

COVERAGE OF CREDIBLE INTERVALS IN BAYESIAN MULTIVARIATE ISOTONIC REGRESSION

BY KANG WANG^a AND SUBHASHIS GHOSAL^b

Department of Statistics, North Carolina State University, ^akwang22@ncsu.edu, ^bsghosal@stat.ncsu.edu

We consider the nonparametric multivariate isotonic regression problem, where the regression function is assumed to be nondecreasing with respect to each predictor. Our goal is to construct a Bayesian credible interval for the function value at a given interior point with assured limiting frequentist coverage. A natural prior on the regression function is given by a random step function with a suitable prior on increasing step-heights, but the resulting posterior distribution is hard to analyze theoretically due to the complicated order restriction on the coefficients. We instead put a prior on unrestricted step-functions, but make inference using the induced posterior measure by an “immersion map” from the space of unrestricted functions to that of multivariate monotone functions. This allows for maintaining the natural conjugacy for posterior sampling. A natural immersion map to use is a projection with respect to a distance function, but in the present context, a block isotonicization map is found to be more useful. The approach of using the induced “immersion posterior” measure instead of the original posterior to make inference provides a useful extension of the Bayesian paradigm, particularly helpful when the model space is restricted by some complex relations. We establish a key weak convergence result for the posterior distribution of the function at a point in terms of some functional of a multiindexed Gaussian process that leads to an expression for the limiting coverage of the Bayesian credible interval. Analogous to a recent result for univariate monotone functions, we find that the limiting coverage is slightly higher than the credibility, the opposite of a phenomenon observed in smoothing problems. Interestingly, the relation between credibility and limiting coverage does not involve any unknown parameter. Hence, by a recalibration procedure, we can get a predetermined asymptotic coverage by choosing a suitable credibility level smaller than the targeted coverage, and thus also shorten the credible intervals.

1. Introduction. Nonparametric inference often involves a regression function or a density function in modeling. Commonly, a smoothness assumption on a function of interest is imposed, but in some applications, qualitative information, such as monotonicity, unimodality and convexity, on the shape of the function may be available. This leads to a control on the complexity of the function space analogous to what a smoothness assumption does, allowing convergence without requiring the latter. Monotonicity is the simplest and the most extensively studied shape restriction, especially in the univariate case. In regression analysis, this problem is commonly referred to as isotonic regression when the conditional mean function of the response variable is assumed to be nondecreasing. Starting from the early works on monotone shape-restricted problems, such as [1, 10], research on non-Bayesian approaches, mainly on the least squares estimator (LSE) and the nonparametric maximum likelihood estimator (MLE), has been fruitful; see [4, 31, 32, 46]. Assuming a nonzero derivative, the pointwise asymptotic distribution of the MLE or the LSE turns out to be the rescaled Chernoff distribution, that is, the minimizer of a quadratically drifted standard two-sided

Received February 2022; revised May 2023.

MSC2020 subject classifications. Primary 62G08; secondary 62F15, 62G05, 62G20.

Key words and phrases. Isotonic regression, credible interval, limiting coverage, Gaussian process, Block isotonicization, immersion posterior.

Brownian motion [11, 32, 45, 56]. The same limiting distribution can also be found in other problems where monotonicity is implied, such as the monotone hazard rate estimation with randomly right-censored observations in survival analysis [39, 40], and many statistical inverse problems, including the current status model and deconvolution problems [33]. Global properties of shape-restricted estimators were also studied extensively [32, 42]. The convergence rates and limiting distributional behaviors of $\mathbb{L}_p$ - and $\mathbb{L}_\infty$ -distance between monotone shape-restricted estimators and the true function were investigated by [26, 27]. Nonasymptotic risk bounds for the LSE under a monotone shape restriction were derived by [5, 18, 57]. Testing for monotonicity was addressed in [25, 29, 30, 35].

The Bayesian approach to shape-restricted problems was also explored, albeit to a lesser extent. Neelon and Dunson [44] applied piecewise linear structures to the regression function, and put monotone restrictions on the priors for the slope values. Cai and Dunson [12] proposed a linear spline model and added an initial Markov random field prior to the coefficients, and then the monotone constraint was incorporated by considering the relation between the nonnegative slopes and coefficients. Wang [55] adopted the free-knot cubic regression spline model, converted the shape restriction to the coefficients and then projected the unconstrained coefficients with conventional priors to the target set, inducing the constrained priors. Shively et al. [52] also used Bayesian splines with constrained normal priors on the coefficients to comply with the monotone shape restriction. Lin and Dunson [43] addressed this problem by using a Gaussian process prior and projected unconstrained posterior samples to the monotone function class by a min-max formula. Chakraborty and Ghosal [14, 15, 17] also used the idea of projection-posterior, making the investigation of frequentist limiting coverage of credible sets possible. Salomond [48] used a mixture representation of a nonincreasing density on $[0, \infty)$ and obtained the nearly minimax posterior contraction rate for both a Dirichlet process and a finite mixture prior on the mixing distribution. Bayesian tests for monotonicity were developed by [15, 17, 49]. A Bayesian credible interval with assured frequentist coverage for a monotone regression quantile, and an accelerated rate of contraction for it using a two-stage sampling were obtained by [16].

Multivariate monotone function estimation was also studied in the literature. Non-Bayesian works focused on the construction of the LSE with respect to various partial orderings on the domain; see [4, 46]. Only the consistency of the isotonic estimator was known until a recent rise in interest in multivariate shape-restricted problems. In a multivariate isotonic regression model, the $\mathbb{L}_2$ -risk of the LSE, respectively for $d = 2$ and for a general dimension d was studied by [19, 37]. They found that the LSE achieved the optimal minimax rate up to logarithmic factors, and adapted to the parametric rate for a piecewise constant true regression function only when $d \leq 2$. Han [36] confirmed that the global empirical risk minimizer is indeed rate-optimal in some set structured models even with rapidly diverging entropy integral, and thus gave a simpler proof for the optimal convergence rate of the LSE in the multivariate isotonic regression. Deng and Zhang [22] investigated a block-estimator proposed by [28] and obtained an $\mathbb{L}_q$ -risk bound. This is minimax rate optimal and adapts to the parametric rate up to a logarithmic factor when the true regression function is piecewise constant. Pointwise distributional limits for the block-estimator were obtained by [38], which lays the foundation for subsequent inference.

The Bayesian approach to multivariate isotonic regression is much less developed. Saarele and Arjas [47] proposed a Bayesian approach to this problem based on marked point processes and resulted in piecewise constant realizations of the regression function. Lin and Dunson [43] mentioned that the method of projecting the Gaussian process posterior can also be applied in the regression surface case. Nonetheless, the theoretical studies presented in those works are either lacking or inadequate.

The construction of confidence regions in function estimation problems was studied by many authors, mostly in smoothness regimes. For shape-restricted problems, confidence regions using limit theory were constructed by [13, 23, 24, 50]. The natural bootstrap method does not lead to a valid confidence interval for the function value at a point, but a modified bootstrap method works [41, 51]. Confidence intervals by inverting the acceptance region of a likelihood ratio test for the value of a monotone function at a point were obtained by [2, 3, 34]. This approach has the advantage that no additional nuisance parameters need to be estimated and plugged into the limit distribution. Deng et al. [21] constructed a confidence interval for multivariate monotone regression from a pivotal limit result for the block-estimator of [22] relying on the limiting distributional theory by [38].

In this paper, we consider a Bayesian approach to the multivariate monotone regression problem. Our objective is to construct a Bayesian credible interval for the function value at an interior point and study its frequentist coverage. As in the univariate problem studied by [14], we ignore the shape restriction at the prior stage and put a prior on random step functions without an order restriction, retaining the posterior conjugacy. We then apply corrections to posterior samples using monotone mapping. This induces a posterior distribution supported on the space of multivariate monotone functions, by which we can obtain credible intervals. However, contrasting with the univariate case, the projection-posterior acquired through the minimization of the empirical $\mathbb{L}_2$ -metric does not possess a limiting distribution. This is due to the fact that the partial sum process, which characterizes the empirical $\mathbb{L}_2$ -projection, is not tight in the limit; see [38]. We also note that the non-Bayesian confidence interval constructed in [21] is also not obtained by distance minimization but by a block max-min procedure. We instead use a map related to the block max-min operation to enforce multivariate monotonicity on posterior samples. As the map immerses a general function into the space of monotone functions, such a map will be referred to as an immersion map, and the induced posterior will be termed an immersion posterior.

The rest of the paper is organized as follows. In the next section, we introduce the notion of an immersion posterior distribution, which will be used to address the problem under study, and is very helpful for similar Bayesian problems with complicated restrictions on the parameter space. In Section 3, we introduce the model and assumptions, construct a prior distribution, and obtain the immersion posterior used to make an inference. Our main results are presented in Section 4. We derive the weak limit of the scaled and centered pointwise immersion posterior distribution. Based on the limit theory, we compute the asymptotic frequentist coverage of credible intervals. Numerical results are present in Section 5. We include all the proofs of the main theorems in Section 6. Proofs of all auxiliary lemmas and propositions are provided in the Supplementary Material.

2. Immersion posterior. Consider a general statistical model with observation $X \sim P_\theta$, where $\theta \in \Theta_0$. Suppose that the parameter space Θ_0 is a complicated subset of a larger, but simpler to represent, set Θ . This is often the case for shape-restricted inference, where structural constraints, such as monotonicity, convexity and log-concavity, are imposed on a regression function or a density function. In differential equation models, the parameter space is implicitly described as the set of solutions of a system of ordinary or partial differential equations involving some unknown parameters. In a vector autoregressive process, the set of autoregression coefficients leading to stationary processes may be the parameter space of interest, but it is described by many complicated constraints. Because of the complicated restrictions on Θ_0 , a prior for θ with support on Θ_0 may be hard to construct, and the corresponding posterior may be difficult to compute. More importantly, the corresponding posterior may be hard to analyze from a frequentist perspective. This may be particularly important for studying delicate properties such as the limiting coverage of a Bayesian credible region.

Often, the distribution P_θ makes sense for any $\theta \in \Theta$, so that Θ_0 can be embedded in Θ , keeping the statistical problem meaningful. For shape-restricted models, this becomes the standard nonparametric regression or the density estimation problem. A differential equation model also embeds in a nonparametric regression model. A prior distribution Π may be specified on Θ , initially disregarding the restriction of θ to Θ_0 . This is typically a standard problem, and often a conjugate prior distribution can be identified. The resulting posterior distribution $\Pi(\cdot|X)$ thus resides in the whole of Θ , and hence is not appropriate to make an inference about θ , which is known to live in Θ_0 . The requirement can be met by considering the random measure induced by a mapping ι from Θ to Θ_0 , in that, we consider the random measure $\Pi^*(B|X) = \Pi(\iota(\theta) \in B|X)$ to make an inference on θ . The map ι immerses θ into the desirable space Θ_0 , and hence will be referred to as the *immersion map*. The induced posterior Π^* will be referred to as the *immersion posterior*. This provides an extension of the Bayesian paradigm since the identity map as the immersion map for the situation $\Theta_0 = \Theta$ reduces the immersion posterior to the classical Bayesian posterior.

The approach has been successfully used in several works including [14–17, 43] for shape-restricted problems, and by [6–9] for differential equation models. These authors used a projection map p obtained by minimizing a certain distance from the posterior sample to the restricted space, and the resulting induced random measure is called the projection posterior distribution. The projection map p satisfies the appealing property $p(\theta) = \theta$ for all $\theta \in \Theta_0$.

While a projection map with respect to an appropriate distance is a natural choice for an immersion map, the restriction to a projection map is unnecessary for the concept to be used. Depending on the aspect to be studied, there may not be a natural distance associated with it. This happens, for instance, if we are interested in studying the posterior distribution of the function value at a given point. It is also not necessary for the immersion map ι to satisfy $\iota(\theta) = \theta$ for all $\theta \in \Theta_0$. Neither Θ_0 needs to be a subset of Θ , nor the immersion map needs to be defined all over Θ . All that is needed is that an alternative parameter space Θ exists where the model distribution P_θ makes sense, a prior Π can be put on Θ such that the posterior distribution can be computed relatively easily, and the random distribution induced by a map ι from the support of the posterior distribution $\Pi(\cdot|X)$ to Θ_0 can be analyzed theoretically to establish some desirable properties. In most situations, the family of measures $\{P_\theta : \theta \in \Theta\}$ is dominated, so the support of the posterior distribution $\Pi(\cdot|X)$ is contained in the support of the prior distribution Π . The immersion map may be allowed to depend on the sample size like a prior distribution may be allowed to depend on the sample size. Even dependence of ι on the data X may be allowed. Although there is no uniqueness in the choice of the immersion map, the main purpose is to increase flexibility in the posterior measure to achieve a targeted asymptotic frequentist property, such as coverage of a credible region. A choice of an immersion map is therefore guided by a desirable frequentist property. Even if Θ_0 and Θ coincide, the flexibility of the immersion posterior may be helpful to satisfy a desirable convergence property of the immersion posterior that the classical Bayesian posterior may lack.

In many applications, the prior distribution may be actually a sequence of prior distributions specified through a sieve indexed by a discrete variable $J = J_n$ depending on the sample size n . Let Θ_J stand for the sieve (typically a finite-dimensional subset of Θ) and Π_J stand for the prior at that stage concentrated on Θ_J . Then the computation of the posterior in the unrestricted space reduces to a finite-dimensional computation, often also aided by posterior conjugacy. It is then typical that the immersion map ι on Θ_J has the range in $\Theta_0 \cap \Theta_J$ so that the computation of the immersion posterior involves finite-dimensional computations only. Most examples from the existing literature, as well as the method used in this paper, fall in this setting.

