in \mathbb{R}^d$.

Proof of Lemma 3. Recall that $\mathcal{R}_\varepsilon(\mathbf{x}) = [x_1, x_1 + \varepsilon] \times \dots \times [x_d, x_d + \varepsilon]$. We can define $\mathcal{R}_\varepsilon(\mathbf{x}, \ell)$, $\ell = 1, \dots, 2^d$, as $\mathcal{R}_\varepsilon(\mathbf{x})$ for the 2^d quadrants centred at $\mathbf{x}$. We then have $\mathcal{L}_n(\mathbf{x}, \ell)$ and $L_n(\mathbf{x}, \ell) = |\mathcal{L}_n(\mathbf{x}, \ell)|$. Also, on the ℓ -th quadrant, we have the sums

$$S_n(\mathbf{x}, \ell) = \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x}, \ell)]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p.$$

If

$$\mathbb{E} \left[\frac{S_n(\mathbf{x}, \ell)}{L_n(\mathbf{x}, \ell)} \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty \text{ for all } \ell,$$

(with the convention $0 \times \infty = 0$ when $L_n(\mathbf{x}, \ell) = 0$), then

$$\mathbb{E} \left[\frac{\sum_{\ell=1}^{2^d} S_n(\mathbf{x}, \ell)}{\sum_{\ell=1}^{2^d} L_n(\mathbf{x}, \ell)} \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

so that

$$\mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

This follows from the fact that

$$\mathbb{E} \left[\frac{A_1 + \dots + A_k}{B_1 + \dots + B_k} \right] \rightarrow 0$$

if $\mathbb{E}[A_i/B_i] \rightarrow 0$ for all i , where the random variables A_i and B_i are nonnegative and satisfy $A_i \leq cB_i$ for some nonnegative c (again, we use the convention $0 \times \infty = 0$). So, we need only concentrate on the first quadrant.

For arbitrary $\varepsilon > 0$, we have

$$\begin{aligned} & \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \right] \\ &= \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \right] \\ & \quad + \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon(\mathbf{x})]} \right] \\ & \leq 2^p \gamma^p \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \right] + \left[\sup_{\mathbf{z} \in \mathbb{R}^d: \|\mathbf{z} - \mathbf{x}\|_\infty \leq \varepsilon} |r(\mathbf{z}) - r(\mathbf{x})| \right]^p \quad (\text{since } |Y| \leq \gamma). \end{aligned}$$

The rightmost term of the latter inequality tends to 0 as $\varepsilon \rightarrow 0$ at points $\mathbf{x}$ at which r is continuous. Thus, the lemma will be proven if we show that, for fixed $\varepsilon > 0$,

$$\mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

To this aim, denote by N_n the (random) number of sample points falling in $\mathcal{R}_\varepsilon(\mathbf{x})$. For $N_n \geq 1$ and each $r = 1, \dots, d$, let $\mathbf{X}_n^{*(r)} = (X_{n,1}^{*(r)}, \dots, X_{n,d}^{*(r)})$ be the observation in $\mathcal{R}_\varepsilon(\mathbf{x})$ whose r -coordinate is the closest to x_r (note that $\mathbf{X}_n^{*(r)}$ is almost surely unique), and consider the set

$$\mathcal{P}_\varepsilon^{(r)} = [x_1 + \varepsilon, +\infty[\times \dots \times [x_{r-1} + \varepsilon, +\infty[\times [x_r, X_{n,r}^{*(r)}] \times [x_{r+1} + \varepsilon, +\infty[\times \dots \times [x_d + \varepsilon, +\infty[$$

(see Fig. 2 for an illustration in dimension 2).

Take finally

$$\mathcal{P}_\varepsilon = \bigcup_{r=1}^d \mathcal{P}_\varepsilon^{(r)},$$

and define the random variable

$$Q_{n,\varepsilon} = \begin{cases} +\infty & \text{if } N_n = 0 \\ \text{the number of sample points falling in } \mathcal{P}_\varepsilon & \text{if } N_n \geq 1. \end{cases}$$

It is shown in Lemma 8 that, for μ -almost all $\mathbf{x}$,

$$Q_{n,\varepsilon} = \mathcal{O}_{\mathbb{P}}(1),$$

i.e., for any $\alpha > 0$, there exists $A > 0$ such that, for all n large enough,

$$\mathbb{P}(Q_{n,\varepsilon} \geq A) \leq \alpha. \quad (10)$$

Now, by definition of the LNN, on the event $[N_n \geq 1]$, we have

$$\mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \leq \mathbf{1}_{[\mathbf{X}_i \in \mathcal{P}_\varepsilon]},$$

and consequently,

$$\sum_{i=1}^n \mathbf{1}_{[\mathbf{X}_i \in \mathcal{L}_n(\mathbf{x})]} \mathbf{1}_{[\mathbf{X}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \leq Q_{n,\varepsilon}. \quad (11)$$

