



An affine invariant k -nearest neighbor regression estimate

G  rard Biau^{a,b,*}, Luc Devroye^c, Vida Dujmovi  ^d, Adam Krzy  ak^e

^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

^d School of Computer Science, Carleton University, 5302 Herzberg Laboratories, 1125 Colonel By Drive, Ottawa, Canada K1S 5B6

^e Department of Computer Science and Software Engineering, Concordia University, Montreal, Canada H3G 1M8

ARTICLE INFO

Article history:

Received 2 January 2012

Available online 17 June 2012

AMS 2000 subject classifications:

62G08

62G05

62G20

Keywords:

Nonparametric estimation

Regression function estimation

Affine invariance

Nearest neighbor methods

Mathematical statistics

ABSTRACT

We design a data-dependent metric in $\mathbb{R}^d$ and use it to define the k -nearest neighbors of a given point. Our metric is invariant under all affine transformations. We show that, with this metric, the standard k -nearest neighbor regression estimate is asymptotically consistent under the usual conditions on k , and minimal requirements on the input data.

   2012 Elsevier Inc. All rights reserved.

1. Introduction

The prediction error of standard nonparametric regression methods may be critically affected by a linear transformation of the coordinate axes. It is typically the case for the popular k -nearest neighbor (k -NN) predictor [11,12,7,5,6], where a mere rescaling of the coordinate axes has a serious impact on the capabilities of this estimate. This is clearly an undesirable feature, especially in applications where the data measurements represent physically different quantities, such as temperature, blood pressure, cholesterol level, and the age of the patient. In this example, a simple change in, say, the unit measure of the temperature parameter will lead to totally different results, and will thus force the statistician to use a somewhat arbitrary preprocessing step prior to the k -NN estimation process. Furthermore, in several practical implementations, one would like, for physical or economical reasons, to supply the freshly collected data to some machine without preprocessing.

In this paper, we discuss a variation of the k -NN regression estimate whose definition is not affected by affine transformations of the coordinate axes. Such a modification could save the user a subjective preprocessing step and would save the manufacturer the trouble of adding input specifications.

The data set we have collected can be regarded as a collection of independent and identically distributed $\mathbb{R}^d \times \mathbb{R}$ -valued random variables $\mathcal{D}_n = \{(\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)\}$, independent of and with the same distribution as a generic pair $(\mathbf{X}, Y)$ satisfying $\mathbb{E}|Y| < \infty$. The space $\mathbb{R}^d$ is equipped with the standard Euclidean norm $\|\cdot\|$. For fixed $\mathbf{x} \in \mathbb{R}^d$, our goal is to

* 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), vida@scs.carleton.ca (V. Dujmovi  ), krzyzak@cs.concordia.ca (A. Krzy  ak).

estimate the regression function $r(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$ using the data $\mathcal{D}_n$. In this context, the usual k -NN regression estimate takes the form

$$r_n(\mathbf{x}; \mathcal{D}_n) = \frac{1}{k_n} \sum_{i=1}^{k_n} Y_{(i)}(\mathbf{x}),$$

where $(\mathbf{X}_{(1)}(\mathbf{x}), Y_{(1)}(\mathbf{x})), \dots, (\mathbf{X}_{(k_n)}(\mathbf{x}), Y_{(k_n)}(\mathbf{x}))$ is a reordering of the data according to increasing distances $\|\mathbf{X}_i - \mathbf{x}\|$ of the $\mathbf{X}_i$'s to $\mathbf{x}$. (If distance ties occur, a tie-breaking strategy must be defined. For example, if $\|\mathbf{X}_i - \mathbf{x}\| = \|\mathbf{X}_j - \mathbf{x}\|$, $\mathbf{X}_i$ may be declared “closer” if $i < j$, i.e., the tie-breaking is done by indices.) For simplicity, we will suppress $\mathcal{D}_n$ in the notation and write $r_n(\mathbf{x})$ instead of $r_n(\mathbf{x}; \mathcal{D}_n)$. Stone [37] showed that, for all $p \geq 1$, $\mathbb{E}[r_n(\mathbf{X}) - r(\mathbf{X})]^p \rightarrow 0$ for all possible distributions of $(\mathbf{X}, Y)$ with $\mathbb{E}|Y|^p < \infty$, whenever $k_n \rightarrow \infty$ and $k_n/n \rightarrow 0$ as $n \rightarrow \infty$. Thus, the k -NN estimate behaves asymptotically well, without exceptions. This property is called L_p universal consistency.

Clearly, any affine transformation of the coordinate axes influences the k -NN estimate through the norm $\|\cdot\|$, thereby illuminating an unpleasant face of the procedure. To illustrate this remark, assume that a nontrivial affine transformation $T: \mathbf{z} \mapsto A\mathbf{z} + \mathbf{b}$ (that is, a nonsingular linear transformation A followed by a translation $\mathbf{b}$) is applied to both $\mathbf{x}$ and $\mathbf{X}_1, \dots, \mathbf{X}_n$. Examples include any number of combinations of rotations, translations, and linear rescalings. Denote by $\mathcal{D}'_n = (T(\mathbf{X}_1), Y_1), \dots, (T(\mathbf{X}_n), Y_n)$ the transformed sample. Then, for such a function T , one has $r_n(\mathbf{x}; \mathcal{D}_n) \neq r_n(T(\mathbf{x}); \mathcal{D}'_n)$ in general, whereas $r(\mathbf{X}) = \mathbb{E}[Y|T(\mathbf{X})]$ since T is bijective. Thus, to continue our discussion, we are looking in essence for a regression estimate r_n with the following property:

$$r_n(\mathbf{x}; \mathcal{D}_n) = r_n(T(\mathbf{x}); \mathcal{D}'_n). \quad (1)$$

We call r_n affine invariant. Affine invariance is indeed a very strong but highly desirable property. In $\mathbb{R}^d$, in the context of k -NN estimates, it suffices to be able to define an affine invariant distance measure, which is necessarily data-dependent. With this objective in mind, we develop in the next section an estimation procedure featuring (1) which in form coincides with the k -NN estimate, and establish its consistency in Section 3. Proofs of the most technical results are gathered in Section 4.

It should be stressed that what we are after in this article is an estimate of r which is invariant by an affine transformation of *both* the query point $\mathbf{x}$ and the original regressors $\mathbf{X}_1, \dots, \mathbf{X}_n$. When the sole regressors are subject to such a transformation, it is then more natural to talk of “affine equivariant” regression estimates rather than of “affine invariant” ones; this is more in line with the terminology used, for example, in [28,29]. These affine invariance and affine equivariance requirements, however, are strictly equivalent.