3. Notation, model, prior and posterior distribution.

3.1. Notation. We summarize the notation we shall use in this paper. The notation $\mathbb{R}$, $\mathbb{N}$ and $\mathbb{Z}$ will stand for the real line, the set of natural numbers and the set of all integers, respectively. The positive half-line with and without 0, and the set of nonnegative integers are respectively denoted by $\mathbb{R}_{\geq 0}$, $\mathbb{R}_{> 0}$ and $\mathbb{Z}_{\geq 0}$. Bold Latin or Greek letters will be used to indicate column vectors and the nonbold letter with a subscript will denote a coordinate of the corresponding vector. For example, a_i is the i th coordinate of $\mathbf{a} \in \mathbb{R}^d$. Let $\mathbf{1}$ denote the d -dimensional all-one column vector and $\mathbf{0}$ the all-zero column vector. Let $\mathbf{A}^T$ denote the transpose of a matrix or a vector $\mathbf{A}$. For an arbitrary set A , the indicator function will be denoted by $\mathbb{1}_A(\cdot)$, and $\#A$ will denote the cardinality of a finite set A . Let $\lceil a \rceil$ stand for the smallest integer greater than or equal to a real number a . The symbol $\lesssim$ will stand for an inequality up to an unimportant constant multiple. For two positive real sequences a_n and b_n , we also use $a_n \ll b_n$, or equivalently, $b_n \gg a_n$ if $a_n = o(b_n)$. For $a, b \in \mathbb{R}$, let $a \wedge b = \min\{a, b\}$ and $a \vee b = \max\{a, b\}$. For $\mathbf{a}, \mathbf{b} \in \mathbb{R}^d$, let $\mathbf{a} \wedge \mathbf{b} = (a_1 \wedge b_1, \dots, a_d \wedge b_d)^T$, $\mathbf{a} \vee \mathbf{b} = (a_1 \vee b_1, \dots, a_d \vee b_d)^T$ and the pointwise product $\mathbf{a} \circ \mathbf{b} = (a_1 b_1, \dots, a_d b_d)^T$. For a vector $\mathbf{a} \in \mathbb{R}^d$, the Euclidean and the maximum norms are respectively denoted by $\|\mathbf{a}\|$ and $\|\mathbf{a}\|_\infty = \max\{|a_k| : 1 \leq k \leq d\}$. Let $[\mathbf{j}_1 : \mathbf{j}_2] = \{\mathbf{j} \in \mathbb{Z}^d : j_{1,k} \leq j_k \leq j_{2,k}, \text{ for all } 1 \leq k \leq d\}$ stand for the lattice with boundaries $\mathbf{j}_1, \mathbf{j}_2 \in \mathbb{Z}^d$.

For a multivariate function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, let $\partial_k^l f(\mathbf{x}) = \partial^l f(\mathbf{x}) / \partial x_k^l$ for $k \in \{1, \dots, d\}$ and $l \in \mathbb{Z}_{\geq 0}$ at a suitable point $\mathbf{x} \in \mathbb{R}^d$. For a multiple index $\mathbf{l} = (l_1, \dots, l_d)^T \in \mathbb{Z}_{\geq 0}^d$, we use $\partial^{\mathbf{l}} = \partial_1^{l_1} \cdots \partial_d^{l_d}$, $\mathbf{l}! = l_1! \cdots l_d!$ and $\mathbf{x}^{\mathbf{l}} = x_1^{l_1} \cdots x_d^{l_d}$. We adopt the coordinatewise partial ordering on $\mathbb{R}^d$, that is, for $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$, $\mathbf{x} \leq \mathbf{y}$ if and only if $x_k \leq y_k$ for all $1 \leq k \leq d$. We say that a function f on $\mathbb{R}^d$ is multivariate monotone if $f(\mathbf{x}) \leq f(\mathbf{y})$ for all $\mathbf{x} \leq \mathbf{y}$. The class of all multivariate monotone functions on $[0, 1]^d$ will be denoted by $\mathcal{M}$. Let $\mathbb{L}_p[\mathbf{a}, \mathbf{b}]$, $1 \leq p \leq \infty$, stand for the Lebesgue $\mathbb{L}_p$ -space on a multivariate interval $[\mathbf{a}, \mathbf{b}]$. Convergence in probability under a measure $\mathbf{P}$ is denoted by $\rightarrow_p$. Distributional equality will be denoted by $=_d$ and weak convergence by $\rightsquigarrow$.

3.2. Model. We observe n independent and identically distributed random samples $\mathbb{D}_n = ((X_1, Y_1), \dots, (X_n, Y_n))$ from the nonparametric multiple regression model,

$$(3.1) \quad Y = f(\mathbf{X}) + \varepsilon,$$

where Y is the response variable, $\mathbf{X}$ is a d -dimensional predictor and ε is a random error with mean 0 and finite variance σ^2 , independent of $\mathbf{X}$. Instead of assuming any global smoothness condition on f , we assume that f is a multivariate monotone function. To construct the likelihood function, we assume that ε is normally distributed, but the actual data-generating process need not be so.

The first assumption is about the local regularity of the true regression function f_0 near a point of interest $\mathbf{x}_0$. This assumption, as in [38], is an essential ingredient to establish the limiting distribution.

ASSUMPTION 1. Let $f_0 \in \mathcal{M}$. For $\mathbf{x}_0 \in (0, 1)^d$ and $1 \leq k \leq d$, let β_k be the order of the first nonzero derivative of f at $\mathbf{x}_0$ along the k th coordinate, that is, $\beta_k = \min_{l \geq 1} \{l : \partial_k^l f_0(\mathbf{x}_0) \neq 0\}$ and $\beta_k = \infty$ if $\partial_k^l f_0(\mathbf{x}_0) = 0$ for all $l \geq 1$. Without loss of generality, we may assume that f_0 depends on its first s arguments locally at $\mathbf{x}_0$, that is, $1 \leq \beta_1, \dots, \beta_s < \infty$, and that $\beta_{s+1} = \dots = \beta_d = \infty$ for some $0 \leq s \leq d$. Define an index set $L = \{\mathbf{l} : 0 < \sum_{k=1}^s l_k / \beta_k \leq 1 \text{ and } l_k = 0, \text{ for } k = s+1, \dots, d\}$. For a positive sequence

$\omega_n \downarrow 0$, set $\mathbf{r}_n = (\omega_n^{1/\beta_1}, \dots, \omega_n^{1/\beta_s}, 1, \dots, 1)^T$. For any $t > 0$,

$$(3.2) \quad \lim_{\omega_n \downarrow 0} \omega_n^{-1} \sup_{\substack{\mathbf{x} \in [0, 1]^d, \\ |x_k - x_{0,k}| \leq t r_{n,k}, \\ 1 \leq k \leq d}} \left| f_0(\mathbf{x}) - f_0(\mathbf{x}_0) - \sum_{l \in L} \frac{\partial^l f_0(\mathbf{x}_0)}{l!} (\mathbf{x} - \mathbf{x}_0)^l \right| = 0.$$

Assumption 1 takes into account varying convergence rates across different coordinates, according to their respective smoothness levels. Each term in the expansion contributes toward approximation rates larger than or equal to ω_n . Let

$$(3.3) \quad L_0 = \left\{ l : 0 < \sum_{k=1}^s l_k / \beta_k < 1 \text{ and } l_k = 0 \text{ for } k = s+1, \dots, d \right\},$$

$$(3.4) \quad L^* = \left\{ l : \sum_{k=1}^s l_k / \beta_k = 1 \text{ and } l_k = 0 \text{ for } k = s+1, \dots, d \right\}.$$

Under Assumption 1, a unique feature for functions in $\mathcal{M}$ is that the derivatives of order $l \in L_0$ are zero (see Lemma 1 of [38]). Only those derivatives corresponding to the index set L^* can be nonzero. Thus, the nonzero terms in the expansion of (3.2) contribute the same approximation rate ω_n . However, Assumption 1 cannot exclude the nonzero mixed derivatives. Additional assumptions will be needed when we want to eliminate the mixed derivative terms.

Next, we make the following assumption on the distributions of the covariate $\mathbf{X}$ and the error ε from the data generating process (3.1).

ASSUMPTION 2. The covariate $\mathbf{X}$ has a density g such that $a_1 \leq g(\mathbf{x}) \leq a_2$ for all $\mathbf{x} \in [0, 1]^d$ and $0 < a_1 \leq a_2 < \infty$. Suppose g is continuous in a neighborhood of the set $\{(x_{0,1}, \dots, x_{0,s}, x_{s+1}, \dots, x_d) : x_k \in [0, 1], \text{ for } s+1 \leq k \leq d\}$. The random error ε , with mean 0 and variance σ_0^2 , has a finite $2(\sum_{k=1}^s \beta_k^{-1} + 1)$ th moment.

3.3. Prior. We put a prior distribution on f through a sieve of piecewise constant functions with gradually refining intervals of constancy, forming a partition of $[0, 1]^d$. For $\mathbf{J} \in \mathbb{Z}_{>0}^d$, let $I_{\mathbf{j}} = \prod_{k=1}^d ((j_k - 1)/J_k, j_k/J_k]$ be a hyperrectangle in $[0, 1]^d$, indexed by a d -dimensional vector $\mathbf{j}$, for $\mathbf{j} \in [\mathbf{1} : \mathbf{J}] \setminus \{\mathbf{1}\}$ and $I_{\mathbf{1}} = \prod_{k=1}^d [0, 1/J_k]$. Then $\{I_{\mathbf{j}}\}_{\mathbf{j} \in [\mathbf{1} : \mathbf{J}]}$ forms a partition of $[0, 1]^d$. We define a class of piecewise constant functions $\mathcal{K}_{\mathbf{J}} := \{f = \sum_{\mathbf{j} \in [\mathbf{1} : \mathbf{J}]} \theta_{\mathbf{j}} \mathbb{1}_{I_{\mathbf{j}}} : \theta_{\mathbf{j}} \in \mathbb{R}\}$. As we follow the immersion posterior approach, we do not initially impose the order restriction. A prior is imposed on $f = \sum_{\mathbf{j} \in [\mathbf{1} : \mathbf{J}]} \theta_{\mathbf{j}} \mathbb{1}_{I_{\mathbf{j}}}$ in $\mathcal{K}_{\mathbf{J}}$ by giving independent Gaussian priors to $\theta_{\mathbf{j}}$, namely,

$$(3.5) \quad \theta_{\mathbf{j}} \sim N(\zeta_{\mathbf{j}}, \sigma^2 \lambda_{\mathbf{j}}^2) \quad \text{independently for all } \mathbf{j} \in [\mathbf{1} : \mathbf{J}],$$

where $\max_{\mathbf{j}} |\zeta_{\mathbf{j}}| < \infty$ and $\min_{\mathbf{j}} \lambda_{\mathbf{j}}^2 \geq b > 0$.

The values of the prior parameters, $\zeta_{\mathbf{j}}$ and $\lambda_{\mathbf{j}}$, will not affect our asymptotic results. However, in practice, when very little prior information is available, it is sensible to choose $\zeta_{\mathbf{j}} = 0$ and $\lambda_{\mathbf{j}}$ large for all $\mathbf{j}$.

3.4. Posterior distribution. We use the Gaussian distribution

$$(3.6) \quad Y_i \sim N\left(\sum_{\mathbf{j} \in [\mathbf{1} : \mathbf{J}]} \theta_{\mathbf{j}} \mathbb{1}\{\mathbf{X}_i \in I_{\mathbf{j}}\}, \sigma^2\right),$$

which leads to, in the unrestricted parameter space, a Gaussian joint likelihood for $(\theta_j : j \in [1 : J])$ without any cross-product terms in the exponent. This gives independent Gaussian posterior distribution for each θ_j , given σ , such that by conjugacy,

$$(3.7) \quad \theta_j | \mathbb{D}_n, \sigma \sim N((N_j \bar{Y}|_{I_j} + \zeta_j \lambda_j^{-2}) / (N_j + \lambda_j^{-2}), \sigma^2 / (N_j + \lambda_j^{-2})),$$

where $N_j = \#\{i : X_i \in I_j\}$ and $\bar{Y}|_{I_j} = \sum_{i=1}^n Y_i \mathbb{1}\{X_i \in I_j\} / N_j$.

The parameter σ^2 can be estimated by maximizing the marginal likelihood function given by

$$(2\pi\sigma^2)^{-n/2} \prod_{j \in [1:J]} (1 + \lambda_j^2 N_j)^{-1/2} \\ \times \exp \left[-\frac{1}{2\sigma^2} \left\{ \sum_{i=1}^n \left(Y_i - \sum_{j: X_i \in I_j} \zeta_j \right)^2 - \sum_{j \in [1:J]} \frac{N_j^2 (\bar{Y}|_{I_j} - \zeta_j)^2}{N_j + \lambda_j^{-2}} \right\} \right],$$

and the resulting estimator

$$(3.8) \quad \hat{\sigma}_n^2 = \frac{1}{n} \left[\sum_{i=1}^n \left(Y_i - \sum_{j: X_i \in I_j} \zeta_j \right)^2 - \sum_{j \in [1:J]} \frac{N_j^2 (\bar{Y}|_{I_j} - \zeta_j)^2}{N_j + \lambda_j^{-2}} \right],$$

may be plugged into the expression (3.7). Alternatively, in a fully Bayesian framework, we can give σ^2 an inverse-Gamma prior $IG(b_1, b_2)$ with parameters $b_1 > 0$, $b_2 > 0$, and obtain that the posterior distribution of σ^2 is given by $IG(b_1 + n/2, b_2 + n\hat{\sigma}_n^2/2)$. It will be shown in Lemma B.5 of [54] that the marginal maximum likelihood estimator of σ^2 as well as the posterior for σ^2 concentrate in a shrinking neighborhood of its true value σ_0^2 . Then it easily follows that the asymptotic behavior of the posterior distribution of f is identical with that when σ is known to be σ_0 . Hence, it suffices to study the asymptotic behavior of the posterior distribution given σ .

The unrestricted posterior distribution of f given σ is induced from (3.7) by the representation $f = \sum_{j \in [1:J]} \theta_j \mathbb{1}_{I_j}$. To obtain the immersion posterior distribution to make an inference, we consider three possible immersion maps.

Define

$$(3.9) \quad \mathcal{M}_J = \left\{ f = \sum_{j \in [1:J]} \theta_j \mathbb{1}_{I_j} : \theta_j \in \mathbb{R} \text{ and } \theta_{j_1} \leq \theta_{j_2} \text{ if } j_1 \leq j_2 \right\},$$

consisting of the coordinatewise nondecreasing functions taking constant value on every I_j .