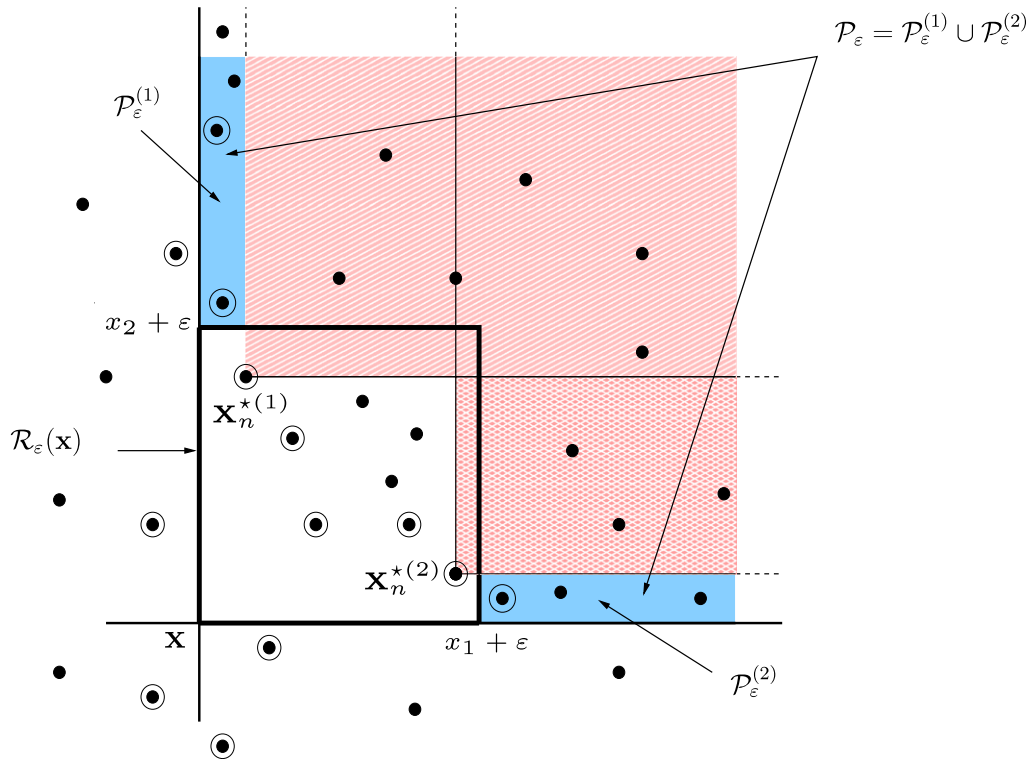


Fig. 2. Notation in dimension $d = 2$. Here $N_n = 8$ and $Q_{n,\varepsilon} = 7$. Note that none of the points in the framed area can be a LNN of $\mathbf{x}$.

Thus, for any $\alpha > 0$ and all n large enough, by (10) and (11),

$$\begin{aligned} \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} \mathbf{1}_{[\mathbf{x}_i \in \mathcal{R}_\varepsilon^c(\mathbf{x})]} \right] &\leq \mathbb{E} \left[\frac{Q_{n,\varepsilon}}{L_n(\mathbf{x})} \mathbf{1}_{[Q_{n,\varepsilon} < A]} \right] + \mathbb{E} \mathbf{1}_{[Q_{n,\varepsilon} \geq A]} \\ &= \mathbb{E} \left[\frac{Q_{n,\varepsilon}}{L_n(\mathbf{x})} \mathbf{1}_{[Q_{n,\varepsilon} < A, L_n(\mathbf{x}) \geq 1]} \right] + \mathbb{P}(Q_{n,\varepsilon} \geq A) \\ &\leq \mathbb{E} \left[\frac{A}{L_n(\mathbf{x})} \mathbf{1}_{[L_n(\mathbf{x}) \geq 1]} \right] + \alpha. \end{aligned}$$

By Theorem 1,

$$L_n(\mathbf{x}) \rightarrow \infty \quad \text{in probability as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x}$. This implies

$$\mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \mathbf{1}_{[L_n(\mathbf{x}) \geq 1]} \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

which concludes the proof of Lemma 3. $\square$

Proof of Theorem 3. Because $|a + b|^p \leq 2^{p-1}(|a|^p + |b|^p)$ for $p \geq 1$, we see that

$$\mathbb{E} |r_n(\mathbf{x}) - r(\mathbf{x})|^p \leq 2^{p-1} \mathbb{E} \left| \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} (Y_i - r(\mathbf{X}_i)) \right|^p + 2^{p-1} \mathbb{E} \left| \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} (r(\mathbf{X}_i) - r(\mathbf{x})) \right|^p.$$

Thus, by Jensen's inequality,

$$\mathbb{E} |r_n(\mathbf{x}) - r(\mathbf{x})|^p \leq 2^{p-1} \mathbb{E} \left| \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} (Y_i - r(\mathbf{X}_i)) \right|^p + 2^{p-1} \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} |r(\mathbf{X}_i) - r(\mathbf{x})|^p \right]. \quad (12)$$