There have been many attempts in the nonparametric literature to achieve affine invariance. One of the most natural ones relates to the so-called transformation-retransformation proposed by Chakraborty et al. [3]. That method and many variants have been discussed in texts such as [9,17] for pattern recognition and regression, respectively, but they have also been used in kernel density estimation (see, e.g., [36]). It is worth noting that, computational issues aside, the transformation step (i.e., premultiplication of the regressors by $\hat{M}_n^{-1}$, where $\hat{M}_n$ is an affine equivariant scatter estimate) may be based on a statistic $\hat{M}_n$ that does not require finiteness of any moment. A typical example is the scatter estimate proposed in [38] or [20]. Rather, our procedure takes ideas from the classical nonparametric literature using concepts such as multivariate ranks. It is close in spirit to the approach of Paindaveine and Van Bever [31], who introduce a class of depth-based classification procedures that are of a nearest neighbor nature.

There are also attempts at getting invariance to other transformations. The most important concept here is that of invariance under monotone transformations of the coordinate axes. In particular, any strategy that uses only the coordinatewise ranks of the $\mathbf{X}_i$'s achieves this. The onus, then, is to show consistency of the methods under the most general conditions possible. For example, using an L_p norm on the d -vectors of differences between ranks, one can show that the classical k -NN regression function estimate is universally consistent in the sense of Stone [37]. This was observed by Olshen [30], and shown by Devroye [8] (see also [15,16,10,2] for related works). Rules based upon statistically equivalent blocks (see, e.g., [1,34,13,14], and [9, Section 21.4]) are other important examples of regression methods invariant with respect to monotone transformations of the coordinate axes. These methods and their generalizations partition the space with sets that contain a fixed number of data points each.

It would be interesting to consider in a future paper the possibility of morphing the input space in more general ways than those suggested in the previous few paragraphs of the present article. It should be possible, in principle, to define appropriate metrics to obtain invariance for interesting large classes of nonlinear transformations, and show consistent asymptotic behaviors.

2. An affine invariant k -NN estimate

The k -NN estimate we are discussing is based upon the notion of empirical distance. Throughout, we assume that the distribution of $\mathbf{X}$ is absolutely continuous with respect to the Lebesgue measure on $\mathbb{R}^d$ and that $n \geq d$. Because of this density assumption, any collection $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}$ ($1 \leq i_1 < i_2 < \dots$

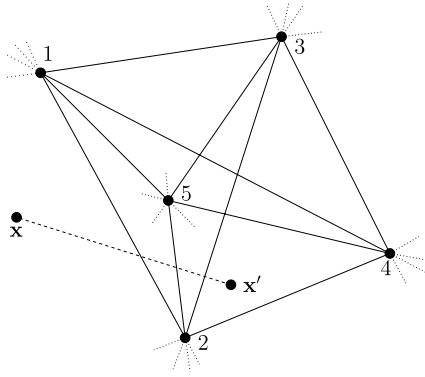


Fig. 1. An example in dimension 2. The empirical distance between $\mathbf{x}$ and $\mathbf{x}'$ is 4. (Note that the hyperplane defined by the pair (3, 5) indeed cuts the segment $(\mathbf{x}, \mathbf{x}')$, so that the distance is 4, not 3.)

With this notation, the empirical distance between d -vectors $\mathbf{x}$ and $\mathbf{x}'$ is defined as

$$\rho_n(\mathbf{x}, \mathbf{x}') = \sum_{1 \leq i_1 < \dots < i_d \leq n} \mathbf{1}_{\{\text{segment } (\mathbf{x}, \mathbf{x}') \text{ intersects the hyperplane } \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})\}}.$$

Put differently, $\rho_n(\mathbf{x}, \mathbf{x}')$ just counts the number of hyperplanes in $\mathbb{R}^d$ passing through d out of the points $\mathbf{X}_1, \dots, \mathbf{X}_n$, that are separating $\mathbf{x}$ and $\mathbf{x}'$. Roughly, “near” points have fewer intersections, see Fig. 1 which depicts an example in dimension 2.

This hyperplane-based concept of distance is known in the multivariate rank tests literature as the empirical lift-interdirection function [27], see also [35,26,18] for companion concepts). It was originally mentioned (but not analyzed) in [19], and independently suggested as an affine invariant alternative to ordinary metrics in the monograph by Devroye et al. [9, Section 11.6]. We speak throughout of distance, even though, for a fixed sample of size n , ρ_n is only defined with probability 1 and is not a distance measure *stricto sensu* (in particular, $\rho_n(\mathbf{x}, \mathbf{x}') = 0$ does not imply that $\mathbf{x} = \mathbf{x}'$). Nevertheless, this empirical distance is invariant under affine transformations $\mathbf{x} \mapsto A\mathbf{x} + \mathbf{b}$, where A is some arbitrary nonsingular linear map and $\mathbf{b}$ any offset vector (see, for instance, [27, Section 2.4]).

Now, fix $\mathbf{x} \in \mathbb{R}^d$ and let $\rho_n(\mathbf{x}, \mathbf{X}_i)$ be the empirical distance between $\mathbf{x}$ and some observation $\mathbf{X}_i$ in the sample $\mathbf{X}_1, \dots, \mathbf{X}_n$. (That is, the number of hyperplanes in $\mathbb{R}^d$ passing through d out of the observations $\mathbf{X}_1, \dots, \mathbf{X}_n$, that are cutting the segment $(\mathbf{x}, \mathbf{X}_i)$.) In this context, the k -NN estimate we are considering still takes the familiar form

$$r_n(\mathbf{x}) = \frac{1}{k_n} \sum_{i=1}^{k_n} Y_{(i)}(\mathbf{x}),$$

with the important difference that now the data set $(\mathbf{X}_1, Y_1), \dots, (\mathbf{X}_n, Y_n)$ is reordered according to increasing values of the empirical distances $\rho_n(\mathbf{x}, \mathbf{X}_i)$, not the original Euclidean metric. By construction, the estimate r_n has the desired affine invariance property and, moreover, it coincides with the standard (Euclidean) estimate in dimension $d = 1$. In the next section, we prove the following theorem. The distribution of the random variable $\mathbf{X}$ is denoted by μ .

Theorem 1 (Pointwise L_p Consistency). Assume that $\mathbf{X}$ has a probability density, that Y is bounded, and that the regression function r is μ -almost surely continuous. Then, for μ -almost all $\mathbf{x} \in \mathbb{R}^d$ and all $p \geq 1$, if $k_n \rightarrow \infty$ and $k_n/n \rightarrow 0$,

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

The following corollary is a consequence of Theorem 1 and the Lebesgue's dominated convergence theorem.