Based on the isotonization procedure introduced in [28], consider transformations $\underline{\iota}$ and $\bar{\iota}$ acting on $f = \sum_{j \in [1:J]} \theta_j \mathbb{1}_{I_j} \in \mathcal{K}_J$ mapping to an element of $\mathcal{M}_J$ defined by

$$(3.10) \quad \underline{\iota}(f)(\mathbf{x}) = \max_{j_1 \leq j_0(\mathbf{x})} \min_{\substack{j_0(\mathbf{x}) \leq j_2 \\ N_{[j_1:j_2]} > 0}} \frac{\sum_{j \in [j_1:j_2]} N_j \theta_j}{N_{[j_1:j_2]}},$$

$$(3.11) \quad \bar{\iota}(f)(\mathbf{x}) = \min_{j_0(\mathbf{x}) \leq j_2} \max_{\substack{j_1 \leq j_0(\mathbf{x}) \\ N_{[j_1:j_2]} > 0}} \frac{\sum_{j \in [j_1:j_2]} N_j \theta_j}{N_{[j_1:j_2]}},$$

where $j_0(\mathbf{x}) = \lceil \mathbf{x} \circ \mathbf{J} \rceil$, $N_{[j_1:j_2]} = \sum_{j \in [j_1:j_2]} N_j$, and $\mathbf{x} \in [0, 1]^d$, for j_1, j_2 in $\mathbb{Z}^d$. The immersion posterior can be derived through the immersion map, ι , which is chosen to be either $\underline{\iota}$ or $\bar{\iota}$. This is determined by examining the resulting induced distribution of

$$(3.12) \quad f_* = \underline{\iota}(f),$$

$$(3.13) \quad f^* = \bar{\iota}(f).$$

It is obvious that $\iota(f) \in \mathcal{M}_J$ and $\iota(f) = f$ if $f \in \mathcal{M}_J$ and $N_j > 0$ for all $j \in [1 : J]$. Generally, $\underline{\iota}(f)(\mathbf{x}) \leq \bar{\iota}(f)(\mathbf{x})$ for any $\mathbf{x} \in [0, 1]^d$, but this may fail to hold if $N_j = 0$ for some j , see [21]. To neutralize the effect stemming from the order of minimization and maximization, we propose using the average of $\underline{\iota}$ and $\bar{\iota}$, leading to another immersion map $\iota = (\underline{\iota} + \bar{\iota})/2$, by which f is mapped to

$$(3.14) \quad \tilde{f} = (f_* + f^*)/2.$$

The projection map for the univariate case is typically computed by the pool adjacent violator algorithm (see Section 2.3 of [4]), which requires $O(J)$ computations for a function with J steps. The computation of f_* or f^* requires no more than $(\prod_{k=1}^d J_k)^3$ operations by the brute-force search utilizing the block max-min or min-max formulas.

3.5. Effect of the immersion map. To see the effect of the immersion map on the posterior distribution of the function value at a point $\mathbf{x}_0 = (0.5, 0.5) \in [0, 1]^2$, we conduct a small simulation study and compare the unrestricted and immersion posterior density for a randomly generated sample of three different sizes $n = 100, 200$ and 500 , and three different regression functions: (i) $f_0(x_1, x_2) = x_1 + x_2$; (ii) $f_0(x_1, x_2) = \sqrt{x_1 + x_2}$; (iii) $f_0(x_1, x_2) = \mathbb{1}\{x_1 < 1/3\} + 2\mathbb{1}\{1/3 \leq x_1 < 2/3\} + 3\mathbb{1}\{x_1 \geq 2/3\}$. Let X_1 and X_2 be distributed independently and uniformly on $[0, 1]$ and error $\varepsilon \sim N(0, \sigma^2)$ with true value of σ to be 0.1 . We choose the number of grid points $J_1 = J_2 = J = \lceil n^{1/4} \log_{10} n \rceil$. The random heights, $\{\theta_{(j_1, j_2)} : j_1, j_2 \leq J\}$, are endowed with the independent Gaussian prior $N(0, 1000\sigma^2)$. The variance σ^2 is estimated using the maximum marginal likelihood method. We plot both the unrestricted posterior density and the estimated immersion posterior density in the same figure. The latter is based on 2000 posterior samples transformed by the immersion map $(\bar{\iota} + \underline{\iota})/2$.

As evident from Figure 1, the immersion posterior density functions exhibit lower variance across all instances, albeit to varying degrees depending on the true regression functions and sample sizes. Furthermore, the modes of the immersion posterior are nearer to the true value. The impacts of the other immersion maps, $\bar{\iota}$ and $\underline{\iota}$, on the posterior were found to be similar.

4. Coverage of credible intervals. Let $\mathbf{x}_0 \in (0, 1)^d$ be fixed. Suppose that we want to make an inference on $f(\mathbf{x}_0)$. For a given $0 < \gamma < 1$, consider a $(1 - \gamma)$ -credible interval with endpoints the $\gamma/2$ and $(1 - \gamma)/2$ quantiles of $f_*(\mathbf{x}_0)$, $f^*(\mathbf{x}_0)$, or $\tilde{f}(\mathbf{x}_0)$ defined in (3.12)–(3.14). To obtain the limiting frequentist coverage of these credible intervals, we obtain the weak limit of the immersion posterior distributions of f for all three immersion maps $\underline{\iota}$, $\bar{\iota}$ and $(\underline{\iota} + \bar{\iota})/2$ at $\mathbf{x} = \mathbf{x}_0$.

Let H_1 and H_2 be two independent centered Gaussian processes indexed by $(\mathbf{u}, \mathbf{v}) \in \mathbb{R}_{\geq 0}^d \times \mathbb{R}_{\geq 0}^d$ with the covariance kernel

$$(4.1) \quad \prod_{k=1}^s (u_k \wedge u'_k + v_k \wedge v'_k) D_s(\mathbf{u} \wedge \mathbf{u}', \mathbf{v} \wedge \mathbf{v}'),$$

where $D_d(\mathbf{u}, \mathbf{v}) = g(\mathbf{x}_0)$, where g is the probability density function of X , and for $s = 0, \dots, d-1$, and $D_s(\mathbf{u}, \mathbf{v})$ is given by

$$(4.2) \quad \int_{\substack{x_k \in [(x_0 - \mathbf{u})_k, (x_0 + \mathbf{v})_k] \cap [0, 1] \\ s+1 \leq k \leq d}} g(x_{0,1}, \dots, x_{0,s}, x_{s+1}, \dots, x_d) dx_{s+1} \cdots dx_d.$$

Additionally, we define a Gaussian process

$$(4.3) \quad \begin{aligned} U(\mathbf{u}, \mathbf{v}) = & \frac{\sigma_0 H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \frac{\sigma_0 H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \\ & + \sum_{l \in L^*} \frac{\partial^l f_0(\mathbf{x}_0)}{(l+1)!} \prod_{k=1}^s \frac{v_k^{l_k+1} - (-u_k)^{l_k+1}}{u_k + v_k} \end{aligned}$$

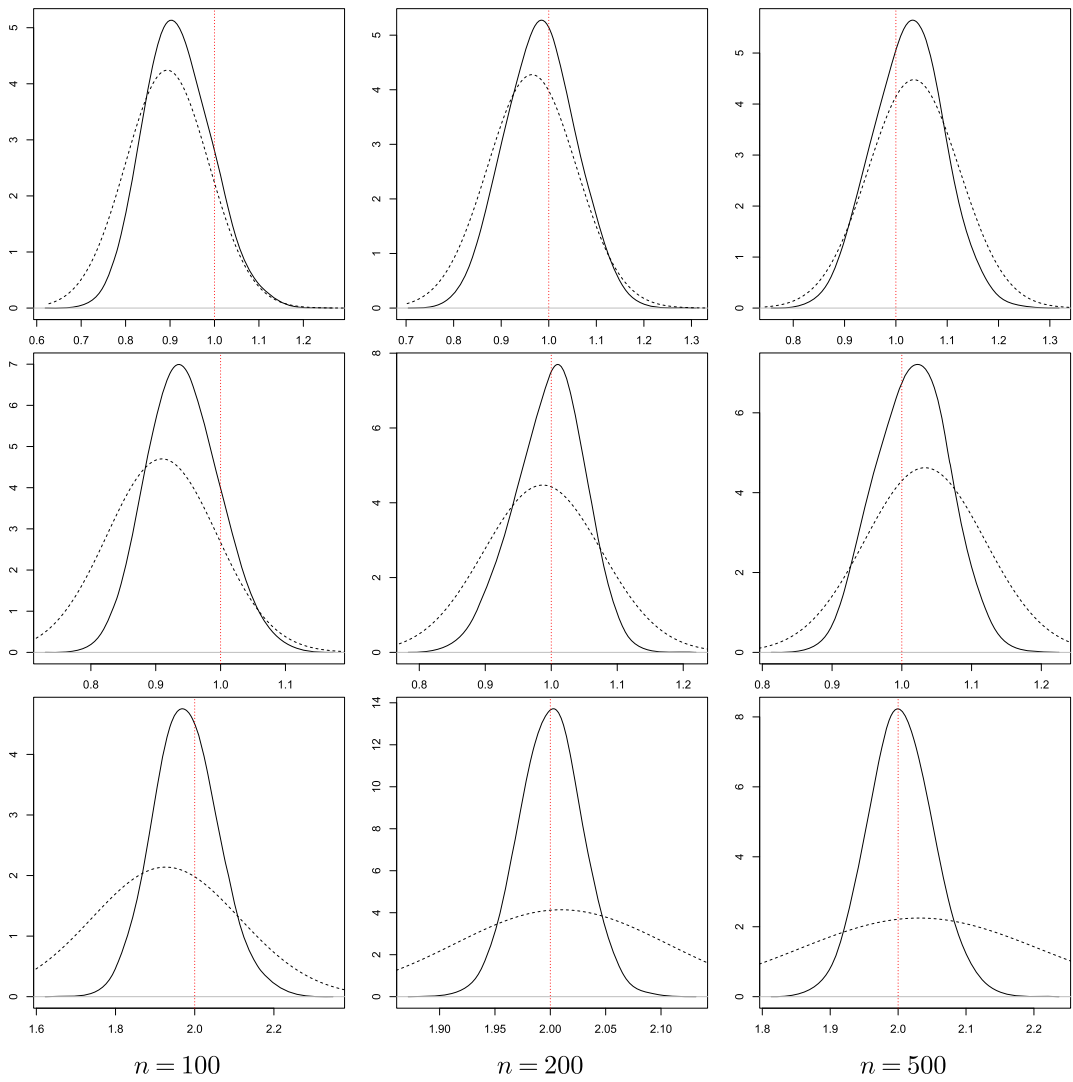


FIG. 1. Unrestricted and immersion posterior density functions of $f(\mathbf{x}_0)$. The solid black line stands for the immersion posterior density, while the black dashed line represents the unrestricted posterior density. The true function value is indicated by the red dotted vertical line. The rows correspond to functions (i), (ii) and (iii), respectively, while the columns represent sample sizes of $n = 100, 200$ and 500 .

indexed by $(\mathbf{u}, \mathbf{v}) \in \mathbb{R}_{\geq 0}^d \times \mathbb{R}_{\geq 0}^d$, and its functionals

$$(4.4) \quad Z_* = \sup_{\substack{\mathbf{u} \geq \mathbf{0} \\ u_k \leq x_{0,k} \\ s+1 \leq k \leq d}} \inf_{\substack{\mathbf{v} \geq \mathbf{0} \\ v_k \leq 1-x_{0,k} \\ s+1 \leq k \leq d}} U(\mathbf{u}, \mathbf{v}), \quad Z^* = \inf_{\substack{\mathbf{v} \geq \mathbf{0} \\ v_k \leq 1-x_{0,k} \\ s+1 \leq k \leq d}} \sup_{\substack{\mathbf{u} \geq \mathbf{0} \\ u_k \leq x_{0,k} \\ s+1 \leq k \leq d}} U(\mathbf{u}, \mathbf{v}).$$

The following result describes the asymptotic behavior of the normalized immersion posterior distributions of $f(\mathbf{x}_0)$. Recall that $\mathbb{D}_n$ represents the data and $r_{n,k}$ in Assumption 1 is the convergence rate along the k th direction through adjusting the overall rate ω_n according to the local smoothness levels. The weak limit of the normalized immersion posterior distribution plays a central role in the study of the limiting coverage of the credible intervals based on the immersion posterior quantiles.

THEOREM 4.1. *Let $\omega_n = n^{-1/(2+\sum_{k=1}^s \beta_k^{-1})}$ and let $\mathbf{r}_n = (\omega_n^{1/\beta_1}, \dots, \omega_n^{1/\beta_s}, 1, \dots, 1)^T$. Suppose that $\mathbf{J}$ satisfies $J_k \gg r_{n,k}^{-1}$ for each $k = 1, \dots, d$, and $\prod_{k=1}^d J_k \ll n\omega_n$. Under As-*

sumptions 1 and 2, for any $z \in \mathbb{R}$, we have

$$(4.5) \quad \Pi(\omega_n^{-1}(f_*(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq z | \mathbb{D}_n) \rightsquigarrow \mathbb{P}(Z_* \leq z | H_1);$$

$$(4.6) \quad \Pi(\omega_n^{-1}(f^*(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq z | \mathbb{D}_n) \rightsquigarrow \mathbb{P}(Z^* \leq z | H_1);$$

$$(4.7) \quad \Pi(\omega_n^{-1}(\tilde{f}(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq z | \mathbb{D}_n) \rightsquigarrow \mathbb{P}((Z_* + Z^*)/2 \leq z | H_1).$$

Furthermore, for any $(z_1, z_2) \in \mathbb{R}^2$,

$$(4.8) \quad \begin{aligned} &\Pi(\omega_n^{-1}(f_*(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq z_1, \omega_n^{-1}(f^*(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq z_2 | \mathbb{D}_n) \\ &\rightsquigarrow \mathbb{P}(Z_* \leq z_1, Z^* \leq z_2 | H_1). \end{aligned}$$

REMARK 1. We make some remarks on Theorem 4.1.

1. The weak limit is understood in the usual sense for random variables since we consider the limiting behavior of the random probability measure of a fixed set $(-\infty, z]$. We refer to the proof technique of [38], which provides the distributional theory for the block-estimator in general multivariate isotonic regression, especially the small and large deviation arguments therein.

2. For the choice of J_k , the lower bound $r_{n,k}^{-1}$ is essential for Theorem 4.1. That eliminates the effect of the roughness of piecewise constant function approximation in view of the local contraction rate. But the upper bound, $n\omega_n$ in Theorem 4.1, is not fundamentally necessary for the validity of the weak limit. Instead, we can set the hyperparameters λ_j large enough, specifically, $\min \lambda_j^2 \gg \omega_n^{-1} \sqrt{n}$, to obtain the limiting theory. The rest of the proof of Theorem 4.1 will not be affected much without such an upper bound except for the treatment of σ^2 . The estimation of σ^2 is not a hard problem and any consistent procedure will work. For any $\beta_k \geq 1$ and any $0 \leq s \leq d$, we observe that $r_k \leq n^{-1/3}$ for all $1 \leq k \leq d$. Without the local smoothness information, we can choose $J_k \gg n^{1/3}$. On the other side, if we admit that $\beta_k = 1$, $1 \leq k \leq d$, is the leading case for the multivariate regression function, we can then choose $J_k \gg n^{1/(2+d)}$. One may raise concerns that selecting J in this manner seems not optimal when the true regression function is less smooth. However, this concern typically does not pose a big issue in practice. When lacking prior information about the regression function, an appropriately large J_k can be chosen and an uninformative prior should be applied to θ_j , such as a normal prior with a large variance. Our empirical study indicates that opting for a larger J can indeed enhance the performance of our method. For practical applications, we recommend selecting J_k no less than 15, particularly when dealing with a smaller sample size. Notably, $\mathbf{J}$ is not a tuning parameter in this context; the immersion posterior is governed by shape restrictions rather than a tuning process like bandwidth selection in kernel smoothing. The choice of $\mathbf{J}$ does not influence the contraction rate or the distributional theory, distinguishing it from typical tuning parameters.