The rightmost term in (12) tends to 0 for μ -almost all $\mathbf{x}$ by Lemma 3. Thus, it remains to show that the first term tends to 0 at μ -almost all $\mathbf{x}$. By successive applications of inequalities of Marcinkiewicz and Zygmund [19] (see also [20, pages 59–60]),

we have for some positive constant C_p depending only on p ,

$$\begin{aligned} \mathbb{E} \left| \frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} (Y_i - r(\mathbf{X}_i)) \right|^p &\leq C_p \mathbb{E} \left[\frac{1}{L_n^2(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} (Y_i - r(\mathbf{X}_i))^2 \right]^{p/2} \\ &\leq (2\gamma)^p C_p \mathbb{E} \left[\frac{1}{L_n^2(\mathbf{x})} \sum_{i=1}^n \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} \right]^{p/2} \quad (\text{since } |Y| \leq \gamma) \\ &= (2\gamma)^p C_p \mathbb{E} \left[\frac{1}{L_n(\mathbf{x})} \sum_{i=1}^n \frac{1}{L_n(\mathbf{x})} \mathbf{1}_{[\mathbf{x}_i \in \mathcal{L}_n(\mathbf{x})]} \right]^{p/2} \\ &= (2\gamma)^p C_p \mathbb{E} \left[\frac{1}{L_n^{p/2}(\mathbf{x})} \right]. \end{aligned}$$

By Theorem 1,

$$L_n(\mathbf{x}) \rightarrow \infty \quad \text{in probability as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x}$. Since $L_n(\mathbf{x}) \geq 1$, this implies

$$\mathbb{E} \left[\frac{1}{L_n^{p/2}(\mathbf{x})} \right] \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

and the proof is complete. $\square$

5.4. Some technical lemmas

Throughout this section, for $\mathbf{x} = (x_1, \dots, x_d)$ and $\varepsilon > 0$, $\mathcal{R}_\varepsilon(\mathbf{x})$ refers to the hyperrectangle

$$\mathcal{R}_\varepsilon(\mathbf{x}) = [x_1, x_1 + \varepsilon] \times \dots \times [x_d, x_d + \varepsilon].$$

Lemma 4. Let $\Phi : (0, \infty) \rightarrow [0, \infty)$ be a nondecreasing function with limit 0 at 0. Then there exists a sequence (ε_n) of positive real numbers such that $n\varepsilon_n^d \rightarrow \infty$ and $n\varepsilon_n^d \Phi(\varepsilon_n) \rightarrow 0$ as $n \rightarrow \infty$.

Proof of Lemma 4. Note first that if such a sequence (ε_n) exists, then $\varepsilon_n \rightarrow 0$ as $n \rightarrow \infty$. Indeed, if this is not the case, then $\varepsilon_n \geq C$ for some positive C and infinitely many n . Consequently, using the fact that Φ is nondecreasing, one obtains $n\varepsilon_n^d \Phi(\varepsilon_n) \geq n\varepsilon_n^d \Phi(C)$ for infinitely many n , and this is impossible.

For any integer $\ell \geq 1$, set $e_\ell = \Phi(1/\ell)$ and observe that the sequence (e_ℓ) is nonincreasing and tends to 0 as $\ell \rightarrow \infty$. Let $\varphi_\ell = \ell^d / \sqrt{e_\ell}$. Clearly, the sequence (φ_ℓ) is nondecreasing and satisfies $\varphi_\ell / \ell^d \rightarrow \infty$ and $[\varphi_\ell / \ell^d] \times \Phi(1/\ell) = \sqrt{e_\ell} \rightarrow 0$ as $\ell \rightarrow \infty$.

For each $n \geq 1$, let ℓ_n be the largest positive integer ℓ such that $\varphi_\ell \leq n$, and let $\varepsilon_n = 1/\ell_n$. Then the sequence (ε_n) satisfies

$$n\varepsilon_n^d \geq \varphi_{\ell_n} / \ell_n^d \rightarrow \infty$$

and

$$n\varepsilon_n^d \Phi(\varepsilon_n) \geq [\varphi_{\ell_n} / \ell_n^d] \times \Phi(1/\ell_n) \rightarrow 0$$

as $n \rightarrow \infty$. $\square$

Lemma 5. Assume that μ has a density f with respect to the Lebesgue measure on $\mathbb{R}^d$. For $\mathbf{x} \in \mathbb{R}^d$, let g_ε be the density defined by

$$g_\varepsilon(\mathbf{y}) = \begin{cases} \frac{\mu(\mathcal{R}_\varepsilon(\mathbf{x}))}{\varepsilon^d} & \text{if } \mathbf{y} \in \mathcal{R}_\varepsilon(\mathbf{x}) \\ f(\mathbf{y}) & \text{otherwise,} \end{cases}$$