Corollary 1 (Global L_p Consistency). Assume that $\mathbf{X}$ has a probability density, that Y is bounded, and that the regression function r is μ -almost surely continuous. Then, for all $p \geq 1$, if $k_n \rightarrow \infty$ and $k_n/n \rightarrow 0$,

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

The conditions of Stone's universal consistency theorem given in [37] are not fulfilled for our estimate. For the standard nearest neighbor estimate, a key result used in the consistency proof by Stone is that a given data point cannot be the nearest neighbor of more than a constant number (say, 3^d) of other points. Such a universal constant does not exist after our transformation is applied. That means that a single data point can have a large influence on the regression function estimate. While this by itself does not imply that the estimate is not universally consistent, it certainly indicates that any such proof will require new insights. The addition of two smoothness constraints, namely that $\mathbf{X}$ has a density (without, however, imposing any continuity conditions on the density itself) and that r is μ -almost surely continuous, is sufficient.

The complexity of our procedure in terms of sample size n and dimension d is quite high. There are $\binom{n}{d}$ possible choices of hyperplanes through d points. This collection of hyperplanes defines an arrangement, or partition, of $\mathbb{R}^d$ into

</

polytopal regions, also called cells or chambers. Within each region, the distance to each data point is constant and, thus, a preprocessing step might consist of setting up a data structure for determining to which cell a given point $\mathbf{x} \in \mathbb{R}^d$ belongs: this is called the point location problem. Meiser [24] showed that such a data structure exists with the following properties: (i) it takes space $\mathcal{O}(n^{d+\varepsilon})$ for any fixed $\varepsilon > 0$, and (ii) point location can be performed in $\mathcal{O}(\log n)$ time. Chazelle's cuttings [4] improve (i) to $\mathcal{O}(n^d)$. Chazelle's processing time for setting up the data structure is $\mathcal{O}(n^d)$. Still in the preprocessing step, one can determine for each cell in the arrangement the distances to all n data points: this can be done by walking across the graph of cells or by brute force. When done naively, the overall set-up complexity is $\mathcal{O}(n^{2d+1})$. For each cell, one might keep a pointer to the k nearest neighbors. Therefore, once set up, the computation of the regression function estimate takes merely $\mathcal{O}(\log n)$ time for point location, and $\mathcal{O}(k)$ time for retrieving the k nearest neighbors.

One could envisage a reduction in the complexity by defining the distances not in terms of all hyperplanes that cut a line segment, but in terms of the number of randomly drawn hyperplanes that make such a cut, where the number of random draws is now a carefully selected number. By the concentration of binomial random variables, such random estimates of the distances are expected to work well, while keeping the complexity reasonable. This idea will be explored elsewhere.

3. Proof of the theorem

Recall, since $\mathbf{X}$ has a probability density with respect to the Lebesgue measure on $\mathbb{R}^d$, that any collection $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}$ ($1 \leq i_1 < i_2 < \dots < i_d \leq n$) of d points among $\mathbf{X}_1, \dots, \mathbf{X}_n$ defines with probability 1 a unique hyperplane $\mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})$ in $\mathbb{R}^d$. Thus, in the sequel, since no confusion is possible, we will freely refer to “the hyperplane $\mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})$ defined by $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}$ ” without further explicit mention of the probability 1 event.

Let us first fix some useful notation. The distribution of the random variable $\mathbf{X}$ is denoted by μ and its density with respect to the Lebesgue measure is denoted by f . For every $\varepsilon > 0$, we let $\mathcal{B}_{\mathbf{x}, \varepsilon} = \{\mathbf{y} \in \mathbb{R}^d : \|\mathbf{y} - \mathbf{x}\| \leq \varepsilon\}$ be the closed Euclidean ball with center at $\mathbf{x}$ and radius ε . We write A^c for the complement of a subset A of $\mathbb{R}^d$. For two random variables Z_1 and Z_2 , the notation

$$Z_1 \leq_{\text{st}} Z_2$$

means that Z_1 is stochastically dominated by Z_2 , that is, for all $t \in \mathbb{R}$,

$$\mathbb{P}\{Z_1 > t\} \leq \mathbb{P}\{Z_2 > t\}.$$

Our first goal is to show that for μ -almost all $\mathbf{x}$, as $k_n/n \rightarrow 0$, the quantity $\max_{i=1, \dots, k_n} \|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\|$ converges to 0 in probability, i.e., for every $\varepsilon > 0$,

$$\lim_{n \rightarrow \infty} \mathbb{P} \left\{ \max_{i=1, \dots, k_n} \|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon \right\} = 0. \quad (2)$$

So, fix such a positive ε . Let δ be a real number in $(0, \varepsilon)$ and γ_n be a positive real number (eventually function of $\mathbf{x}$ and ε) to be determined later. To prove identity (2), we use the following decomposition, which is valid for all $\mathbf{x} \in \mathbb{R}^d$:

$$\begin{aligned} \mathbb{P} \left\{ \max_{i=1, \dots, k_n} \|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon \right\} &\leq \mathbb{P} \left\{ \min_{\substack{i=1, \dots, n \\ \mathbf{x}_i \in \mathcal{B}_{\mathbf{x}, \varepsilon}^c}} \rho_n(\mathbf{x}, \mathbf{X}_i) < \gamma_n \right\} + \mathbb{P} \left\{ \max_{\substack{i=1, \dots, n \\ \mathbf{x}_i \in \mathcal{B}_{\mathbf{x}, \delta}} \rho_n(\mathbf{x}, \mathbf{X}_i) \geq \gamma_n \right\} \\ &\quad + \mathbb{P} \{ \text{Card} \{i = 1, \dots, n : \|\mathbf{X}_i - \mathbf{x}\| \leq \delta\} < k_n \} \\ &:= \mathbf{A} + \mathbf{B} + \mathbf{C}. \end{aligned} \quad (3)$$

The convergence to 0 of each of the three terms above – from which identity (2) immediately follows – are separately analyzed in the next three paragraphs.

Analysis of A. As for now, taking an affine geometry point of view, we keep $\mathbf{x}$ fixed and see it as the origin of the space. Recall that each point in the Euclidean space $\mathbb{R}^d$ (with the origin at $\mathbf{x}$) may be described by its hyperspherical coordinates (see, e.g., [25, Chapter 1]), which consist of a nonnegative radial coordinate r and $d - 1$ angular coordinates $\theta_1, \dots, \theta_{d-1}$, where θ_{d-1} ranges over $[0, 2\pi]$ and the other angles range over $[0, \pi]$ (adaptation of this definition to the cases $d = 1$ and $d = 2$ is clear). For a $(d - 1)$ -dimensional vector $\Theta = (\theta_1, \dots, \theta_{d-1})$ of hyperspherical angles, we let $\mathcal{B}_{\mathbf{x}, \varepsilon}(\Theta)$ be the unique closed ball anchored at $\mathbf{x}$ in the direction Θ and with diameter ε (see Fig. 2 which

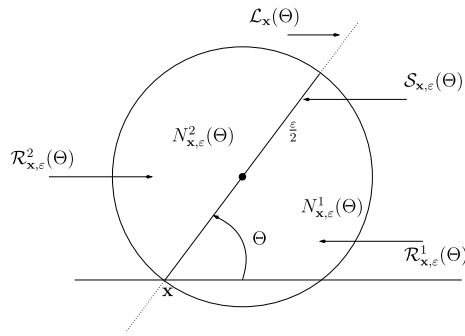


Fig. 2. The ball $\mathcal{B}_{\mathbf{x},\varepsilon}(\Theta)$ and related notation. Illustration in dimension 2.