3. It is also important to note that we employ a working normal model to derive the posterior distribution. The validity of this method remains intact even when the model is misspecified. The finite-moment condition for the random error ε can be relaxed to the second order, as in [38], by selecting a sufficiently large λ_j^2 as in the last point.

The covariance kernels of the processes H_1 and H_2 depend on g , and the limiting Gaussian process also involves the derivative values of f_0 at $\mathbf{x}_0$. A considerable simplification happens in some special cases where the parameters appear through a scale parameter in the kernel. It will be seen shortly that this fact has a far-reaching implication in that the limiting coverage of a credible interval constructed from the immersion posterior is free of the unknown parameters of the model. If L^* defined by (3.4) only contains $\beta_k \mathbf{e}_k$ for $k = 1, \dots, s$, where

$\mathbf{e}_k$ denotes the standard unit vector in $\mathbb{R}^d$ with one in the k th component and zero elsewhere, then the limiting processes in Theorem 4.1 can be further simplified by self-similarity. A factor depending on f_0 comes out as a multiplicative constant, and the remaining factor is only a known functional of H_1 and H_2 . The case $s = d$ stands for the regular case that all directional derivatives of f_0 at $\mathbf{x}_0$ are positive at a certain order. Then the covariance kernel further simplifies as a completely known function and a factor involving derivatives of the regression function and predictor density g . The result is precisely formulated in the result below.

PROPOSITION 4.1. *If $L^* = \{\beta_k \mathbf{e}_k : 1 \leq k \leq s\}$, then*

$$\begin{aligned} & \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\sigma_0 H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \frac{\sigma_0 H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ & \quad \left. + \sum_{k=1}^s \left[\frac{\partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \cdot \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right] \right\} \\ & =_d A_\beta \cdot \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \frac{H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ & \quad \left. + \sum_{k=1}^s \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right\}, \end{aligned}$$

where $A_\beta = (\sigma_0^2 \prod_{k=1}^s (\frac{\partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k+1)!})^{1/\beta_k})^{1/(2+\sum_{k=1}^s \beta_k^{-1})}$.

Furthermore, if $s = d$, then the above expression further simplifies to

$$\tilde{A}_\beta \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\tilde{H}_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \frac{\tilde{H}_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \sum_{k=1}^d \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right\},$$

where $\tilde{A}_\beta = (\frac{\sigma_0^2}{g(\mathbf{x}_0)} \prod_{k=1}^d (\frac{\partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k+1)!})^{1/\beta_k})^{1/(2+\sum_{k=1}^d \beta_k^{-1})}$, and $\tilde{H}_1$ and $\tilde{H}_2$ are two independent centered Gaussian processes with covariance kernel given by $\prod_{k=1}^d (u_k \wedge u'_k + v_k \wedge v'_k)$, $(\mathbf{u}, \mathbf{v}), (\mathbf{u}', \mathbf{v}') \in \mathbb{R}_{\geq 0}^d \times \mathbb{R}_{\geq 0}^d$.

The same conclusion also applies to the inf sup-functional obtained by switching the positions of the supremum and the infimum.

REMARK 2 (Univariate case). We specialize to the univariate case $s = d = 1$, with a general β , expanding from the case $\beta = 1$ studied by [14]. Then

$$(4.9) \quad \tilde{H}_i(u, v) =_d W_i(v) + W_i(-u) =_d W_i(v) - W_i(-u), \quad (u, v) \in \mathbb{R}_{\geq 0}^2,$$

where W_1, W_2 are two independent standard two-sided Brownian motions starting from 0. Observe that the sup-inf functional

$$\begin{aligned} & \sup_{u>0} \inf_{v>0} \left\{ \frac{\tilde{H}_1(u, v)}{u + v} + \frac{\tilde{H}_2(u, v)}{u + v} + \frac{v^{\beta+1} - (-u)^{\beta+1}}{u + v} \right\} \\ & =_d \sup_{u>0} \inf_{v>0} \left\{ \frac{(W_1(v) + W_2(v) + v^{\beta+1}) - (W_1(-u) + W_2(-u) + u^{\beta+1})}{v - (-u)} \right\}, \end{aligned}$$

coincides with the slope of the greatest convex minorant of the process $W_1(t) + W_2(t) + t^{\beta+1}$. By the switching relation (cf. [33], page 56), for any $z \in \mathbb{R}$,

$$\begin{aligned} & \mathbb{P} \left(\tilde{A}_\beta \sup_{u>0} \inf_{v>0} \left\{ \frac{\tilde{H}_1(u, v)}{u + v} + \frac{\tilde{H}_2(u, v)}{u + v} + \frac{v^{\beta+1} - (-u)^{\beta+1}}{u + v} \right\} \leq z \right) \\ & = \mathbb{P}(\arg \min \{W_1(t) + W_2(t) + t^{\beta+1} - \tilde{A}_\beta^{-1} z t : t \in \mathbb{R}\} \geq 0). \end{aligned}$$

If $\beta = 1$, the last display can be further simplified by applying the change of variable, $t = s + z/(2\tilde{A}_1)$, and is equal to

$$\mathbb{P}(2\tilde{A}_1 \arg \min\{W_1(s) + W_2(s) + s^2 : s \in \mathbb{R}\} \leq z),$$

with $\tilde{A}_1 = (\sigma_0^2 f'(x_0)/(2g(x_0)))^{1/3}$. This reproduces the main result of [14].

Now we are ready for the evaluation of the limiting coverage of an immersion posterior credible interval for $f(\mathbf{x}_0)$. Let

$$(4.10) \quad Q_{n,\gamma}^{(1)} = \inf\{z : \Pi(f_*(\mathbf{x}_0) \leq z | \mathbb{D}_n) \geq 1 - \gamma\}$$

stand for the $(1 - \gamma)$ -quantile of $f_*(\mathbf{x}_0)$. Similarly, let $Q_{n,\gamma}^{(2)}$ and $Q_{n,\gamma}^{(3)}$ stand for that of $f^*(\mathbf{x}_0)$ and $\tilde{f}(\mathbf{x}_0)$, respectively. Let $\tilde{U}(\mathbf{u}, \mathbf{v})$ stand for the Gaussian process

$$(4.11) \quad \frac{\tilde{H}_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \frac{\tilde{H}_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \sum_{k=1}^d \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k}$$

indexed by $(\mathbf{u}, \mathbf{v}) \in \mathbb{R}_{\geq 0}^d \times \mathbb{R}_{\geq 0}^d$.

The following result gives the ultimate conclusion of the paper about asymptotic coverage of credible intervals for the regression function value at an interior point.

THEOREM 4.2. *Under the assumed setup, Assumptions 1 and 2, and the condition that $L^* = \{\beta_k \mathbf{e}_k : 1 \leq k \leq s\}$, the asymptotic coverage of the quantile-based one-sided credible interval $(-\infty, Q_{n,\gamma}^{(1)}]$ is given by*

$$\begin{aligned} & \mathbb{P}\left(\mathbb{P}\left(\sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \frac{H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \right. \right. \\ & \quad \left. \left. \left. + \sum_{k=1}^s \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right\} \leq 0 \middle| H_1 \right) \leq 1 - \gamma\right). \end{aligned}$$

If $Q_{n,\gamma}^{(1)}$ is replaced by $Q_{n,\gamma}^{(2)}$, the above limit is changed by swapping the order of the supremum and infimum operations. If $Q_{n,\gamma}^{(1)}$ is replaced by $Q_{n,\gamma}^{(3)}$, the above limit is changed by replacing the expression on the right-hand side with the average of the sup inf and inf sup operations.

Moreover, if $s = d$:

- (i) $\mathbb{P}_0(f_0(\mathbf{x}_0) \leq Q_{n,\gamma}^{(1)}) \rightarrow \mathbb{P}(Z_B^{(1)} \leq 1 - \gamma)$;
- (ii) $\mathbb{P}_0(f_0(\mathbf{x}_0) \leq Q_{n,\gamma}^{(2)}) \rightarrow \mathbb{P}(Z_B^{(2)} \leq 1 - \gamma)$;
- (iii) $\mathbb{P}_0(f_0(\mathbf{x}_0) \leq Q_{n,\gamma}^{(3)}) \rightarrow \mathbb{P}(Z_B^{(3)} \leq 1 - \gamma)$,

where $Z_B^{(1)} = \mathbb{P}(\sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \tilde{U}(\mathbf{u}, \mathbf{v}) \leq 0 | \tilde{H}_1)$, $Z_B^{(2)} = \mathbb{P}(\inf_{\mathbf{v} \geq \mathbf{0}} \sup_{\mathbf{u} \geq \mathbf{0}} \tilde{U}(\mathbf{u}, \mathbf{v}) \leq 0 | \tilde{H}_1)$ and $Z_B^{(3)} = \mathbb{P}(\frac{1}{2} \{\sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \tilde{U}(\mathbf{u}, \mathbf{v}) + \inf_{\mathbf{v} \geq \mathbf{0}} \sup_{\mathbf{u} \geq \mathbf{0}} \tilde{U}(\mathbf{u}, \mathbf{v})\} \leq 0 | \tilde{H}_1)$.

PROOF. We observe that $f_0(\mathbf{x}_0) \leq Q_{n,\gamma}^{(1)}$ if and only if

$$\Pi(f_*(\mathbf{x}_0) \leq f_0(\mathbf{x}_0) | \mathbb{D}_n) = \Pi(\omega_n^{-1}(f_*(\mathbf{x}_0) - f_0(\mathbf{x}_0)) \leq 0 | \mathbb{D}_n) \leq 1 - \gamma.$$

Hence, by Theorem 4.1 and Proposition 4.1, as the multiplicative positive constant in the limiting process can be dropped because the interval $(-\infty, 0]$ remains invariant under a scale-change, the first conclusion follows immediately. The special cases follow from the second part of Proposition 4.1. $\square$

REMARK 3. For $d = 1$, $Z_B^{(1)}$, $Z_B^{(2)}$ and $Z_B^{(3)}$ all coincide, and may be simply denoted by Z_B as in [14].

The distributions of $Z_B^{(1)}$ and $Z_B^{(2)}$ are related, as shown next.

PROPOSITION 4.2. For any $z \in [0, 1]$, we have $P(Z_B^{(1)} \leq z) = P(Z_B^{(2)} \geq 1 - z)$, and $Z_B^{(3)}$ is symmetrically distributed about $1/2$.

From Theorem 4.2 and Proposition 4.2, it follows that the limiting coverage of a one-sided Bayesian credible interval for $f(\mathbf{x}_0)$ using one of the three proposed immersion posteriors can be evaluated, is free of the true regression function (and also is free of the density g of the predictor if $s = d$, and hence depends only on the credibility level), but in general, need not be equal to the credibility. Nevertheless, a targeted limiting coverage can be obtained by starting with a certain credibility level that can be explicitly computed by back-calculation. As in the univariate monotone problems studied by [14, 15], numerical calculations show that the required credibility to obtain a specific limiting coverage is less than the targeted coverage, the opposite of the phenomenon [20] observed for smoothing problems. However, unlike in the univariate case where the limiting Bayes–Chernoff distribution determining the asymptotic coverage of the credible interval is symmetric, the corresponding random variables $Z_B^{(1)}$ and $Z_B^{(2)}$ for the posterior based on the immersion maps $\underline{\iota}$ and $\bar{\iota}$ appearing in the multivariate case are not symmetric. This has implications for the limiting coverage of a two-sided credible interval, which is more commonly used in practice. For instance, for $0 < \gamma < 1/2$, a two-sided $(1 - \gamma)$ -credible interval $[Q_{n,1-\gamma/2}, Q_{n,\gamma/2}]$ based on the immersion posterior using the map $\underline{\iota}$, the limiting coverage is given by $P(Z_B^{(1)} \leq 1 - \gamma/2) - P(Z_B^{(1)} \leq \gamma/2)$. The corresponding limit for the immersion posterior using the map $\bar{\iota}$ is $P(Z_B^{(2)} \leq 1 - \gamma/2) - P(Z_B^{(2)} \leq \gamma/2)$. Interestingly, a separate table for the distribution function of $Z_B^{(2)}$ is not needed, as it can be obtained from that of $Z_B^{(1)}$ in view of Proposition 4.2. The symmetry of $Z_B^{(3)}$, however, implies that the credibility level $1 - \gamma$ needed to make the asymptotic coverage of an equal-tailed $(1 - \gamma)$ -credible interval $1 - \alpha$ is obtained by choosing $1 - \gamma = 1 - 2F_{Z_B^{(3)}}^{-1}(\alpha/2)$, which is readily obtained once the cumulative distribution function $F_{Z_B^{(3)}}$ of $Z_B^{(3)}$ is tabulated.

5. Numerical results.

5.1. *Distribution of Z_B .* In this section, we present tables detailing the distribution and quantiles of Z_B for the case $d = 1$ when $\beta = 1, 3, 5$, as well as those for $Z_B^{(1)}$, $Z_B^{(2)}$, $Z_B^{(3)}$ for the case $d = 2$ when $\beta = (1, 1), (1, 3), (3, 3)$. The distributions of these variables are simulated using the Monte Carlo method, with the Gaussian processes concerned being generated by discrete approximation. The quantile table can function as a recalibration reference to achieve the exact frequentist asymptotic coverage.

5.1.1. *Case $d = 1$.* First, we generate approximations to the Gaussian processes $\tilde{H}_1$ and $\tilde{H}_2$. Let $\tilde{H}$ denote either $\tilde{H}_1$ or $\tilde{H}_2$. To approximate $\tilde{H}$, we generate $14m$ independent standard Gaussian random variables, specifically $\{\zeta_j : j = 1, \dots, 7m\}$ and $\{\zeta'_j : j = 1, \dots, 7m\}$, where $m = 50$. Then $\tilde{H}$ can be approximated as follows:

$$(5.1) \quad \tilde{H}(u, v) \approx \frac{1}{\sqrt{m}} \left[\sum_{j=1}^{\lceil mu \rceil} \zeta_j + \sum_{j=1}^{\lceil mv \rceil} \zeta'_j \right],$$

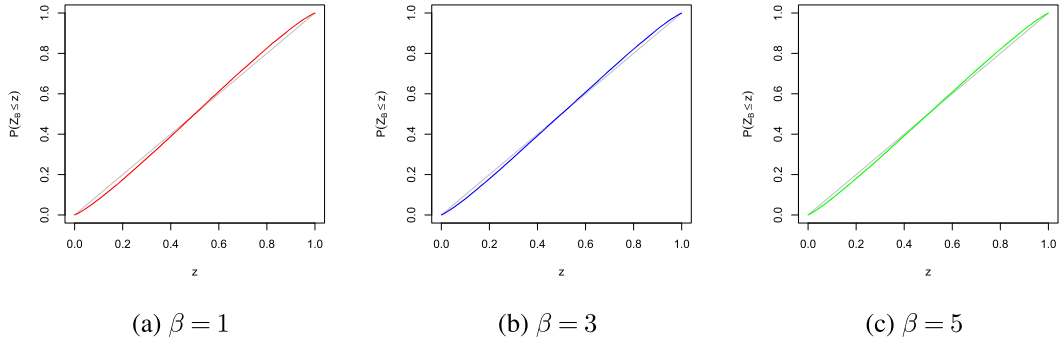


FIG. 2. Distribution functions of Z_B .

for $u, v \in [0, 7]$. Given each instance of $\tilde{H}_1$, we generate 500 realizations of $\tilde{H}_2$. For each realization, we compute the sup-inf functional. The proportion of nonpositive outcomes then serves as a sample value of Z_B . We repeat the generation process 50,000 times to obtain the approximate distribution function of Z_B .