and let $\mathbf{Z}_1, \dots, \mathbf{Z}_n$ be independent random vectors distributed according to g_ε . Let (ε_n) be a sequence of positive real numbers such that $\varepsilon_n \rightarrow 0$ and $n\varepsilon_n^d \rightarrow \infty$ as $n \rightarrow \infty$. Then, denoting by $L'_{\varepsilon_n}(\mathbf{x})$ the number of LNN of $\mathbf{x}$ in the sample $\{\mathbf{Z}_1, \dots, \mathbf{Z}_n\}$ falling in $\mathcal{R}_{\varepsilon_n}(\mathbf{x})$, one has

$$L'_{\varepsilon_n}(\mathbf{x}) \rightarrow \infty \quad \text{in probability as } n \rightarrow \infty,$$

at μ -almost all $\mathbf{x}$.

Proof of Lemma 5. To lighten notation a bit, we set $p_\varepsilon(\mathbf{x}) = \mu(\mathcal{R}_\varepsilon(\mathbf{x}))$. Choose $\mathbf{x}$ in a set of μ -measure 1 such that $\mu(\mathcal{R}_\varepsilon(\mathbf{x})) > 0$ for all $\varepsilon > 0$ and $np_{\varepsilon_n}(\mathbf{x}) \rightarrow \infty$ as $n \rightarrow \infty$ (by Corollary 2 this is possible).

The number of sample points falling in $\mathcal{R}_{\varepsilon_n}(\mathbf{x})$ is distributed according to some binomial random variable N_n with parameters n and $p_{\varepsilon_n}(\mathbf{x})$. Thus, we may write, for all $A > 0$,

$$\begin{aligned} \mathbb{P}(N_n < A) &\leq \mathbb{P}(N_n < np_{\varepsilon_n}(\mathbf{x})/2) \quad (\text{for all } n \text{ large enough}) \\ &= \mathbb{P}(N_n - np_{\varepsilon_n}(\mathbf{x}) < -np_{\varepsilon_n}(\mathbf{x})/2) \\ &\leq 4/(np_{\varepsilon_n}(\mathbf{x})) \quad (\text{by Chebyshev's inequality}), \end{aligned}$$

from which we deduce that $N_n \rightarrow \infty$ in probability as $n \rightarrow \infty$. This implies that

$$\mathbb{E}\left[\frac{1}{(\log N_n)^{d-1}} \mathbf{1}_{[N_n \geq 2]}\right] \rightarrow 0 \quad \text{as } n \rightarrow \infty. \quad (13)$$

Now, denote by K_m the number of maxima in a sequence of m i.i.d. points chosen uniformly at random from $(0, 1)^d$. Using the fact that the $\mathbf{Z}_i$'s which fall in $\mathcal{R}_{\varepsilon_n}(\mathbf{x})$ are uniformly distributed on $\mathcal{R}_{\varepsilon_n}(\mathbf{x})$, we note that $L'_{\varepsilon_n}(\mathbf{x})$ and K_{N_n} have the same distribution. Therefore, the theorem will be proven if we show that $K_{N_n} \rightarrow \infty$ in probability as $n \rightarrow \infty$.

A straightforward adaptation of the arguments in Barndorff-Nielsen and Sobel [4] and Bai et al. [3,2] shows that there exist two positive constants Δ_1 and Δ_2 such that, on the event $[N_n \geq 2]$,

$$\mathbb{E}[K_{N_n} | N_n] \geq \Delta_1 (\log N_n)^{d-1} \quad (14)$$

and

$$\mathbb{V}[K_{N_n} | N_n] \leq \Delta_2 (\log N_n)^{d-1}. \quad (15)$$

Fix $A > 0$ and $\alpha > 0$, and let the event $\mathcal{E}_n$ be defined as

$$\mathcal{E}_n = \left[N_n < e^{(2A/\Delta_1)^{1/(d-1)}} \vee 2 \right].$$

Since $N_n \rightarrow \infty$ in probability, one has $\mathbb{P}(\mathcal{E}_n) \leq \alpha$ for all n large enough. Using (14), we may write, conditionally on N_n ,

$$\mathbb{P}(K_{N_n} < A | N_n) \leq \mathbb{P}(K_{N_n} < \mathbb{E}[K_{N_n} | N_n]/2 | N_n) \mathbf{1}_{\mathcal{E}_n^c} + \mathbf{1}_{\mathcal{E}_n}.$$

Thus, by Chebyshev's inequality and inequalities (14)–(15),

$$\mathbb{P}(K_{N_n} < A | N_n) \leq \frac{\Delta}{(\log N_n)^{d-1}} \mathbf{1}_{\mathcal{E}_n^c} + \mathbf{1}_{\mathcal{E}_n}$$

for some positive constant Δ . Taking expectations on both sides, we finally obtain, for all n large enough,

$$\mathbb{P}(K_{N_n} < A) \leq \mathbb{E}\left[\frac{\Delta}{(\log N_n)^{d-1}} \mathbf{1}_{[N_n \geq 2]}\right] + \alpha,$$

which, together with (13), completes the proof of the lemma. $\square$

Lemma 6. Let $\Delta \in (0, 1)$. Then, for all $n \geq 1$,

$$n \int_{[0, \Delta]^d} (1 - \prod y_i)^{n-1} d\mathbf{y} = \frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}_\Delta((\log n)^{d-2}),$$

where the notation $\mathcal{O}_\Delta$ means that the constant in the $\mathcal{O}$ term depends on Δ .