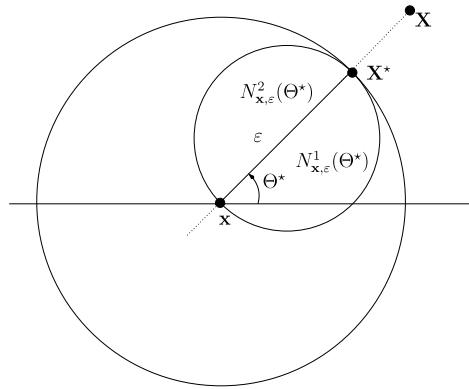


Fig. 3. The ball $\mathcal{B}_{\mathbf{x},\varepsilon}(\Theta^*)$ in dimension 2.

Proposition 1. For μ -almost all $\mathbf{x} \in \mathbb{R}^d$ and all $\varepsilon > 0$ small enough,

$$\mathbb{P} \left\{ \min_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\varepsilon}^c}} \rho_n(\mathbf{x}, \mathbf{X}_i) < \gamma_n \right\} \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

provided

$$\gamma_n = n^d \left(\frac{V_d}{2^{2d+1}} \varepsilon^d f(\mathbf{x}) \right)^{2^{d-1}}.$$

Proof of Proposition 1. Set

$$p_{\mathbf{x},\varepsilon} = \min_{j=1,\dots,2^{d-1}} \inf_{\Theta} \mu \left\{ \mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta) \right\},$$

where the infimum is taken over all possible hyperspherical angles Θ . We know, according to technical [Lemma 1](#), that for μ -almost all $\mathbf{x}$ and all $\varepsilon > 0$ small enough,

$$p_{\mathbf{x},\varepsilon} \geq \frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x}) > 0. \quad (4)$$

Thus, in the rest of the proof, we fix such an $\mathbf{x}$ and assume that ε is small enough so that the inequalities above are satisfied.

Let $\mathbf{X}^*$ be defined as the intersection of the line $(\mathbf{x}, \mathbf{X})$ with $\mathcal{B}_{\mathbf{x},\varepsilon}$, and let Θ^* be the (random) hyperspherical angle corresponding to $\mathbf{X}^*$ (see [Fig. 3](#) for an example in dimension 2). Denote by $N_{\mathbf{x},\varepsilon}(\Theta^*)$ the number of hyperplanes passing through d out of the observations $\mathbf{X}_1, \dots, \mathbf{X}_n$ and cutting the segment $\mathcal{S}_{\mathbf{x},\varepsilon}(\Theta^*)$. We have

$$\begin{aligned} \mathbb{P} \left\{ \min_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\varepsilon}^c}} \rho_n(\mathbf{x}, \mathbf{X}_i) < \gamma_n \right\} &\leq n \mathbb{P} \left\{ \rho_n(\mathbf{x}, \mathbf{X}^*) < \gamma_n \right\} \\ &= n \mathbb{P} \left\{ N_{\mathbf{x},\varepsilon}(\Theta^*) < \gamma_n \right\} \end{aligned}$$

$$\leq n \mathbb{P} \left\{ \frac{N_{\mathbf{x},\varepsilon}^1(\Theta^*) \cdots N_{\mathbf{x},\varepsilon}^{2^{d-1}}(\Theta^*)}{n^{2^{d-1}-d}} < \gamma_n \right\},$$

where the last inequality follows from technical Lemma 2. Thus,

$$\begin{aligned} \mathbb{P} \left\{ \min_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\varepsilon}^c}} \rho_n(\mathbf{x}, \mathbf{X}_i) < \gamma_n \right\} &\leq n \sum_{j=1}^{2^{d-1}} \mathbb{P} \left\{ N_{\mathbf{x},\varepsilon}^j(\Theta^*) < \left(\gamma_n n^{2^{d-1}-d} \right)^{1/2^{d-1}} \right\} \\ &= n \sum_{j=1}^{2^{d-1}} \mathbb{P} \left\{ N_{\mathbf{x},\varepsilon}^j(\Theta^*) < \gamma_n^{1/2^{d-1}} n^{1-d/2^{d-1}} \right\}. \end{aligned}$$

Clearly, conditionally on Θ^* , each $N_{\mathbf{x},\varepsilon}^j(\Theta^*)$ satisfies

$$\text{Binomial}(n, p_{\mathbf{x},\varepsilon}) \leq_{\text{st}} N_{\mathbf{x},\varepsilon}^j(\Theta^*)$$

and consequently, by inequality (4),

$$\text{Binomial} \left(n, \frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x}) \right) \leq_{\text{st}} N_{\mathbf{x},\varepsilon}^j(\Theta^*).$$

Thus, for each $j = 1, \dots, 2^{d-1}$, by Hoeffding's inequality for binomial random variables [21], we are led to

$$\begin{aligned} \mathbb{P} \left\{ N_{\mathbf{x},\varepsilon}^j(\Theta^*) < \gamma_n^{1/2^{d-1}} n^{1-d/2^{d-1}} \right\} &= \mathbb{E} \left[\mathbb{P} \left\{ N_{\mathbf{x},\varepsilon}^j(\Theta^*) < \gamma_n^{1/2^{d-1}} n^{1-d/2^{d-1}} \mid \Theta^* \right\} \right] \\ &\leq \exp \left[-2 \left(\gamma_n^{1/2^{d-1}} n^{1-d/2^{d-1}} - n \frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x}) \right)^2 / n \right] \end{aligned}$$

as soon as $\gamma_n^{1/2^{d-1}} n^{1-d/2^{d-1}} < n \frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x})$. Therefore, taking

$$\gamma_n = n^d \left(\frac{V_d}{2^{2d+1}} \varepsilon^d f(\mathbf{x}) \right)^{2^{d-1}},$$

we obtain

$$\mathbb{P} \left\{ \min_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\varepsilon}^c}} \rho_n(\mathbf{x}, \mathbf{X}_i) < \gamma_n \right\} \leq 2^{d-1} n \exp \left[-n \left(\frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x}) \right)^2 / 2 \right].$$

The upper bound goes to 0 as $n \rightarrow \infty$. $\square$

Analysis of B. Consistency of the second term in inequality (3) is established in the following proposition.