In Figure 2, we draw the simulated distribution functions of Z_B with $\beta = 1, 3$ and 5. We give the values of $P(Z_B \leq z)$ with different smoothness levels for selected z values in Table 1 and the values of the quantiles of Z_B 's distribution in Table 2.

5.1.2. Case $d = 2$. To approximate $\tilde{H}(\mathbf{u}, \mathbf{v})$, for $\mathbf{u}, \mathbf{v} \in \mathbb{R}^2$, we generate 4 random matrices $\zeta^{(1)}, \zeta^{(2)}, \zeta^{(3)}$ and $\zeta^{(4)}$ with independent standard Gaussian random variables. The dimensions of these 4 matrices are $\lceil mt_1 \rceil \times \lceil mt_2 \rceil$, $\lceil ms_1 \rceil \times \lceil mt_2 \rceil$, $\lceil ms_1 \rceil \times \lceil ms_2 \rceil$ and $\lceil mt_1 \rceil \times \lceil ms_2 \rceil$, for $m = 5$ and $t_1 = t_2 = s_1 = s_2 = 5$. $\tilde{H}(\mathbf{u}, \mathbf{v})$ is then approximated by

$$\frac{1}{m} \left(\sum_{i=1}^{\lceil mv_1 \rceil} \sum_{j=1}^{\lceil mv_2 \rceil} \zeta_{ij}^{(1)} + \sum_{i=1}^{\lceil mu_1 \rceil} \sum_{j=1}^{\lceil mv_2 \rceil} \zeta_{ij}^{(2)} + \sum_{i=1}^{\lceil mu_1 \rceil} \sum_{j=1}^{\lceil mu_2 \rceil} \zeta_{ij}^{(3)} + \sum_{i=1}^{\lceil mv_1 \rceil} \sum_{j=1}^{\lceil mu_2 \rceil} \zeta_{ij}^{(4)} \right),$$

for $u_1, u_2, v_1, v_2 \in [0, 5]$.

To get a sample of any one of $Z_B^{(1)}, Z_B^{(2)}$ or $Z_B^{(3)}$, we first generate a sample of $\tilde{H}_1$. Given this sample, we generate 500 realizations of $\tilde{H}_2$. We then compute the three functionals that define $Z_B^{(1)}, Z_B^{(2)}$ and $Z_B^{(3)}$. The conditional probabilities are approximated by the frequency

TABLE 1
Values of $P(Z_B \leq z)$

z	0.700	0.750	0.800	0.850	0.900	0.950	0.975	0.990	0.995
$\beta = 1$	0.719	0.772	0.826	0.875	0.923	0.965	0.985	0.994	0.997
$\beta = 3$	0.715	0.768	0.821	0.870	0.921	0.963	0.983	0.994	0.997
$\beta = 5$	0.716	0.768	0.820	0.869	0.919	0.962	0.983	0.993	0.997

TABLE 2
Values of $q = \inf\{z : P(Z_B \leq z) \geq p\}$

p	0.700	0.750	0.800	0.850	0.900	0.950	0.975	0.990	0.995
$\beta = 1$	0.683	0.730	0.777	0.825	0.878	0.932	0.964	0.994	0.997
$\beta = 3$	0.687	0.734	0.781	0.829	0.882	0.935	0.966	0.986	0.992
$\beta = 5$	0.686	0.734	0.782	0.831	0.882	0.936	0.966	0.986	0.994

TABLE 3
Values of $P(Z_B \leq z)$ for various z and β , and $Z_B = Z_B^{(1)}, Z_B^{(2)}, Z_B^{(3)}$

z	$\beta = (1, 1)$			$\beta = (3, 1)$			$\beta = (3, 3)$		
	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$
0.700	0.705	0.752	0.725	0.704	0.741	0.721	0.708	0.735	0.718
0.750	0.762	0.803	0.778	0.760	0.791	0.773	0.762	0.787	0.771
0.800	0.817	0.851	0.832	0.814	0.842	0.827	0.817	0.838	0.825
0.850	0.871	0.898	0.880	0.868	0.889	0.877	0.868	0.885	0.874
0.900	0.921	0.939	0.927	0.917	0.932	0.924	0.918	0.930	0.922
0.950	0.966	0.975	0.968	0.964	0.971	0.966	0.964	0.970	0.965
0.975	0.985	0.989	0.987	0.983	0.987	0.986	0.984	0.986	0.985
0.990	0.995	0.997	0.995	0.995	0.997	0.995	0.995	0.996	0.994
0.995	0.997	0.998	0.998	0.998	0.998	0.998	0.997	0.998	0.997

of nonpositive functional values. This process is repeated 50,000 times for $\beta = (1, 1), (3, 1)$ and $(3, 3)$ to estimate the distribution of $Z_B^{(1)}, Z_B^{(2)}$ or $Z_B^{(3)}$.

Since in higher-dimensional cases, $Z_B^{(1)}$ and $Z_B^{(2)}$ are not equal in distribution and their distribution functions are not symmetric about 0.5, we give both the values of $P(Z_B^{(1)} \leq z)$ and $P(Z_B^{(2)} \leq z)$ for some selected z values in Table 3. The corresponding distribution functions are plotted in Figure 3. We present the quantiles of $Z_B^{(3)}$ with different smoothness levels in Table 4.

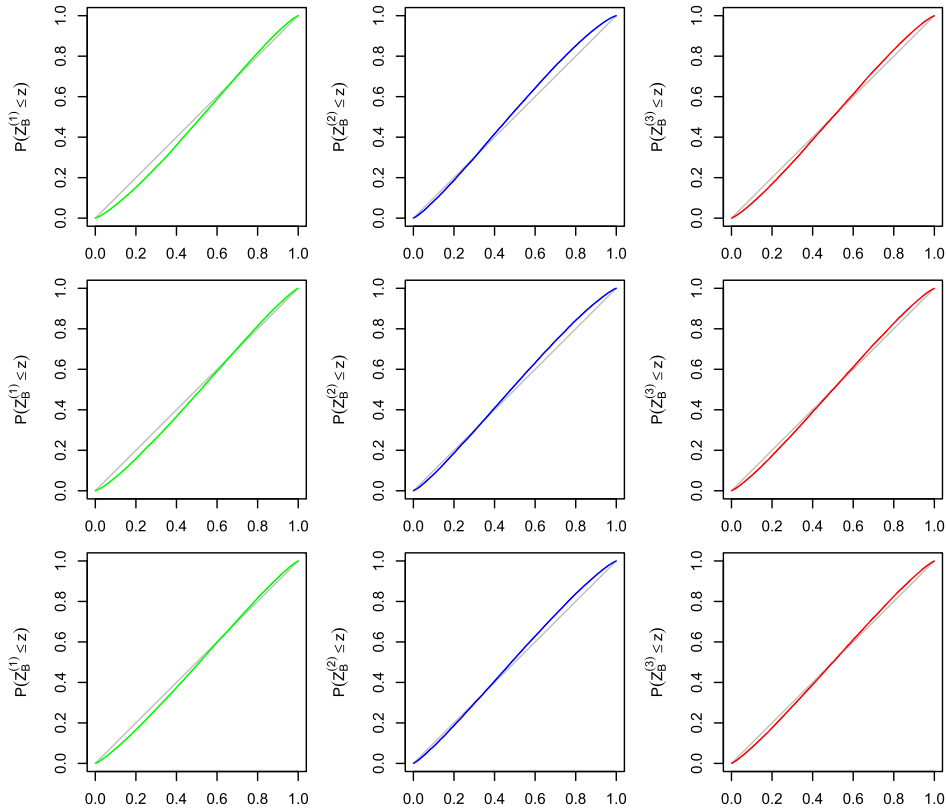


FIG. 3. Distribution functions of $Z_B^{(1)}, Z_B^{(2)}$ or $Z_B^{(3)}$. The three plots in the first row are for $\beta = (1, 1)$; the second row is for $\beta = (3, 1)$; the last row is for $\beta = (3, 3)$.

TABLE 4
Values of $q = \inf\{z : P(Z_B \leq z) \geq p\}$ for various p , and $Z_B = Z_B^{(1)}, Z_B^{(2)}, Z_B^{(3)}$

p	$\beta = (1, 1)$			$\beta = (3, 1)$			$\beta = (3, 3)$		
	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$	$Z_B^{(1)}$	$Z_B^{(2)}$	$Z_B^{(3)}$
0.700	0.697	0.653	0.677	0.699	0.665	0.681	0.695	0.669	0.684
0.750	0.741	0.699	0.724	0.743	0.711	0.728	0.741	0.715	0.732
0.800	0.787	0.749	0.771	0.789	0.759	0.776	0.787	0.763	0.778
0.850	0.833	0.801	0.819	0.835	0.811	0.823	0.833	0.815	0.825
0.900	0.881	0.855	0.872	0.883	0.865	0.876	0.883	0.869	0.878
0.950	0.933	0.917	0.928	0.937	0.925	0.931	0.935	0.927	0.933
0.975	0.963	0.951	0.959	0.965	0.957	0.962	0.965	0.959	0.964
0.990	0.983	0.977	0.982	0.985	0.981	0.984	0.985	0.983	0.984
0.995	0.991	0.987	0.990	0.991	0.989	0.992	0.993	0.991	0.992

5.2. *Comparison with Deng, Han and Zhang's method.* For pointwise inference in multivariate isotonic regression, Deng et al. [21] constructed the confidence interval by the asymptotic distribution of pivotal statistics. Their method will be referred to as DHZ in the following. Let $\hat{u}(\mathbf{x}_0)$ and $\hat{v}(\mathbf{x}_0)$ be such that

$$\begin{aligned}\hat{f}^-(\mathbf{x}_0) &= \max_{\mathbf{u} \leq \mathbf{x}_0} \min_{\substack{v \geq \mathbf{x}_0 \\ \#\{i: X_i \in [\mathbf{u}: \mathbf{v}]\} > 0}} \bar{Y}|_{[\mathbf{u}: \mathbf{v}]} = \min_{\substack{v \geq \mathbf{x}_0 \\ \#\{i: X_i \in [\hat{\mathbf{u}}(\mathbf{x}_0): \mathbf{v}]\} > 0}} \bar{Y}|_{[\hat{\mathbf{u}}(\mathbf{x}_0): \mathbf{v}]}, \\ \hat{f}^+(\mathbf{x}_0) &= \min_{v \geq \mathbf{x}_0} \max_{\substack{\mathbf{u} \leq \mathbf{x}_0 \\ \#\{i: X_i \in [\mathbf{u}: \mathbf{v}]\} > 0}} \bar{Y}|_{[\mathbf{u}: \mathbf{v}]} = \max_{\substack{\mathbf{u} \leq \mathbf{x}_0 \\ \#\{i: X_i \in [\mathbf{u}: \hat{\mathbf{v}}(\mathbf{x}_0)]\} > 0}} \bar{Y}|_{[\mathbf{u}: \hat{\mathbf{v}}(\mathbf{x}_0)]},\end{aligned}$$

and $\hat{f}(\mathbf{x}_0) = (\hat{f}^-(\mathbf{x}_0) + \hat{f}^+(\mathbf{x}_0))/2$. Under the same data generating conditions as in Theorem 4.2 and additionally assuming $\mathbf{X}$ is uniform distributed, Deng et al. [21] showed that

$$\frac{\sqrt{\#\{i: X_i \in [\hat{\mathbf{u}}(\mathbf{x}_0): \hat{\mathbf{v}}(\mathbf{x}_0)]\}}}{\sigma} (\hat{f}(\mathbf{x}_0) - f(\mathbf{x}_0)) \rightsquigarrow K_\beta,$$

where K_β is a universal distribution that depends solely on the local regularity β . Let $1 - \gamma \in (0.5, 1)$ be the confidence level. They proposed the following confidence interval for $f_0(\mathbf{x}_0)$:

$$(5.2) \quad \left[\hat{f}(\mathbf{x}_0) - \frac{c_\gamma \hat{\sigma}}{\sqrt{\#\{i: X_i \in [\hat{\mathbf{u}}(\mathbf{x}_0): \hat{\mathbf{v}}(\mathbf{x}_0)]\}}}, \hat{f}(\mathbf{x}_0) + \frac{c_\gamma \hat{\sigma}}{\sqrt{\#\{i: X_i \in [\hat{\mathbf{u}}(\mathbf{x}_0): \hat{\mathbf{v}}(\mathbf{x}_0)]\}} \right],$$

where c_γ is the critical value obtained by simulating the limiting distribution of K_β and $\hat{\sigma}$ is a consistent estimator of σ .

We consider five regression functions: (1) $f_1(x_1, x_2) = (x_1 + x_2)^2$; (2) $f_2(x_1, x_2) = \sqrt{x_1 + x_2}$; (3) $f_3(x_1, x_2) = x_1 x_2$; (4) $f_4(x_1, x_2) = e^{x_1 + x_2}$; (5) $f_5(x_1, x_2) = e^{x_1 x_2}$. Set $\varepsilon_i \sim N(0, 1)$ and $X_1, X_2 \sim \text{Unif}(0, 1)$, mutually independent, for $i = 1, \dots, n$. We consider sample sizes $n = 200, 500, 1000$ and 2000 . To construct credible intervals, we choose $J = \lceil n^{1/3} \log(\log n) \rceil$. We compare the coverage and length of our immersion credible interval (IB), the recalibrated credible interval (IB(adj)) and the DHZ's confidence interval under two credible/confidence levels 0.95 and 0.90. The coverage percentage and the average length are calculated over 2000 replications. The result is summarized in Table 5.