Proof of Lemma 6. The proof starts with the observation (see for example Bai et al. [3]) that

$$n \int_{[0, 1]^d} (1 - \prod y_i)^{n-1} d\mathbf{y} = \frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}((\log n)^{d-2}). \quad (16)$$

To show the result, we proceed by induction on $d \geq 2$. For $d = 2$, we may write

$$\begin{aligned} n \int_{[0, \Delta]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 &= n \int_{[0, 1]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 - n \int_{[0, 1]^2 \setminus [0, \Delta]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 \\ &= \log n + \mathcal{O}(1) - n \int_{[0, 1]^2 \setminus [0, \Delta]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 \quad (\text{by identity (16)}). \end{aligned}$$

Observing that

$$n \int_{[0,1]^2 \setminus [0,\Delta]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 \leq 2n \int_{[0,1]} (1 - \Delta y)^{n-1} dy \leq 2/\Delta$$

yields

$$n \int_{[0,\Delta]^2} (1 - y_1 y_2)^{n-1} dy_1 dy_2 = \log n + \mathcal{O}_\Delta(1),$$

as desired. Having disposed of this preliminary step, assume that, for all positive $\Delta \in (0, 1)$,

$$n \int_{[0,\Delta]^d} (1 - \prod y_i)^{n-1} d\mathbf{y} = \frac{(\log n)^{d-1}}{(d-1)!} + \mathcal{O}_\Delta((\log n)^{d-2}). \quad (17)$$

Then, for $d+1$,

$$\begin{aligned} n \int_{[0,\Delta]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} &= n \int_{[0,1]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} - n \int_{[0,1]^{d+1} \setminus [0,\Delta]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} \\ &= \frac{(\log n)^d}{d!} + \mathcal{O}((\log n)^{d-1}) - n \int_{[0,1]^{d+1} \setminus [0,\Delta]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} \quad (\text{by identity (16)}). \end{aligned}$$

With respect to the rightmost term, we note that

$$\begin{aligned} n \int_{[0,1]^{d+1} \setminus [0,\Delta]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} &\leq nd \int_{[0,1]^d} (1 - \Delta \prod y_i)^{n-1} d\mathbf{y} \\ &= n(d/\Delta) \int_{[0,\Delta^{1/d}]^d} (1 - \prod y_i)^{n-1} d\mathbf{y} \\ &= \frac{d(\log n)^{d-1}}{\Delta(d-1)!} + \mathcal{O}_\Delta((\log n)^{d-2}) \quad (\text{by induction hypothesis (17)}) \\ &= \mathcal{O}_\Delta((\log n)^{d-1}). \end{aligned}$$

Putting all pieces together, we obtain

$$n \int_{[0,\Delta]^{d+1}} (1 - \prod y_i)^{n-1} d\mathbf{y} = \frac{(\log n)^d}{d!} + \mathcal{O}_\Delta((\log n)^{d-1}),$$

as desired. $\square$

For a better understanding of the next two lemmas, the reader should refer to Fig. 2.

Lemma 7. Assume that μ has a density f with respect to the Lebesgue measure on $\mathbb{R}^d$. Fix $\mathbf{x} = (x_1, \dots, x_d)$, $\varepsilon > 0$, and denote by N_n the (random) number of sample points falling in $\mathcal{R}_\varepsilon(\mathbf{x})$. For $N_n \geq 1$ and each $r = 1, \dots, d$, let $\mathbf{X}_n^{*(r)} = (X_{n,1}^{*(r)}, \dots, X_{n,d}^{*(r)})$ be the observation in $\mathcal{R}_\varepsilon(\mathbf{x})$ whose r -coordinate is the closest to x_r . Define the random variables

$$M_{n,r} = \begin{cases} +\infty & \text{if } N_n = 0 \\ X_{n,r}^{*(r)} - x_r & \text{if } N_n \geq 1. \end{cases}$$