Proposition 2. For μ -almost all $\mathbf{x} \in \mathbb{R}^d$, all $\varepsilon > 0$ and all $\delta > 0$ small enough,

$$\mathbb{P} \left\{ \max_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\delta}}} \rho_n(\mathbf{x}, \mathbf{X}_i) \geq \gamma_n \right\} \rightarrow 0 \quad \text{as } n \rightarrow \infty,$$

provided

$$\gamma_n = n^d \left(\frac{V_d}{2^{2d+1}} \varepsilon^d f(\mathbf{x}) \right)^{2^{d-1}}. \quad (5)$$

Proof of Proposition 2. Fix $\mathbf{x}$ in a set of μ -measure 1 such that $f(\mathbf{x}) > 0$ and denote by $N_{\mathbf{x},\delta}$ the number of hyperplanes that cut the ball $\mathcal{B}_{\mathbf{x},\delta}$. Clearly,

$$\mathbb{P} \left\{ \max_{\substack{i=1,\dots,n \\ \mathbf{X}_i \in \mathcal{B}_{\mathbf{x},\delta}}} \rho_n(\mathbf{x}, \mathbf{X}_i) \geq \gamma_n \right\} \leq \mathbb{P} \{ N_{\mathbf{x},\delta} \geq \gamma_n \}.$$

Observe that, with probability 1,

$$N_{\mathbf{x},\delta} = \sum_{1 \leq i_1 < \dots < i_d \leq n} \mathbf{1}_{\{\mathcal{H}(\mathbf{x}_{i_1}, \dots, \mathbf{x}_{i_d}) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset\}},$$

whence, since $\mathbf{X}_1, \dots, \mathbf{X}_n$ are identically distributed,

$$\begin{aligned}\mathbb{E}[N_{\mathbf{x},\delta}] &= \binom{n}{d} \mathbb{P}\{\mathcal{H}(\mathbf{X}_1, \dots, \mathbf{X}_d) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset\} \\ &\leq \frac{n^d}{d!} \mathbb{P}\{\mathcal{H}(\mathbf{X}_1, \dots, \mathbf{X}_d) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset\}.\end{aligned}$$

Consequently, given the choice (5) for γ_n and the result of technical Lemma 3, it follows that

$$\mathbb{E}[N_{\mathbf{x},\delta}] < \gamma_n/2$$

for all δ small enough, independently of n . Thus, using the bounded difference inequality [23], we obtain, still with the choice

$$\begin{aligned}\gamma_n &= n^d \left(\frac{V_d}{2^{2d+1}} \varepsilon^d f(\mathbf{x}) \right)^{2^{d-1}}, \\ \mathbb{P}\{N_{\mathbf{x},\delta} \geq \gamma_n\} &\leq \mathbb{P}\{N_{\mathbf{x},\delta} - \mathbb{E}[N_{\mathbf{x},\delta}] \geq \gamma_n/2\} \\ &\leq \exp\left(-2 \frac{(\gamma_n/2)^2}{n^{2d-1}}\right) \\ &= \exp\left[-\left(\frac{V_d}{2^{2d+1}} \varepsilon^d f(\mathbf{x})\right)^{2^d} n/2\right].\end{aligned}$$

This upper bound goes to zero as n tends to infinity, and this concludes the proof of the proposition. $\square$

Analysis of C. To achieve the proof of identity (2), it remains to show that the third and last term of (3) converges to 0. This is done in the following proposition.

Proposition 3. Assume that $k_n/n \rightarrow 0$ as $n \rightarrow \infty$. Then, for μ -almost all $\mathbf{x} \in \mathbb{R}^d$ and all $\delta > 0$,

$$\mathbb{P}\{\text{Card}\{i = 1, \dots, n : \|\mathbf{X}_i - \mathbf{x}\| \leq \delta\} < k_n\} \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

Proof of Proposition 3. Recall that the collection of all $\mathbf{x}$ with $\mu(\mathcal{B}_{\mathbf{x},\tau}) > 0$ for all $\tau > 0$ is called the support of μ , and note that it may alternatively be defined as the smallest closed subset of $\mathbb{R}^d$ of μ -measure 1 [32, Chapter 2]. Thus, fix $\mathbf{x}$ in the support of μ and set

$$p_{\mathbf{x},\delta} = \mathbb{P}\{\mathbf{X} \in \mathcal{B}_{\mathbf{x},\delta}\},$$

so that $p_{\mathbf{x},\delta} > 0$. Then the following chain of inequalities is valid:

$$\begin{aligned}\mathbb{P}\{\text{Card}\{i = 1, \dots, n : \|\mathbf{X}_i - \mathbf{x}\| \leq \delta\} < k_n\} &= \mathbb{P}\{\text{Binomial}(n, p_{\mathbf{x},\delta}) < k_n\} \\ &\leq \mathbb{P}\{\text{Binomial}(n, p_{\mathbf{x},\delta}) \leq np_{\mathbf{x},\delta}/2\} \\ &\quad (\text{for all } n \text{ large enough, since } k_n/n \text{ tends to } 0) \\ &\leq \exp(-np_{\mathbf{x},\delta}^2/2),\end{aligned}$$

where the last inequality follows from Hoeffding's inequality [21]. This terminates the proof of Proposition 3. $\square$

We have proved so far that, for μ -almost all $\mathbf{x}$, as $k_n/n \rightarrow 0$, the quantity $\max_{i=1,\dots,k_n} \|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\|$ converges to 0 in probability. By the elementary inequality

$$\mathbb{E}\left[\frac{1}{k_n} \sum_{i=1}^{k_n} \mathbf{1}_{\{\|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon\}}\right] \leq \mathbb{P}\left\{\max_{i=1,\dots,k_n} \|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon\right\},$$

it immediately follows that, for such an $\mathbf{x}$,

$$\mathbb{E}\left[\frac{1}{k_n} \sum_{i=1}^{k_n} \mathbf{1}_{\{\|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon\}}\right] \rightarrow 0 \tag{6}$$

provided $k_n/n \rightarrow 0$. We are now ready to complete the proof of Theorem 1.