The unadjusted credible intervals generally overcover the true function value for larger sample sizes, whereas the recalibrated credible intervals provide more accurate coverage to different extents for different functions. DHZ's method yields more precise coverage at the given confidence level when the sample sizes are relatively smaller. However, our credible

TABLE 5
Coverage percentage (C) and length (L) comparison

<i>f</i>	Level	<i>n</i>	IB		IB(adj)		DHZ	
			C	L	C	L	C	L
<i>f</i> ₁	0.05	200	93.6	0.903 (0.145)	90.0	0.805 (0.132)	92.4	1.138 (0.600)
		500	98.8	0.777 (0.111)	97.7	0.692 (0.101)	95.0	0.959 (0.490)
		1000	97.0	0.630 (0.086)	94.4	0.562 (0.078)	94.8	0.827 (0.435)
		2000	97.4	0.535 (0.072)	95.4	0.476 (0.066)	94.9	0.686 (0.324)
	0.10	200	88.1	0.761 (0.126)	81.5	0.668 (0.112)	86.1	0.898 (0.473)
		500	96.9	0.656 (0.097)	94.4	0.576 (0.087)	90.0	0.757 (0.387)
		1000	92.8	0.532 (0.075)	88.8	0.467 (0.068)	88.8	0.652 (0.343)
		2000	94.3	0.451 (0.063)	90.4	0.397 (0.057)	89.7	0.541 (0.256)
<i>f</i> ₂	0.05	200	91.0	0.503 (0.089)	87.0	0.447 (0.081)	95.2	0.722 (0.339)
		500	96.6	0.380 (0.061)	93.8	0.338 (0.055)	95.4	0.546 (0.303)
		1000	94.9	0.308 (0.047)	91.7	0.274 (0.043)	94.8	0.439 (0.252)
		2000	95.9	0.253 (0.039)	93.3	0.225 (0.035)	95.3	0.357 (0.175)
	0.10	200	85.0	0.423 (0.077)	79.0	0.371 (0.068)	89.8	0.570 (0.268)
		500	92.7	0.320 (0.052)	88.0	0.280 (0.047)	90.3	0.431 (0.239)
		1000	90.0	0.259 (0.040)	84.7	0.227 (0.036)	88.9	0.346 (0.199)
		2000	91.8	0.213 (0.033)	87.0	0.186 (0.030)	90.6	0.281 (0.138)
<i>f</i> ₃	0.05	200	91.8	0.476 (0.084)	87.2	0.423 (0.076)	94.7	0.740 (0.410)
		500	96.0	0.371 (0.061)	93.4	0.329 (0.055)	95.0	0.532 (0.246)
		1000	95.0	0.293 (0.046)	91.6	0.260 (0.042)	95.4	0.433 (0.207)
		2000	95.6	0.242 (0.037)	93.0	0.215 (0.034)	94.8	0.353 (0.165)
	0.10	200	84.9	0.400 (0.072)	79.6	0.350 (0.064)	89.4	0.584 (0.323)
		500	92.2	0.311 (0.053)	87.8	0.273 (0.047)	89.7	0.419 (0.194)
		1000	90.0	0.246 (0.040)	84.7	0.216 (0.036)	90.3	0.341 (0.163)
		2000	91.6	0.204 (0.033)	86.9	0.178 (0.029)	89.8	0.279 (0.131)
<i>f</i> ₄	0.05	200	97.0	1.260 (0.188)	94.4	1.122 (0.170)	89.2	1.234 (0.621)
		500	99.8	1.086 (0.133)	99.5	0.968 (0.120)	92.8	1.077 (0.538)
		1000	99.0	0.869 (0.100)	97.9	0.774 (0.092)	94.4	0.927 (0.437)
		2000	99.7	0.728 (0.083)	98.2	0.649 (0.075)	94.8	0.800 (0.405)
	0.10	200	92.8	1.063 (0.162)	87.9	0.932 (0.144)	83.4	0.974 (0.490)
		500	99.3	0.917 (0.115)	98.4	0.805 (0.103)	87.0	0.850 (0.424)
		1000	97.0	0.733 (0.088)	93.8	0.644 (0.079)	89.1	0.731 (0.345)
		2000	97.2	0.615 (0.072)	94.9	0.540 (0.064)	89.3	0.631 (0.320)
<i>f</i> ₅	0.05	200	94.0	0.540 (0.093)	90.2	0.480 (0.083)	95.4	0.798 (0.398)
		500	97.4	0.432 (0.068)	95.6	0.384 (0.062)	94.2	0.597 (0.316)
		1000	96.5	0.338 (0.051)	93.4	0.301 (0.047)	95.1	0.491 (0.265)
		2000	96.9	0.280 (0.042)	94.4	0.249 (0.038)	95.5	0.401 (0.195)
	0.10	200	88.2	0.454 (0.079)	82.6	0.398 (0.070)	90.3	0.629 (0.314)
		500	94.4	0.364 (0.059)	90.6	0.319 (0.053)	89.1	0.471 (0.249)
		1000	92.2	0.285 (0.045)	87.8	0.249 (0.040)	91.1	0.387 (0.209)
		2000	93.4	0.236 (0.036)	89.2	0.207 (0.033)	90.1	0.317 (0.154)

intervals are generally shorter and exhibit less variation compared to DHZ’s confidence intervals. The variation observed in our method across different regression functions may be attributed to the roughness of the partition used. In practical applications, a slightly larger *J* can be set, provided that the credible intervals can be computed within a reasonable time.

6. Proofs.

6.1. *Proof of Theorem 4.1.* For $\mathbf{t} \in \mathbb{R}^d$, let $j(\mathbf{t}) = \lceil (\mathbf{x}_0 + \mathbf{t} \circ \mathbf{r}_n) \circ \mathbf{J} \rceil$. Let

$$(6.1) \quad f_{*,c}(\mathbf{x}_0) = \max_{\substack{c^{-\gamma} \mathbf{1} \leq \mathbf{u} \leq c \mathbf{1}, \\ u_k \leq x_{0,k}, \\ s+1 \leq k \leq d}} \min_{\substack{c^{-\gamma} \mathbf{1} \leq \mathbf{v} \leq c \mathbf{1}, \\ v_k \leq 1-x_{0,k}, \\ s+1 \leq k \leq d}} \frac{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j \theta_j}{N_{[j(-\mathbf{u}):j(\mathbf{v})]}},$$

where γ is a positive constant to be determined later. We also introduce the notation $W_n^* = \omega_n^{-1}(f_*(\mathbf{x}_0) - f_0(\mathbf{x}_0))$, $W_{n,c}^* = \omega_n^{-1}(f_{*,c}(\mathbf{x}_0) - f_0(\mathbf{x}_0))$,

$$\begin{aligned} W_c &= \sup_{\substack{c^{-\gamma} \mathbf{1} \leq \mathbf{u} \leq c \mathbf{1}, \\ u_k \leq x_{0,k}, \\ s+1 \leq k \leq d}} \inf_{\substack{c^{-\gamma} \mathbf{1} \leq \mathbf{v} \leq c \mathbf{1}, \\ v_k \leq 1-x_{0,k}, \\ s+1 \leq k \leq d}} \left\{ \frac{\sigma_0 H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ &\quad \left. + \frac{\sigma_0 H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \sum_{l \in L^*} \frac{\partial^l f_0(\mathbf{x}_0)}{(l+1)!} \prod_{k=1}^s \frac{v_k^{l_k+1} - (-u_k)^{l_k+1}}{u_k + v_k} \right\}, \\ W &= \sup_{\substack{\mathbf{u} > \mathbf{0}, \\ u_k \leq x_{0,k}, \\ s+1 \leq k \leq d}} \inf_{\substack{\mathbf{v} \geq \mathbf{0}, \\ v_k \leq 1-x_{0,k}, \\ s+1 \leq k \leq d}} \left\{ \frac{\sigma_0 H_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ &\quad \left. + \frac{\sigma_0 H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} + \sum_{l \in L^*} \frac{\partial^l f_0(\mathbf{x}_0)}{(l+1)!} \prod_{k=1}^s \frac{v_k^{l_k+1} - (-u_k)^{l_k+1}}{u_k + v_k} \right\}. \end{aligned}$$

The proof of the theorem is carried out in several steps using Lemma B.1 of [54], presented as lemmas below.

LEMMA 6.1. *Under the conditions of Theorem 4.1, for every $c > 0$ and $\gamma > 0$, $\mathcal{L}(W_{n,c}^* | \mathbb{D}_n)$ converges weakly to $\mathcal{L}(W_c | H_1)$ as random probability measures.*

PROOF. For every $\mathbf{u}, \mathbf{v} \geq \mathbf{0}$, we can write

$$(6.2) \quad \frac{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j \theta_j}{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j} - f_0(\mathbf{x}_0) = A_n(\mathbf{u}, \mathbf{v}; \boldsymbol{\theta}) + A'_n(\mathbf{u}, \mathbf{v}) + B_n(\mathbf{u}, \mathbf{v}),$$

and then $W_{n,c}^* = \max_{c^{-\gamma} \mathbf{1} \leq \mathbf{u} \leq c \mathbf{1}} \min_{c^{-\gamma} \mathbf{1} \leq \mathbf{v} \leq c \mathbf{1}} \{A_n(\mathbf{u}, \mathbf{v}; \boldsymbol{\theta}) + A'_n(\mathbf{u}, \mathbf{v}) + B_n(\mathbf{u}, \mathbf{v})\}$, where

$$(6.3) \quad A_n(\mathbf{u}, \mathbf{v}; \boldsymbol{\theta}) = \omega_n^{-1} \frac{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j (\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n])}{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j},$$

$$(6.4) \quad A'_n(\mathbf{u}, \mathbf{v}) = \omega_n^{-1} \frac{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j (\mathbb{E}[\theta_j | \mathbb{D}_n] - \bar{Y}_{I_j})}{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j},$$

$$(6.5) \quad B_n(\mathbf{u}, \mathbf{v}) = \omega_n^{-1} (\bar{Y}_{I_{[j(-\mathbf{u}):j(\mathbf{v})]}} - f_0(\mathbf{x}_0)).$$

Since the max-min functional is continuous on the space $\mathbb{L}_\infty([c^{-\gamma} \mathbf{1}, c \mathbf{1}] \times [c^{-\gamma} \mathbf{1}, c \mathbf{1}])$, it suffices to show that $A_n + A'_n + B_n$ converges weakly in $\mathbb{L}_\infty([c^{-\gamma} \mathbf{1}, c \mathbf{1}] \times [c^{-\gamma} \mathbf{1}, c \mathbf{1}])$, conditional on the data $\mathbb{D}_n$. By Lemma B.2 of [54] and Lemma 6.2, we prove the weak convergence of A_n . We show that A'_n converges to zero uniformly in Lemma 6.3. The convergence of B_n is completed by combining Lemma B.2 of [54], Lemma 6.4 and Lemma 6.5. $\square$

LEMMA 6.2. *Under the conditions of Theorem 4.1, for every $c > 0$, let $\mathbb{H}_{2,n}(\mathbf{u}, \mathbf{v}; \boldsymbol{\theta}) = \omega_n \sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n])$. Then $\mathbb{H}_{2,n}$ converges weakly to a centered Gaussian process H_2 in $\mathbb{L}_\infty([\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}])$ for every $c > 0$ in $\mathbb{P}_0$ -probability.*

PROOF. By (3.7), Lemmas B.2, B.4 and B.5 of [54], the covariance kernel of $\mathbb{H}_{2,n}$ given $(\mathbb{D}_n, \sigma_n^2)$, is given by $\omega_n^2 \sigma_n^2 \sum_{j \in [j(-\mathbf{u} \wedge \mathbf{u}'):j(\mathbf{v} \wedge \mathbf{v}')] } N_j^2 / (N_j + \lambda_j^{-2})$, which converges in $\mathbb{P}_0$ -probability to $\sigma_0^2 \prod_{k=1}^s (u_k \wedge u'_k + v_k \wedge v'_k) D_s(\mathbf{u} \wedge \mathbf{u}', \mathbf{v} \wedge \mathbf{v}')$. Thus finite-dimensional distributions of $\mathbb{H}_{2,n}$ converge weakly to those of a centered Gaussian process $\sigma_0 H_2$ in $\mathbb{P}_0$ -probability.

Next, we need to show that $\mathcal{L}(\mathbb{H}_{2,n}(\mathbf{u}, \mathbf{v}; \boldsymbol{\theta}) : (\mathbf{u}, \mathbf{v}) \in [\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}])$ is tight on $\mathbb{L}_\infty([\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}])$ for any $c > 0$ in $\mathbb{P}_0$ -probability. In view of Theorem 18.14 of [53], we need to verify that, for every $\epsilon > 0$ and $\eta > 0$, there exists a finite partition $\{T_p : p \leq K\}$ of $[\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}]$ with K depending only on ϵ and η such that

$$\mathbb{P}\left(\sup_{(\mathbf{u}_1, \mathbf{v}_1), (\mathbf{u}_2, \mathbf{v}_2) \in T_p} \{|\mathbb{H}_{2,n}(\mathbf{u}_1, \mathbf{v}_1) - \mathbb{H}_{2,n}(\mathbf{u}_2, \mathbf{v}_2)| : 1 \leq p \leq K\} > \epsilon | \mathbb{D}_n\right) < \eta$$

with $\mathbb{P}_0$ -probability tending to 1. Let $\delta > 0$, to be determined later, which depends only on ϵ and η . Let $0 = s_0 < s_1 < \dots < s_l = c$ with $(s_{l-1}, s_l]$ of equal length at least δ and $l \leq 2c/\delta$. We choose a partition $\{T_p : p \leq K\}$ of $[\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}]$ to be

$$(6.6) \quad \mathcal{P}(\delta) = \left\{ \prod_{k=1}^d (s_{t_k-1}, s_{t_k}] \times \prod_{k=1}^d (s_{r_k-1}, s_{r_k}] : t_k, r_k \in \{1, \dots, l\} \right\},$$

with cardinality $K = \#\mathcal{P}(\delta) = l^{2d}$. It suffices to verify that, for any $p \leq K$,

$$\mathbb{P}\left(\sup_{(\mathbf{u}_1, \mathbf{v}_1), (\mathbf{u}_2, \mathbf{v}_2) \in T_p} \{|\mathbb{H}_{2,n}(\mathbf{u}_1, \mathbf{v}_1) - \mathbb{H}_{2,n}(\mathbf{u}_2, \mathbf{v}_2)|\} > \epsilon | \mathbb{D}_n\right) < \eta \left(\frac{\delta}{2c}\right)^{2d}.$$

Let $\mathcal{J}(\mathbf{u}, \mathbf{v}) = [j(-\mathbf{u}) : j(\mathbf{v})]$. For $(\mathbf{u}_1, \mathbf{v}_1), (\mathbf{u}_2, \mathbf{v}_2)$, we write $\mathbb{H}_{2,n}(\mathbf{u}_1, \mathbf{v}_1) - \mathbb{H}_{2,n}(\mathbf{u}_2, \mathbf{v}_2)$ as the difference of the sums of $\omega_n N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n])$ over the sets $\mathcal{J}(\mathbf{u}_1, \mathbf{v}_1) \setminus \mathcal{J}(\mathbf{u}_1 \wedge \mathbf{u}_2, \mathbf{v}_1 \wedge \mathbf{v}_2)$ and $\mathcal{J}(\mathbf{u}_2, \mathbf{v}_2) \setminus \mathcal{J}(\mathbf{u}_1 \wedge \mathbf{u}_2, \mathbf{v}_1 \wedge \mathbf{v}_2)$, after canceling out the common terms. Thus, its absolute value can be bounded by the sum of the corresponding absolute values over these two index sets. To verify tightness, it then suffices to show that

$$\mathbb{P}\left(\max \left\{ \omega_n \left| \sum_{\mathcal{J}(\mathbf{u}, \mathbf{v}) \setminus \mathcal{J}(s_{t-1}, s_{r-1})} N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n]) \right| : (\mathbf{u}, \mathbf{v}) \in T_p \right\} > \frac{\epsilon}{4} | \mathbb{D}_n \right)$$

is bounded by $\eta(\delta/(2c))^{2d}/4$, with $T_p = \prod_{k=1}^d (s_{t_k-1}, s_{t_k}] \times \prod_{k=1}^d (s_{r_k-1}, s_{r_k}]$, for any $s_t = (s_{t_1}, \dots, s_{t_d})$ and $s_r = (s_{r_1}, \dots, s_{r_d})$.