Then, for μ -almost all $\mathbf{x}$,

$$M_{n,r} = \mathcal{O}_{\mathbb{P}}\left(\frac{1}{n}\right),$$

i.e., for any $\alpha > 0$, there exists $A > 0$ such that, for all n large enough,

$$\mathbb{P}\left(M_{n,r} \geq \frac{A}{n}\right) \leq \alpha.$$

Proof of Lemma 7. Note first that $\mathbf{X}_n^{*(r)}$ is almost surely uniquely defined. Choose $\mathbf{x}$ in a set of μ -measure 1 such that $\mu(\mathcal{R}_\varepsilon(\mathbf{x})) > 0$ and set $p_\varepsilon(\mathbf{x}) = \mu(\mathcal{R}_\varepsilon(\mathbf{x}))$. For any $r = 1, \dots, d$, let $\mathcal{T}_\varepsilon^r(\mathbf{x})$ be the $d-1$ -dimensional rectangle defined by

$$\mathcal{T}_\varepsilon^r(\mathbf{x}) = \{\mathbf{y} = (y_1, \dots, y_{r-1}, y_{r+1}, \dots, y_d) \in \mathbb{R}^{d-1} : x_j \leq y_j \leq x_j + \varepsilon, j \neq r\},$$

and let

$$f_{\varepsilon, \mathbf{x}}^{(r)}(z) = \frac{\mathbf{1}_{[0 \leq z \leq \varepsilon]}}{\mu(\mathcal{R}_\varepsilon(\mathbf{x}))} \int_{T_\varepsilon^{(r)}(\mathbf{x})} f(y_1, \dots, y_{r-1}, z, y_{r+1}, \dots, y_d) dy_1 \cdots dy_{r-1} dy_{r+1} \cdots dy_d$$

be the marginal density of the distribution μ conditioned by the event $[\mathbf{X} \in \mathcal{R}_\varepsilon(\mathbf{x})]$. Note that we can still choose $\mathbf{x}$ in a set of μ -measure 1 such that, for any $r = 1, \dots, d$, $f_{\varepsilon, \mathbf{x}}^{(r)}(x_r) > 0$ and $f_{\varepsilon, \mathbf{x}}^{(r)}(z)$ satisfies (3) at x_r , i.e.,

$$\int_{x_r}^{x_r+t} f_{\varepsilon, \mathbf{x}}^{(r)}(z) dz = t f_{\varepsilon, \mathbf{x}}^{(r)}(x_r) + t \zeta_r(t), \quad \text{with } \lim_{t \rightarrow 0^+} \zeta_r(t) = 0.$$

Since N_n is binomial with parameters n and $p_{\varepsilon_n}(\mathbf{x})$, we have for any $r = 1, \dots, d$ and $t > 0$,

$$\begin{aligned} \mathbb{P}(M_{n,r} \geq t) &= \mathbb{E}[\mathbb{P}(M_{n,r} > t | N_n)] \\ &\leq \mathbb{E}[\mathbf{1}_{[N_n > 0]} \mathbb{P}(M_{n,r} > t | N_n)] + \mathbb{P}(N_n = 0) \\ &\leq \mathbb{E}\left[\left(1 - \int_{x_r}^{x_r+t} f_{\varepsilon, \mathbf{x}}^{(r)}(z) dz\right)^{N_n}\right] + (1 - p_\varepsilon(\mathbf{x}))^n \\ &= \left[\left(1 - \int_{x_r}^{x_r+t} f_{\varepsilon, \mathbf{x}}^{(r)}(z) dz\right) p_\varepsilon(\mathbf{x}) + 1 - p_\varepsilon(\mathbf{x})\right]^n + (1 - p_\varepsilon(\mathbf{x}))^n \\ &= \left(1 - p_\varepsilon(\mathbf{x}) \int_{x_r}^{x_r+t} f_{\varepsilon, \mathbf{x}}^{(r)}(z) dz\right)^n + (1 - p_\varepsilon(\mathbf{x}))^n \\ &\leq \exp\left(-np_\varepsilon(\mathbf{x}) \int_{x_r}^{x_r+t} f_{\varepsilon, \mathbf{x}}^{(r)}(z) dz\right) + \exp(-np_\varepsilon(\mathbf{x})) \\ &= \exp(-ntp_\varepsilon(\mathbf{x}) (f_{\varepsilon, \mathbf{x}}^{(r)}(x_r) + \zeta_r(t))) + \exp(-np_\varepsilon(\mathbf{x})). \end{aligned}$$

This shows that

$$M_{n,r} = \mathcal{O}_{\mathbb{P}}\left(\frac{1}{n}\right),$$

as desired. $\square$

With the notation of Lemma 7, we define the random variable

$$Q_{n,\varepsilon} = \begin{cases} +\infty & \text{if } N_n = 0 \\ \text{the number of sample points falling in } \mathcal{P}_\varepsilon & \text{if } N_n \geq 1, \end{cases}$$

where, in the second statement,

$$\mathcal{P}_\varepsilon = \bigcup_{r=1}^d \mathcal{P}_\varepsilon^{(r)}$$

and

$$\mathcal{P}_\varepsilon^{(r)} = [x_1 + \varepsilon, +\infty[\times \cdots \times [x_{r-1} + \varepsilon, +\infty[\times [x_r, X_{n,r}^{*(r)}] \times [x_{r+1} + \varepsilon, +\infty[\times \cdots \times [x_d + \varepsilon, +\infty[.$$