Fix $\mathbf{x}$ in a set of μ -measure 1 such that consistency (6) holds and r is continuous at $\mathbf{x}$ (this is possible by the assumption on r). Because $|a + b|^p \leq 2^{p-1}(|a|^p + |b|^p)$ for $p \geq 1$, we see that

$$\mathbb{E}|r_n(\mathbf{x}) - r(\mathbf{x})|^p \leq 2^{p-1} \mathbb{E}\left|\frac{1}{k_n} \sum_{i=1}^{k_n} [Y_{(i)}(\mathbf{x}) - r(\mathbf{X}_{(i)}(\mathbf{x}))]\right|^p + 2^{p-1} \mathbb{E}\left|\frac{1}{k_n} \sum_{i=1}^{k_n} [r(\mathbf{X}_{(i)}(\mathbf{x})) - r(\mathbf{x})]\right|^p.$$

Thus, by Jensen's inequality,

$$\begin{aligned}\mathbb{E} |r_n(\mathbf{x}) - r(\mathbf{x})|^p &\leq 2^{p-1} \mathbb{E} \left| \frac{1}{k_n} \sum_{i=1}^{k_n} [Y_{(i)}(\mathbf{x}) - r(\mathbf{X}_{(i)}(\mathbf{x}))] \right|^p + 2^{p-1} \mathbb{E} \left[\frac{1}{k_n} \sum_{i=1}^{k_n} |r(\mathbf{X}_{(i)}(\mathbf{x})) - r(\mathbf{x})|^p \right] \\ &:= 2^{p-1} \mathbf{I}_n + 2^{p-1} \mathbf{J}_n.\end{aligned}$$

Firstly, for arbitrary $\varepsilon > 0$, we have

$$\mathbf{J}_n = \mathbb{E} \left[\frac{1}{k_n} \sum_{i=1}^{k_n} |r(\mathbf{X}_{(i)}(\mathbf{x})) - r(\mathbf{x})|^p \mathbf{1}_{\{\|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon\}} \right] + \mathbb{E} \left[\frac{1}{k_n} \sum_{i=1}^{k_n} |r(\mathbf{X}_{(i)}(\mathbf{x})) - r(\mathbf{x})|^p \mathbf{1}_{\{\|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| \leq \varepsilon\}} \right],$$

whence

$$\mathbf{J}_n \leq 2^p \zeta^p \mathbb{E} \left[\frac{1}{k_n} \sum_{i=1}^{k_n} \mathbf{1}_{\{\|\mathbf{X}_{(i)}(\mathbf{x}) - \mathbf{x}\| > \varepsilon\}} \right] + \left[\sup_{\mathbf{y} \in \mathbb{R}^d: \|\mathbf{y} - \mathbf{x}\| \leq \varepsilon} |r(\mathbf{y}) - r(\mathbf{x})| \right]^p \quad (\text{since } |Y| \leq \zeta).$$

The first term on the right-hand side of the latter inequality tends to 0 by (6) as $k_n/n \rightarrow 0$, whereas the rightmost one can be made arbitrarily small as $\varepsilon \rightarrow 0$ since r is continuous at $\mathbf{x}$. This proves that $\mathbf{J}_n \rightarrow 0$ as $n \rightarrow \infty$.

Next, by successive applications of inequalities of Marcinkiewicz and Zygmund [22] (see also [33, pp. 59–60]), we have for some positive constant C_p depending only on p ,

$$\begin{aligned}\mathbf{I}_n &\leq C_p \mathbb{E} \left[\frac{1}{k_n^2} \sum_{i=1}^{k_n} |Y_{(i)}(\mathbf{x}) - r(\mathbf{X}_{(i)}(\mathbf{x}))|^2 \right]^{p/2} \\ &\leq \frac{(2\zeta)^p C_p}{k_n^{p/2}} \quad (\text{since } |Y| \leq \zeta).\end{aligned}$$

Consequently, $\mathbf{I}_n \rightarrow 0$ as $k_n \rightarrow \infty$, and this concludes the proof of the theorem.

4. Some technical lemmas

The notation of this section is identical to that of Section 3. In particular, it is assumed throughout that $\mathbf{X}$ has a probability density f with respect to the Lebesgue measure λ on $\mathbb{R}^d$. This requirement implies that any collection $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}$ ($1 \leq i_1 < i_2 < \dots < i_d \leq n$) of d points among $\mathbf{X}_1, \dots, \mathbf{X}_n$ define with probability 1 a unique hyperplane $\mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})$ in $\mathbb{R}^d$. Recall finally that, for $\mathbf{x} \in \mathbb{R}^d$ and $\varepsilon > 0$, we set

$$p_{\mathbf{x},\varepsilon} = \min_{j=1,\dots,2^{d-1}} \inf_{\Theta} \mu \{ \mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta) \},$$

where the infimum is taken over all possible hyperspherical angles Θ , and the regions $\mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta)$, $j = 1, \dots, 2^{d-1}$, define a partition of the ball $\mathcal{B}_{\mathbf{x},\varepsilon}(\Theta)$. Recall also that the numbers of sample points falling in each of these regions are denoted by $N_{\mathbf{x},\varepsilon}^1(\Theta), \dots, N_{\mathbf{x},\varepsilon}^{2^{d-1}}(\Theta)$. For a better understanding of the next lemmas, the reader should refer to Figs. 2 and 3.

Lemma 1. For μ -almost all $\mathbf{x} \in \mathbb{R}^d$ and all $\varepsilon > 0$ small enough,

$$p_{\mathbf{x},\varepsilon} \geq \frac{V_d}{2^{2d}} \varepsilon^d f(\mathbf{x}) > 0.$$

Proof of Lemma 1. We let $\mathbf{x}$ be a Lebesgue point of f , that is, an $\mathbf{x}$ such that for any collection $\mathcal{A}$ of subsets of $\mathcal{B}_{\mathbf{0},1}$ with the property that for all $A \in \mathcal{A}$, $\lambda(A) \geq c\lambda(\mathcal{B}_{\mathbf{0},1})$ for some fixed $c > 0$,

$$\limsup_{\varepsilon \rightarrow 0} \sup_{A \in \mathcal{A}} \left| \frac{\int_{\mathbf{x} + \varepsilon A} f(\mathbf{y}) d\mathbf{y}}{\lambda\{\mathbf{x} + \varepsilon A\}} - f(\mathbf{x}) \right| = 0, \quad (7)$$

where $\mathbf{x} + \varepsilon A = \{\mathbf{y} \in \mathbb{R}^d : (\mathbf{y} - \mathbf{x})/\varepsilon \in A\}$. As f is a density, we know that μ -almost all $\mathbf{x}$ satisfy this property (see,

and observing that

$$\lambda \{ \mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta) \} = \frac{V_d}{2^{d-1}} \left(\frac{\varepsilon}{2} \right)^d,$$

we may write, for each $j = 1, \dots, 2^{d-1}$,

$$\begin{aligned} \left| \frac{2^{d-1} p_{\mathbf{x},\varepsilon}^j}{V_d(\varepsilon/2)^d} - f(\mathbf{x}) \right| &= \left| \inf_{\Theta} \frac{\mu \{ \mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta) \}}{\lambda \{ \mathcal{R}_{\mathbf{x},\varepsilon}^j(\Theta) \}} - f(\mathbf{x}) \right| \\ &= \left| \inf_{A \in \mathcal{A}} \frac{\int_{\mathbf{x}+\varepsilon A} f(\mathbf{y}) d\mathbf{y}}{\lambda \{ \mathbf{x} + \varepsilon A \}} - f(\mathbf{x}) \right| \\ &\leq \sup_{A \in \mathcal{A}} \left| \frac{\int_{\mathbf{x}+\varepsilon A} f(\mathbf{y}) d\mathbf{y}}{\lambda \{ \mathbf{x} + \varepsilon A \}} - f(\mathbf{x}) \right|. \end{aligned}$$