Let $S_{(-j(-\mathbf{u}), j(\mathbf{v}))} = \sum_{\mathcal{J}(\mathbf{u}, \mathbf{v}) \setminus \mathcal{J}(s_{t-1}, s_{r-1})} N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n])$, a collection of random variables indexed by a $2d$ -dimensional vector in a finite-index set. The negative sign in front of $j(-\mathbf{u})$ in the subscript of S is to make the σ -fields,

$$\mathcal{F}_j^{(k)} = \begin{cases} \sigma\langle N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n]) : -(j(s_{t-1}))_k < -(j(-\mathbf{u}))_k \leq j \rangle & \text{if } k \leq d, \\ \sigma\langle N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n]) : (j(s_{r-1}))_{k-d} < (j(\mathbf{v}))_{k-d} \leq j \rangle & \text{if } k > d. \end{cases}$$

increase with respect to each of the first d components in the subscript. In the sum above, all j are in $\mathcal{J}(s_t, s_r) \setminus \mathcal{J}(s_{t-1}, s_{r-1})$. We note that for every $k \leq 2d$, the random sequence $\{S_{(j_1, \dots, j_{k-1}, j, j_{k+1}, \dots, j_{2d})}, \mathcal{F}_j^{(k)}\}$ is a martingale. Applying Lemma B.6 of [54] with $p = 4d + 2$, we can get an upper bound of the probability of the maximal deviation needed to verify tightness to be a constant multiple of

$$(6.7) \quad (\omega_n/\epsilon)^{(4d+2)} \mathbb{E}\left(\left| \sum_{\mathcal{J}(s_t, s_r) \setminus \mathcal{J}(s_{t-1}, s_{r-1})} N_j(\theta_j - \mathbb{E}[\theta_j | \mathbb{D}_n]) \right|^{4d+2} | \mathbb{D}_n \right).$$

Observe that $\#\mathcal{J}(s_t, s_r) \leq \prod_k (r_{n,k} J_k(s_{t_k} + s_{r_k}) + 2)$, $\#\mathcal{J}(s_{t-1}, s_{r-1}) \geq \prod_k r_{n,k} J_k(s_{t_k} + s_{r_k} - 2\delta)$. As $\delta \leq s_{t_k}, s_{r_k} \leq c$ and $J_k \gg r_{n,k}^{-1}$, it follows that the cardinality of the index set $\mathcal{J}(s_t, s_r) \setminus \mathcal{J}(s_{t-1}, s_{r-1})$ is bounded by a multiple of

$$\prod_{k=1}^d r_{n,k} J_k \left(\prod_{k=1}^d (s_{t_k} + s_{r_k}) - \prod_{k=1}^d (s_{t_k} + s_{r_k} - 2\delta) \right) \leq (2d\delta)(2c)^{d-1} \prod_{k=1}^d r_{n,k} J_k,$$

where the last inequality follows from Lemma B.7 of [54].

The variance $\sigma_n^2 N_j^2 / (N_j + \lambda_j^{-2}) \lesssim n(\prod_{k=1}^d J_k)^{-1}$ with P_0 -probability tending to 1 by Lemma B.4 of [54]. Hence, (6.7) is bounded by a constant multiple of

$$\begin{aligned} & \epsilon^{-(4d+2)} \omega_n^{4d+2} \left(\frac{n\#\mathcal{J}(s_t, s_r) \setminus \mathcal{J}(s_{t-1}, s_{r-1})}{\prod_{k=1}^d J_k} \right)^{2d+1} \\ & \lesssim \epsilon^{-(4d+2)} \omega_n^{4d+2} \left(\prod_{k=1}^d r_{n,k} \right)^{2d+1} n^{2d+1} \delta^{2d+1}, \end{aligned}$$

which simplifies to $\epsilon^{-(4d+2)} \delta^{2d+1}$. With δ chosen a sufficiently small constant multiple of $\eta \epsilon^{4d+2}$, the tightness condition is verified. $\square$

LEMMA 6.3. *Under the conditions of Theorem 4.1, $A'_n(\mathbf{u}, \mathbf{v})$ converges to 0 in P_0 -probability uniformly in $(\mathbf{u}, \mathbf{v})$.*

PROOF. Let $E_n = \{a_1 n / (2 \prod_{k=1}^d J_k) \leq N_j \leq 2a_2 n / (\prod_{k=1}^d J_k)\}$ for some $a_1, a_2 > 0$ and $\bar{\epsilon}|_{I_j} = \sum_{i \in I_j} \epsilon_i / N_j$. By Lemma B.4 of [54], we have for every $T > 0$,

$$(6.8) \quad P_0\left(\max_j |\bar{\epsilon}|_{I_j} > T\right) \leq \sum_j P_0(|\bar{\epsilon}|_{I_j} > T | E_n) + P_0(E_n^c).$$

By Assumption 2 and the Marcinkiewicz—Zygmund inequality,

$$E(|\bar{\epsilon}|_{I_j}|^{2(\sum_{k=1}^s \beta_k^{-1} + 1)} | E_n) \lesssim \left(a_1 n / \left(2 \prod_{k=1}^d J_k \right) \right)^{-(\sum_{k=1}^s \beta_k^{-1} + 1)}.$$

Then (6.8) is bounded by a constant multiple of $(\prod_{k=1}^d J_k)^{\sum_{k=1}^s \beta_k^{-1} + 2} n^{-(\sum_{k=1}^s \beta_k^{-1} + 1)} + o(1)$, which tends to zero because $\prod_{k=1}^d J_k \ll n\omega_n = n^{(\sum_{k=1}^s \beta_k^{-1} + 1)/(\sum_{k=1}^s \beta_k^{-1} + 2)}$.

On the other hand, $\max_j |f_0(\bar{X}_i)|_{I_j} \leq f_0(\mathbf{1})$. Thus, $\max_j |\bar{Y}|_{I_j} = O_{P_0}(1)$. Because $E[\theta_j | \mathbb{D}_n] = (N_j \bar{Y}|_{I_j} + \zeta_j \lambda_j^{-2}) / (N_j + \lambda_j^{-2})$, on the event E_n ,

$$\begin{aligned} (6.9) \quad |A'_n(\mathbf{u}, \mathbf{v})| &= \omega_n^{-1} \left| \frac{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} \lambda_j^{-2} N_j (N_j + \lambda_j^{-2})^{-1} (\zeta_j - \bar{Y}|_{I_j})}{\sum_{j \in [j(-\mathbf{u}):j(\mathbf{v})]} N_j} \right| \\ &\lesssim \omega_n^{-1} \left(\max_j |\bar{Y}|_{I_j} + \zeta_j \right) \left(\min_j N_j \right)^{-1}, \end{aligned}$$

which is of the order of $(n\omega_n)^{-1} \prod_{k=1}^d J_k$ in P_0 -probability. As $\prod_{k=1}^d J_k \ll n\omega_n$ and $P_0(E_n) \rightarrow 1$, we can conclude $A'_n(\mathbf{u}, \mathbf{v}) \rightarrow_{P_0} 0$ uniformly for any $\mathbf{u} \geq \mathbf{0}$ and $\mathbf{v} \geq \mathbf{0}$ provided that $\mathbf{x}_0 - \mathbf{u} \circ \mathbf{r}_n$ and $\mathbf{x}_0 + \mathbf{v} \circ \mathbf{r}_n$ in $[0, 1]^d$. $\square$

To establish the weak convergence of B_n in $\mathbb{L}_\infty([\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}])$, write

$$(6.10) \quad B_n(\mathbf{u}, \mathbf{v}) = \omega_n^{-1} (\bar{\epsilon}|_{I_{[j(-\mathbf{u}):j(\mathbf{v})]}} + \overline{f_0(\bar{X})}|_{I_{[j(-\mathbf{u}):j(\mathbf{v})]}} - f_0(\mathbf{x}_0)).$$

LEMMA 6.4. Let $Z_{ni}(\mathbf{u}, \mathbf{v}) = \omega_n \varepsilon_i \mathbb{1}_{\{X_i \in I_{[j(-\mathbf{u}):j(\mathbf{v})]}\}}$ and $\mathbb{H}_{1,n}(\mathbf{u}, \mathbf{v}) = \sum_{i=1}^n Z_{ni}(\mathbf{u}, \mathbf{v})$. Under the conditions of Theorem 4.1, $\mathbb{H}_{1,n}(\mathbf{u}, \mathbf{v}) \rightsquigarrow \sigma_0 H_1(\mathbf{u}, \mathbf{v})$ in $\mathbb{L}_\infty([\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}])$.

LEMMA 6.5. Under the conditions of Theorem 4.1, for any $c > 0$, uniformly in $(\mathbf{u}, \mathbf{v}) \in [\mathbf{0}, c\mathbf{1}] \times [\mathbf{0}, c\mathbf{1}]$, we have

$$\omega_n^{-1}(\overline{f_0(\mathbf{X})}|_{I_{[j(-\mathbf{u}):j(\mathbf{v})]}} - f_0(\mathbf{x}_0)) \rightarrow_{P_0} \sum_{l \in L^*} \frac{\partial^l f_0(\mathbf{x}_0)}{(l+1)!} \prod_{k=1}^s \frac{v_k^{l_k+1} - (-u_k)^{l_k+1}}{u_k + v_k}.$$

LEMMA 6.6. Under the conditions of Theorem 4.1, for any $M_n \uparrow \infty$, $\Pi(|f_*(\mathbf{x}_0) - f_0(\mathbf{x}_0)| > M_n \omega_n |D_n) \rightarrow 0$ in P_0 -probability.

With the aid of Lemma 6.6, the second condition of Proposition B.1 of [54] is verified by Lemma 6.7 in the following.

LEMMA 6.7. Let $\mathbf{u}^*$ and $\mathbf{v}^*$ be any pair indexes such that

$$(6.11) \quad f_*(\mathbf{x}_0) = \max_{\mathbf{u} \geq \mathbf{0}} \min_{\mathbf{v} \geq \mathbf{0}} \frac{\sum_{[j(-\mathbf{u}):j(\mathbf{v})]} N_j \theta_j}{\sum_{[j(-\mathbf{u}):j(\mathbf{v})]} N_j} = \frac{\sum_{[j(-\mathbf{u}^*):j(\mathbf{v}^*)]} N_j \theta_j}{\sum_{[j(-\mathbf{u}^*):j(\mathbf{v}^*)]} N_j}.$$

Let $\omega_n = n^{-1/(2+\sum_{k=1}^s \beta_k^{-1})}$ and let $\mathbf{r}_n = (\omega_n^{1/\beta_1}, \dots, \omega_n^{1/\beta_s}, 1, \dots, 1)^T$. Suppose that $\mathbf{J}$ satisfies $J_k \gg r_{n,k}^{-1}$, for each $k = 1, \dots, d$, and $\prod_{k=1}^d J_k \ll n\omega_n$. Under Assumptions 1 and 2, there exists $\gamma > 0$ such that

$$\lim_{c \rightarrow \infty} \limsup_{n \rightarrow \infty} \Pi(c^{-\gamma} \leq \min_{1 \leq k \leq d} \{v_k^*\} \leq \max_{1 \leq k \leq d} \{v_k^*\} \leq c |D_n) = 1,$$

in P_0 -probability.

The proofs of Lemmas 6.4–6.7 are provided in [54].

The proof of Theorem 4.1 can now be completed. Using arguments similar to Proposition 7 of [38], it can be verified that $P(W_c \neq W) \rightarrow 0$ as $c \rightarrow \infty$. Hence, the proof follows by an application of Lemma B.1 of [54].

6.2. *Proof of Proposition 4.1.* This can be shown by the self-similarity property of Gaussian processes H_1 and H_2 : for $\mathbf{t} \in \mathbb{R}_{>0}^d$ such that $t_{s+1} = \dots = t_d = 1$, we have that $H_i(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v}) =_d (\prod_{j=1}^s t_j)^{1/2} H_i(\mathbf{u}, \mathbf{v})$, $i = 1, 2$. By the choice of $\mathbf{t}$, multiplying a vector coordinatewise by $\mathbf{t}$ does not change the last $d - s$ coordinates, and thus $D_s(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v}) = D_s(\mathbf{u}, \mathbf{v})$. Then, since a scaling of the domain does not alter suprema and infima, the expression in the limiting distribution is equal to

$$\begin{aligned} & \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\sigma_0 H_1(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v}) + \sigma_0 H_2(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v})}{\prod_{k=1}^s (t_k u_k + t_k v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ & \quad \left. + \sum_{k=1}^s \left[\frac{\partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \cdot \frac{(t_k v_k)^{\beta_k+1} - (-t_k u_k)^{\beta_k+1}}{t_k u_k + t_k v_k} \right] \right\} \\ & =_d \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \left(\sigma_0^{-2} \prod_{j=1}^s t_j \right)^{-1/2} \frac{H_1(\mathbf{u}, \mathbf{v}) + H_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^s (u_k + v_k) D_s(\mathbf{u}, \mathbf{v})} \right. \\ & \quad \left. + \sum_{k=1}^s \left[\frac{t_k^{\beta_k} \partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \cdot \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right] \right\}. \end{aligned}$$

By equating $(\sigma_0^{-2} \prod_{j=1}^s t_j)^{-1/2}$ to $t_k^{\beta_k} \partial_k^{\beta_k} f_0(\mathbf{x}_0)/(\beta_k + 1)!$ for each $k = 1, \dots, s$, we can find the solution t_k to the system of equations, and also the common factor A_β as stated in the proposition.