Lemma 8. Assume that μ has a density f with respect to the Lebesgue measure on $\mathbb{R}^d$. For μ -almost all $\mathbf{x}$,

$$Q_{n,\varepsilon} = \mathcal{O}_{\mathbb{P}}(1).$$

Proof of Lemma 8. For $N_n \geq 1$, denote by $Q_{n,\varepsilon}^{(r)}$ the number of sample points falling in $\mathcal{P}_\varepsilon^{(r)}$, and set $Q_{n,\varepsilon}^{(r)} = +\infty$ otherwise. Then, clearly,

$$Q_{n,\varepsilon} = \sum_{r=1}^d Q_{n,\varepsilon}^{(r)}.$$

Therefore, the result will be proven if we show that, for μ -almost all $\mathbf{x}$ and all $r = 1, \dots, d$,

$$Q_{n,\varepsilon}^{(r)} = \mathcal{O}_{\mathbb{P}}(1).$$

We fix $\mathbf{x}$ for which Lemma 7 is satisfied and fix $r \in \{1, \dots, d\}$.

Let $\alpha > 0$. According to Lemma 7, there exists $A > 0$ such that, for all n large enough,

$$\mathbb{P}\left(M_{n,r} \geq \frac{A}{n}\right) \leq \alpha.$$

Denoting by $\mathcal{E}_n$ the event

$$\left[M_{n,r} < \frac{A}{n} \right],$$

we obtain, for all $t > 0$,

$$\begin{aligned} \mathbb{P}(Q_{n,\varepsilon}^{(r)} \geq t) &= \mathbb{E}[\mathbb{P}(Q_{n,\varepsilon}^{(r)} \geq t | M_{n,r})] \\ &\leq \mathbb{E}[\mathbf{1}_{\mathcal{E}_n} \mathbb{P}(Q_{n,\varepsilon}^{(r)} \geq t | M_{n,r})] + \mathbb{P}(\mathcal{E}_n^c) \\ &\leq \mathbb{E}[\mathbf{1}_{\mathcal{E}_n} \mathbb{P}(Q_{n,\varepsilon}^{(r)} \geq t | M_{n,r})] + \alpha \quad (\text{for all } n \text{ large enough}) \\ &\leq \frac{\mathbb{E}[\mathbf{1}_{\mathcal{E}_n} \mathbb{E}[Q_{n,\varepsilon}^{(r)} | M_{n,r}]]}{t} + \alpha \quad (\text{by Markov's inequality}). \end{aligned}$$

With respect to the first term in the last inequality we may write, using the definition of $\mathcal{E}_n$,

$$\begin{aligned} \mathbf{1}_{\mathcal{E}_n} \mathbb{E}[Q_{n,\varepsilon}^{(r)} | M_{n,r}] &= n \mathbf{1}_{\mathcal{E}_n} \int_{x_1+\varepsilon}^{\infty} \cdots \int_{x_{r-1}+\varepsilon}^{\infty} \int_{x_r}^{x_r+M_{n,r}} \int_{x_{r+1}+\varepsilon}^{\infty} \cdots \int_{x_d+\varepsilon}^{\infty} f(\mathbf{y}) d\mathbf{y} \\ &\leq n \int_{x_1+\varepsilon}^{\infty} \cdots \int_{x_{r-1}+\varepsilon}^{\infty} \int_{x_r}^{x_r+A/n} \int_{x_{r+1}+\varepsilon}^{\infty} \cdots \int_{x_d+\varepsilon}^{\infty} f(\mathbf{y}) d\mathbf{y}. \end{aligned}$$

Let

$$g_{\varepsilon,\mathbf{x}}^{(r)}(z) = \int_{x_1+\varepsilon}^{\infty} \cdots \int_{x_{r-1}+\varepsilon}^{\infty} \int_{x_{r+1}+\varepsilon}^{\infty} \cdots \int_{x_d+\varepsilon}^{\infty} f(y_1, \dots, y_{r-1}, z, y_{r+1}, \dots, y_d) dy_1 \cdots dy_{r-1} dy_{r+1} \cdots dy_d,$$

and observe that we can still choose $\mathbf{x}$ in a set of μ -measure 1 such that $g_{\varepsilon,\mathbf{x}}^{(r)}(z)$ satisfies (3), i.e.,

$$\int_{x_r}^{x_r+t} g_{\varepsilon,\mathbf{x}}^{(r)}(z) dz = t g_{\varepsilon,\mathbf{x}}^{(r)}(x_r) + t \zeta_r(t), \quad \text{with } \lim_{t \rightarrow 0^+} \zeta_r(t) = 0.$$

Thus, for $\delta > 0$, we can take n large enough to ensure

$$n \int_{x_r}^{x_r+A/n} g_{\varepsilon,\mathbf{x}}^{(r)}(z) dz \leq A(1 + \delta) g_{\varepsilon,\mathbf{x}}^{(r)}(x_r).$$

Putting all pieces together, we obtain, for any $t > 0$, $\delta > 0$, $\alpha > 0$, and all n large enough,

$$\mathbb{P}(Q_{n,\varepsilon}^{(r)} \geq t) \leq \frac{1}{t} A(1 + \delta) g_{\varepsilon,\mathbf{x}}^{(r)}(x_r) + \alpha.$$

This shows that

$$Q_{n,\varepsilon}^{(r)} = \mathcal{O}_{\mathbb{P}}(1). \quad \square$$

Acknowledgments

We thank Jérôme Droniou and Arnaud Guyader for stimulating discussions. We also thank two referees for valuable comments and insightful suggestions.