The conclusion follows from identity (7). $\square$

Lemma 2. Fix $\mathbf{x} \in \mathbb{R}^d$, $\varepsilon > 0$ and $\Theta \in [0, \pi]^{d-2} \times [0, 2\pi)$. Let $N_{\mathbf{x},\varepsilon}(\Theta)$ be the number of hyperplanes passing through d out of the observations $\mathbf{X}_1, \dots, \mathbf{X}_n$ and cutting the segment $S_{\mathbf{x},\varepsilon}(\Theta)$. Then, with probability 1,

$$N_{\mathbf{x},\varepsilon}(\Theta) \geq \frac{N_{\mathbf{x},\varepsilon}^1(\Theta) \cdots N_{\mathbf{x},\varepsilon}^{2^{d-1}}(\Theta)}{n^{2^{d-1}-d}}.$$

Proof of Lemma 2. If one of the $N_{\mathbf{x},\varepsilon}^j(\Theta)$ ($j = 1, \dots, 2^{d-1}$) is zero, then the result is trivial. Thus, in the rest of the proof, we suppose that each $N_{\mathbf{x},\varepsilon}^j(\Theta)$ is positive and note that this implies $n \geq 2^{d-1}$.

Pick sequentially 2^{d-1} observations, say $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{2^{d-1}}}$, in the 2^{d-1} regions $\mathcal{R}_{\mathbf{x},\varepsilon}^1(\Theta), \dots, \mathcal{R}_{\mathbf{x},\varepsilon}^{2^{d-1}}(\Theta)$. By construction, the polytope defined by these 2^{d-1} points cuts the axis $\mathcal{L}_{\mathbf{x},\varepsilon}(\Theta)$. Consequently, with probability 1, any hyperplane drawn according to d out of these 2^{d-1} points cuts the segment $S_{\mathbf{x},\varepsilon}(\Theta)$. The desired result follows by observing that there are exactly $N_{\mathbf{x},\varepsilon}^1(\Theta) \cdots N_{\mathbf{x},\varepsilon}^{2^{d-1}}(\Theta)$ such polytopes. $\square$

Lemma 3. For $1 \leq i_1 < \dots < i_d \leq n$, let $\mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})$ be the hyperplane passing through d out of the observations $\mathbf{X}_1, \dots, \mathbf{X}_n$. Then, for all $\mathbf{x} \in \mathbb{R}^d$,

$$\mathbb{P} \{ \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset \} \rightarrow 0 \quad \text{as } \delta \downarrow 0.$$

Proof of Lemma 3. Given two hyperplanes $\mathcal{H}$ and $\mathcal{H}'$ in $\mathbb{R}^d$, we denote by $\Phi(\mathcal{H}, \mathcal{H}')$ the (dihedral) angle between $\mathcal{H}$ and $\mathcal{H}'$. Recall that $\Phi(\mathcal{H}, \mathcal{H}') \in [0, \pi]$ and that it is defined as the angle between the corresponding normal vectors.

Fix $1 \leq i_1 < \dots < i_d \leq n$. Let $\mathcal{E}_\delta$ be the event

$$\mathcal{E}_\delta = \{ \|\mathbf{X}_{i_j} - \mathbf{x}\| > \delta : j = 1, \dots, d-1 \},$$

and let $\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}})$ be the hyperplane passing through $\mathbf{x}$ and the $d-1$ points $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}$. Clearly, on $\mathcal{E}_\delta$, the event $\{ \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset \}$ is the same as

$$\{ \Phi(\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}), \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})) \leq \Phi_\delta \},$$

where Φ_δ is the angle formed by $\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}})$ and the hyperplane going through $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}$ and tangent to $\mathcal{B}_{\mathbf{x},\delta}$ (see Fig. 4 for an example in dimension 2).

Thus, with this notation, we may write

$$\begin{aligned} \mathbb{P} \{ \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \cap \mathcal{B}_{\mathbf{x},\delta} \neq \emptyset \} &\leq \mathbb{P} \{ \mathcal{E}_\delta^c \} + \mathbb{P} \{ \mathcal{H}(\mathbf{X}_{i_1}, \$$

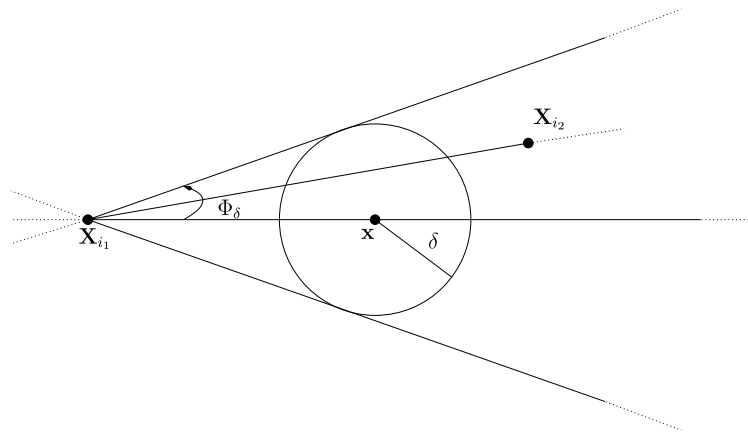


Fig. 4. The hyperplanes $\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}})$ and $\mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d})$, and the angle Φ_δ . Illustration in dimension 2.

Thus, writing

$$\begin{aligned} & \mathbb{P} \left\{ \Phi \left(\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}), \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \right) \leq \Phi_\delta, \mathcal{E}_\delta \right\} \\ &= \mathbb{E} \left[\mathbf{1}_{\mathcal{E}_\delta} \mathbb{P} \left\{ \Phi \left(\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}), \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \right) \leq \Phi_\delta \mid \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}} \right\} \right] \end{aligned}$$

and noting that, on the event $\mathcal{E}_\delta$, for fixed $\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}$, $\Phi_\delta \downarrow 0$ as $\delta \downarrow 0$, we conclude by the Lebesgue's dominated convergence theorem that

$$\mathbb{P} \left\{ \Phi \left(\mathcal{H}(\mathbf{x}, \mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_{d-1}}), \mathcal{H}(\mathbf{X}_{i_1}, \dots, \mathbf{X}_{i_d}) \right) \leq \Phi_\delta \right\} \rightarrow 0 \quad \text{as } \delta \downarrow 0. \quad \square$$

Acknowledgments

We thank two anonymous referees for valuable comments and insightful suggestions.

The first author's research was partially supported by the French National Research Agency under grant ANR-09-BLAN-0051-02 "CLARA". Research carried out within the INRIA project "CLASSIC" hosted by Ecole Normale Supérieure and CNRS. The second author's research was sponsored by NSERC Grant A3456 and FQRNT Grant 90-ER-0291. The third author's research was sponsored by NSERC Grant RGPIN 402438-2011. The fourth author's research was sponsored by NSERC Grant N00118.

References

- [1] T. Anderson, Some nonparametric multivariate procedures based on statistically equivalent blocks, in: P. Krishnaiah (Ed.), *Multivariate Analysis*, Academic Press, New York, 1966, pp. 5–27.
- [2] G. Biau, L. Devroye, On the layered nearest neighbour estimate, the bagged nearest neighbour estimate and the random forest method in regression and classification, *Journal of Multivariate Analysis* 101 (2010) 2499–2518.
- [3] B. Chakraborty, P. Chaudhuri, H. Oja, Operating transformation retransformation on spatial median and angle test, *Statistica Sinica* 8 (1998) 767–784.
- [4] B. Chazelle, Cutting hyperplanes for divide-and-conquer, *Discrete and Computational Geometry* 9 (1993) 145–158.
- [5] T. Cover, Estimation by the nearest neighbor rule, *IEEE Transactions on Information Theory* 14 (1968) 50–55.
- [6] T. Cover, Rates of convergence for nearest neighbor procedures, in: *Proceedings of the Hawaii International Conference on Systems Sciences*, Honolulu, 1968, pp. 413–415.
- [7] T. Cover, P. Hart, Nearest neighbor pattern classification, *IEEE Transactions on Information Theory* 13 (1967) 21–27.
- [8] L. Devroye, A universal k -nearest neighbor procedure in discrimination, in: B. Dasarthy (Ed.), *Nearest Neighbor Pattern Classification Techniques*, IEEE Computer Society Press, Los Alamitos, 1991, pp. 101–106.
- [9] L. Devroye, L. Györfi, G. Lugosi, *A Probabilistic Theory of Pattern Recognition*, Springer-Verlag, New York, 1996.
- [10] L. Devroye, A. Krzyżak, New multivariate product density estimators, *Journal of Multivariate Analysis* 82 (2002) 88–110.
- [11] 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.
- [12] E. Fix, J. Hodges, Discriminatory analysis: small sample performance, Technical Report 11, Project Number 21-49-004, USAF School of Aviation Medicine, Randolph Field, Texas, 1952.
- [13] M. Gessaman, A consistent nonparametric multivariate density estimator based on statistically equivalent blocks, *The Annals of Mathematical Statistics* 41 (1970) 1344–1346.
- [14] M. Gessaman, P. Gessaman, A comparison of some multivariate discrimination procedures, *Journal of the American Statistical Association* 67 (1972) 468–472.
- [15] L. Gordon, R. Olshen, Asymptotically efficient solutions to the classification problem, *The Annals of Statistics* 6 (1978) 515–533.
- [16] L. Gordon, R. Olshen, Consistent nonparametric regression from recursive partitioning schemes, *Journal of Multivariate Analysis* 10 (1980) 611–627.
- [17] L. Györfi, M. Kohler, A. Krzyżak, H. Walk, *A Distribution-Free Theory of Nonparametric Regression*, Springer-Verlag, New York, 2002.
- [18] M. Hallin, D. Paindaveine, Optimal tests for multivariate location based on interdirections and pseudo-Mahalanobis ranks, *The Annals of Statistics* 30 (2002) 1103–1133.
- [19] T. Hettmansperger, J. Möttönen, H. Oja, The geometry of the affine invariant multivariate sign and rank methods, *Journal of Nonparametric Statistics* 11 (1998) 271–285.

- [20] T. Hettmansperger, R. Randles, A practical affine equivariant multivariate median, *Biometrika* 89 (2002) 851–860.
- [21] W. Hoeffding, Probability inequalities for sums of bounded random variables, *Journal of the American Statistical Association* 58 (1963) 13–30.
- [22] J. Marcinkiewicz, A. Zygmund, Sur les fonctions indépendantes, *Fundamenta Mathematicae* 29 (1937) 60–90.
- [23] C. McDiarmid, On the method of bounded differences, in: J. Siemons (Ed.), *Surveys in Combinatorics*, 1989, in: *London Mathematical Society Lecture Note Series*, vol. 141, Cambridge University Press, 1989, pp. 148–188.
- [24] S. Meiser, Point location in arrangements of hyperplanes, *Information and Computation* 106 (1993) 286–303.
- [25] K. Miller, *Multidimensional Gaussian Distributions*, Wiley, New York, 1964.
- [26] H. Oja, Affine invariant multivariate sign and rank tests and corresponding estimates: a review, *Scandinavian Journal of Statistics* 26 (1999) 319–343.
- [27] H. Oja, D. Paindaveine, Optimal signed-rank tests based on hyperplanes, *Journal of Statistical Planning and Inference* 135 (2005) 300–323.
- [28] E. Ollila, T. Hettmansperger, H. Oja, Estimates of regression coefficients based on sign covariance matrix, *Journal of the Royal Statistical Society Series B* 64 (2002) 447–466.
- [29] E. Ollila, H. Oja, V. Koivunen, Estimates of regression coefficients based on lift rank covariance matrix, *Journal of the American Statistical Association* 98 (2003) 90–98.
- [30] R. Olshen, Comments on a paper by C.J. Stone, *The Annals of Statistics* 5 (1977) 632–633.
- [31] D. Paindaveine, G. Van Bever, Nonparametric consistent depth-based classifiers, ECARES working paper 2012-014, Brussels, 2012.
- [32] K. Parthasarathy, *Probability Measures on Metric Spaces*, AMS Chelsea Publishing, Providence, 2005.
- [33] V. Petrov, *Sums of Independent Random Variables*, Springer-Verlag, Berlin, 1975.
- [34] C. Quesenberry, M. Gessaman, Nonparametric discrimination using tolerance regions, *The Annals of Mathematical Statistics* 39 (1968) 664–673.
- [35] R. Randles, A distribution-free multivariate sign test based on interdirections, *Journal of the American Statistical Association* 84 (1989) 1045–1050.
- [36] M. Samanta, A note on uniform strong convergence of bivariate density estimates, *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete* 28 (1974) 85–88.
- [37] C. Stone, Consistent nonparametric regression, *The Annals of Statistics* 5 (1977) 595–645.
- [38] D. Tyler, A distribution-free M -estimator of multivariate scatter, *The Annals of Statistics* 15 (1987) 234–251.
- [39] R. Wheeden, A. Zygmund, *Measure and Integral. An Introduction to Real Analysis*, Marcel Dekker, New York, 1977.