G rard Biau's research was supported by the French "Agence Nationale pour la Recherche" under grant ANR-09-BLAN-0051-02 "CLARA". The research at the ENS is carried out within the INRIA project "CLASSIC" hosted by Ecole Normale Sup rieure and CNRS. Luc Devroye's research was sponsored by NSERC Grant A3456 and FQRNT Grant 90-ER-0291.

References

- [1] Y. Amit, D. Geman, Shape quantization and recognition with randomized trees, *Neural Computation* 9 (1997) 1545–1588.
- [2] Z.-D. Bai, C.-C. Chao, H.-K. Hwang, W.-Q. Liang, On the variance of the number of maxima in random vectors and its applications, *The Annals of Applied Probability* 8 (1998) 886–895.
- [3] Z.-D. Bai, L. Devroye, H.-K. Hwang, T.-S. Tsai, Maxima in hypercubes, *Random Structures and Algorithms* 27 (2005) 290–309.
- [4] O. Barndorff-Nielsen, M. Sobel, On the distribution of admissible points in a vector random sample, *Theory of Probability and its Applications* 11 (1966) 249–269.
- [5] G. Biau, L. Devroye, G. Lugosi, Consistency of random forests and other averaging classifiers, *Journal of Machine Learning Research* 9 (2008) 2015–2033.
- [6] L. Breiman, Bagging predictors, *Machine Learning* 24 (1996) 123–140.
- [7] L. Breiman, Some infinite theory for predictor ensembles, Technical Report 577, Statistics Department, UC Berkeley, 2000.
<http://www.stat.berkeley.edu/~breiman>.

- [8] L. Breiman, Random forests, *Machine Learning* 45 (2001) 5–32.
- [9] L. Breiman, Consistency for a simple model of random forests, Technical Report 670, Statistics Department, UC Berkeley, 2004.
<http://www.stat.berkeley.edu/~breiman>.
- [10] T.M. Cover, P.E. Hart, Nearest neighbour pattern classification, *IEEE Transactions on Information Theory* 13 (1967) 21–27.
- [11] A. Cutler, G. Zhao, Pert—Perfect random tree ensembles, *Computing Science and Statistics* 33 (2001) 490–497.
- [12] L. Devroye, The uniform convergence of nearest neighbour regression function estimators and their application in optimization, *IEEE Transactions on Information Theory* 24 (1978) 142–151.
- [13] L. Devroye, L. Györfi, G. Lugosi, *A Probabilistic Theory of Pattern Recognition*, Springer-Verlag, New York, 1996.
- [14] T.G. Dietterich, An experimental comparison of three methods for constructing ensembles of decision trees: Bagging, boosting, and randomization, *Machine Learning* 40 (2000) 139–157.
- [15] W. Doeblin, Exposé de la théorie des chaînes simples constantes de Markov à un nombre fini d'états, *Revue Mathématique de l'Union Interbalkanique* 2 (1937) 77–105.
- [16] E. Fix, J. Hodges, Discriminatory analysis. Nonparametric discrimination: Consistency properties, Technical Report 4, Project Number 21-49-004, USAF School of Aviation Medicine, Randolph Field, Texas, 1951.
- [17] L. Györfi, M. Kohler, A. Krzyżak, H. Walk, *A Distribution-Free Theory of Nonparametric Regression*, Springer-Verlag, New York, 2002.
- [18] Y. Lin, Y. Jeon, Random forests and adaptive nearest neighbors, *Journal of the American Statistical Association* 101 (2006) 578–590.
- [19] J. Marcinkiewicz, A. Zygmund, Sur les fonctions indépendantes, *Fundamenta Mathematicae* 29 (1937) 60–90.
- [20] V.V. Petrov, *Sums of Independent Random Variables*, Springer-Verlag, Berlin, 1975.
- [21] S.T. Rachev, L. Rüschendorf, *Mass Transportation Problems, Volume I: Theory*, Springer, New York, 1998.
- [22] B.M. Steele, Exact bootstrap k -nearest neighbor learners, *Machine Learning* 74 (2009) 235–255.
- [23] C.J. Stone, Consistent nonparametric regression, *The Annals of Statistics* 5 (1977) 595–645.
- [24] R.L. Wheeden, A. Zygmund, *Measure and Integral. An Introduction to Real Analysis*, Marcel Dekker, New York, 1977.