If $s = d$, then $D_d(\mathbf{u}, \mathbf{v}) = g(\mathbf{x}_0)$ and $H_i(\mathbf{u}, \mathbf{v}) =_d \sqrt{g(\mathbf{x}_0)} \tilde{H}_i(\mathbf{u}, \mathbf{v})$. For $\mathbf{t} \in \mathbb{R}_{>0}^d$, $\tilde{H}_i(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v}) =_d (\prod_{j=1}^d t_j)^{1/2} \tilde{H}_i(\mathbf{u}, \mathbf{v})$, $i = 1, 2$. Hence by self-similarity, the last expression reduces to

$$\begin{aligned} & \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\sigma_0}{\sqrt{g(\mathbf{x}_0)}} \left(\frac{\tilde{H}_1(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v})}{\prod_{k=1}^d (t_k u_k + t_k v_k)} + \frac{\tilde{H}_2(\mathbf{t} \circ \mathbf{u}, \mathbf{t} \circ \mathbf{v})}{\prod_{k=1}^d (t_k u_k + t_k v_k)} \right) \right. \\ & \quad \left. + \sum_{k=1}^d \left[\frac{\partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \cdot \frac{(t_k v_k)^{\beta_k+1} - (-t_k u_k)^{\beta_k+1}}{t_k u_k + t_k v_k} \right] \right\} \\ & =_d \sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \left\{ \sqrt{\frac{\sigma_0^2}{g(\mathbf{x}_0) \prod_{j=1}^d t_j}} \left(\frac{\tilde{H}_1(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \frac{\tilde{H}_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} \right) \right. \\ & \quad \left. + \sum_{k=1}^d \left[\frac{t_k^{\beta_k} \partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \cdot \frac{v_k^{\beta_k+1} - (-u_k)^{\beta_k+1}}{u_k + v_k} \right] \right\}. \end{aligned}$$

By exploring the equation system for t_k as follows:

$$\sqrt{\frac{\sigma_0^2}{g(\mathbf{x}_0) \prod_{j=1}^d t_j}} = \frac{t_k^{\beta_k} \partial_k^{\beta_k} f_0(\mathbf{x}_0)}{(\beta_k + 1)!} \quad \text{for } k = 1, \dots, d,$$

we can find the common factor $\tilde{A}_\beta$ in a similar way of solving a set of equations.

PROOF OF PROPOSITION 4.2. For $0 \leq z \leq 1$,

$$\begin{aligned} \mathbb{P}(Z_B^{(1)} \leq z) &= \mathbb{P}(1 - Z_B^{(1)} \geq 1 - z) \\ &= \mathbb{P}\left(\mathbb{P}\left(-\sup_{\mathbf{u} \geq \mathbf{0}} \inf_{\mathbf{v} \geq \mathbf{0}} \tilde{U}(\mathbf{u}, \mathbf{v}) \leq 0 \mid \tilde{H}_1\right) \geq 1 - z\right) \\ &= \mathbb{P}\left(\mathbb{P}\left(\inf_{\mathbf{u} \geq \mathbf{0}} \sup_{\mathbf{v} \geq \mathbf{0}} [-\tilde{U}(\mathbf{u}, \mathbf{v})] \leq 0 \mid \tilde{H}_1\right) \geq 1 - z\right). \end{aligned}$$

Note that $\tilde{H}_i(\mathbf{u}, \mathbf{v}) =_d \tilde{H}_i(\mathbf{v}, \mathbf{u})$ and $\tilde{H}_i =_d -\tilde{H}_i$ for $i = 1, 2$. Denote $\tilde{H}_1^* = -\tilde{H}_1$. Then we have

$$\begin{aligned} & \mathbb{P}\left(\inf_{\mathbf{u} \geq \mathbf{0}} \sup_{\mathbf{v} \geq \mathbf{0}} [-\tilde{U}(\mathbf{u}, \mathbf{v})] \leq 0 \mid \tilde{H}_1\right) \\ & =_d \mathbb{P}\left(\inf_{\mathbf{u} \geq \mathbf{0}} \sup_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\tilde{H}_1^*(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \frac{-\tilde{H}_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} \right. \right. \\ & \quad \left. \left. + \sum_{k=1}^d \frac{-v_k^{\beta_k+1} + (-u_k)^{\beta_k+1}}{u_k + v_k} \right\} \leq 0 \mid -\tilde{H}_1^*\right) \\ & =_d \mathbb{P}\left(\inf_{\mathbf{u} \geq \mathbf{0}} \sup_{\mathbf{v} \geq \mathbf{0}} \left\{ \frac{\tilde{H}_1^*(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \frac{\tilde{H}_2(\mathbf{u}, \mathbf{v})}{\prod_{k=1}^d (u_k + v_k)} + \sum_{k=1}^d \frac{u_k^{\beta_k+1} - v_k^{\beta_k+1}}{u_k + v_k} \right\} \leq 0 \mid \tilde{H}_1^*\right) \end{aligned}$$

$$\begin{aligned}
&= {}_d P \left(\inf_{u \geq 0} \sup_{v \geq 0} \left\{ \frac{\tilde{H}_1^*(v, u)}{\prod_{k=1}^d (u_k + v_k)} + \frac{\tilde{H}_2(v, u)}{\prod_{k=1}^d (u_k + v_k)} + \sum_{k=1}^d \frac{u_k^{\beta_k+1} - v_k^{\beta_k+1}}{u_k + v_k} \right\} \leq 0 \mid \tilde{H}_1^* \right) \\
&= P \left(\inf_{u \geq 0} \sup_{v \geq 0} \tilde{U}(v, u) \leq 0 \mid \tilde{H}_1^* \right).
\end{aligned}$$

Hence, $P(Z_B^{(1)} \leq z) = P(Z_B^{(2)} \geq 1 - z)$. The symmetry of the distribution of $Z_B^{(3)}$ holds by similar arguments. $\square$

Acknowledgments. The authors would like to thank the anonymous referees, an Associate Editor and the Editor for their constructive comments that improved the quality of this paper.

Funding. The second author was supported in part by NSF Grant number DMS-1916419.

SUPPLEMENTARY MATERIAL

Supplement: Additional proofs and supporting lemmas (DOI: [10.1214/23-AOS2298SUPP](https://doi.org/10.1214/23-AOS2298SUPP); .pdf). We provide remaining proofs and all the supporting lemmas and propositions in the Supplementary Material [54].

REFERENCES

- [1] AYER, M., BRUNK, H. D., EWING, G. M., REID, W. T. and SILVERMAN, E. (1955). An empirical distribution function for sampling with incomplete information. *Ann. Math. Stat.* **26** 641–647. [MR0073895](https://doi.org/10.1214/aoms/1177728423) <https://doi.org/10.1214/aoms/1177728423>
- [2] BANERJEE, M. (2007). Likelihood based inference for monotone response models. *Ann. Statist.* **35** 931–956. [MR2341693](https://doi.org/10.1214/009053606000001578) <https://doi.org/10.1214/009053606000001578>
- [3] BANERJEE, M. and WELLNER, J. A. (2001). Likelihood ratio tests for monotone functions. *Ann. Statist.* **29** 1699–1731. [MR1891743](https://doi.org/10.1214/aos/1015345959) <https://doi.org/10.1214/aos/1015345959>
- [4] BARLOW, R. E., BARTHOLOMEW, D. J., BREMNER, J. M. and BRUNK, H. D. (1972). *Statistical Inference Under Order Restrictions. The Theory and Application of Isotonic Regression*. Wiley Series in Probability and Mathematical Statistics. Wiley, London–Sydney. [MR0326887](https://doi.org/10.1214/aos/1015345959)
- [5] BELLEC, P. C. (2018). Sharp oracle inequalities for least squares estimators in shape restricted regression. *Ann. Statist.* **46** 745–780. [MR3782383](https://doi.org/10.1214/17-AOS1566) <https://doi.org/10.1214/17-AOS1566>
- [6] BHAUMIK, P. and GHOSAL, S. (2015). Bayesian two-step estimation in differential equation models. *Electron. J. Stat.* **9** 3124–3154. [MR3453972](https://doi.org/10.1214/15-EJS1099) <https://doi.org/10.1214/15-EJS1099>
- [7] BHAUMIK, P. and GHOSAL, S. (2017). Efficient Bayesian estimation and uncertainty quantification in ordinary differential equation models. *Bernoulli* **23** 3537–3570. [MR3654815](https://doi.org/10.3150/16-BEJ856) <https://doi.org/10.3150/16-BEJ856>
- [8] BHAUMIK, P. and GHOSAL, S. (2017). Bayesian inference for higher-order ordinary differential equation models. *J. Multivariate Anal.* **157** 103–114. [MR3641739](https://doi.org/10.1016/j.jmva.2017.03.003) <https://doi.org/10.1016/j.jmva.2017.03.003>
- [9] BHAUMIK, P., SHI, W. and GHOSAL, S. (2022). Two-step Bayesian methods for generalized regression driven by partial differential equations. *Bernoulli* **28** 1625–1647. [MR4411505](https://doi.org/10.3150/21-bej1363) <https://doi.org/10.3150/21-bej1363>
- [10] BRUNK, H. D. (1955). Maximum likelihood estimates of monotone parameters. *Ann. Math. Stat.* **26** 607–616. [MR0073894](https://doi.org/10.1214/aoms/1177728420) <https://doi.org/10.1214/aoms/1177728420>
- [11] BRUNK, H. D. (1970). Estimation of isotonic regression. In *Nonparametric Techniques in Statistical Inference (Proc. Sympos., Indiana Univ., Bloomington, Ind., 1969)* 177–197. Cambridge Univ. Press, London. [MR0277070](https://doi.org/10.1214/aoms/1177728420)
- [12] CAI, B. and DUNSON, D. B. (2007). Bayesian multivariate isotonic regression splines: Applications to carcinogenicity studies. *J. Amer. Statist. Assoc.* **102** 1158–1171. [MR2412540](https://doi.org/10.1198/016214506000000942) <https://doi.org/10.1198/016214506000000942>
- [13] CAI, T. T., LOW, M. G. and XIA, Y. (2013). Adaptive confidence intervals for regression functions under shape constraints. *Ann. Statist.* **41** 722–750. [MR3099119](https://doi.org/10.1214/12-AOS1068) <https://doi.org/10.1214/12-AOS1068>

- [14] CHAKRABORTY, M. and GHOSAL, S. (2021). Coverage of credible intervals in nonparametric monotone regression. *Ann. Statist.* **49** 1011–1028. MR4255117 <https://doi.org/10.1214/20-aos1989>
- [15] CHAKRABORTY, M. and GHOSAL, S. (2021). Convergence rates for Bayesian estimation and testing in monotone regression. *Electron. J. Stat.* **15** 3478–3503. MR4280172 <https://doi.org/10.1214/21-ejs1861>
- [16] CHAKRABORTY, M. and GHOSAL, S. (2021). Bayesian inference on monotone regression quantile: Coverage and rate acceleration. Preprint.
- [17] CHAKRABORTY, M. and GHOSAL, S. (2022). Rates and coverage for monotone densities using projection-posterior. *Bernoulli* **28** 1093–1119. MR4388931 <https://doi.org/10.3150/21-bej1379>
- [18] CHATTERJEE, S., GUNTUBOYINA, A. and SEN, B. (2015). On risk bounds in isotonic and other shape restricted regression problems. *Ann. Statist.* **43** 1774–1800. MR3357878 <https://doi.org/10.1214/15-AOS1324>
- [19] CHATTERJEE, S., GUNTUBOYINA, A. and SEN, B. (2018). On matrix estimation under monotonicity constraints. *Bernoulli* **24** 1072–1100. MR3706788 <https://doi.org/10.3150/16-BEJ865>
- [20] COX, D. D. (1993). An analysis of Bayesian inference for nonparametric regression. *Ann. Statist.* **21** 903–923. MR1232525 <https://doi.org/10.1214/aos/1176349157>
- [21] DENG, H., HAN, Q. and ZHANG, C.-H. (2021). Confidence intervals for multiple isotonic regression and other monotone models. *Ann. Statist.* **49** 2021–2052. MR4319240 <https://doi.org/10.1214/20-aos2025>
- [22] DENG, H. and ZHANG, C.-H. (2020). Isotonic regression in multi-dimensional spaces and graphs. *Ann. Statist.* **48** 3672–3698. MR4185824 <https://doi.org/10.1214/20-AOS1947>
- [23] DÜMBGEN, L. (2003). Optimal confidence bands for shape-restricted curves. *Bernoulli* **9** 423–449. MR1997491 <https://doi.org/10.3150/bj/1065444812>
- [24] DÜMBGEN, L. and JOHNS, R. B. (2004). Confidence bands for isotonic median curves using sign tests. *J. Comput. Graph. Statist.* **13** 519–533. MR2063998 <https://doi.org/10.1198/1061860043506>
- [25] DÜMBGEN, L. and SPOKOINY, V. G. (2001). Multiscale testing of qualitative hypotheses. *Ann. Statist.* **29** 124–152. MR1833961 <https://doi.org/10.1214/aos/996986504>
- [26] DUROT, C. (2007). On the L_p -error of monotonicity constrained estimators. *Ann. Statist.* **35** 1080–1104. MR2341699 <https://doi.org/10.1214/009053606000001497>
- [27] DUROT, C., KULIKOV, V. N. and LOPUHAÄ, H. P. (2012). The limit distribution of the L_∞ -error of Grenander-type estimators. *Ann. Statist.* **40** 1578–1608. MR3015036 <https://doi.org/10.1214/12-AOS1015>
- [28] FOKIANOS, K., LEUCHT, A. and NEUMANN, M. H. (2020). On integrated L^1 convergence rate of an isotonic regression estimator for multivariate observations. *IEEE Trans. Inf. Theory* **66** 6389–6402. MR4173546 <https://doi.org/10.1109/TIT.2020.3013390>
- [29] GHOSAL, S., SEN, A. and VAN DER VAART, A. W. (2000). Testing monotonicity of regression. *Ann. Statist.* **28** 1054–1082. MR1810919 <https://doi.org/10.1214/aos/1015956707>
- [30] GIJBELS, I., HALL, P., JONES, M. C. and KOCH, I. (2000). Tests for monotonicity of a regression mean with guaranteed level. *Biometrika* **87** 663–673. MR1789816 <https://doi.org/10.1093/biomet/87.3.663>
- [31] GRENANDER, U. (1956). On the theory of mortality measurement. II. *Skand. Aktuarietidskr.* **39** 125–153. MR0093415 <https://doi.org/10.1080/03461238.1956.10414944>
- [32] GROENEBOOM, P. (1985). Estimating a monotone density. In *Proceedings of the Berkeley Conference in Honor of Jerzy Neyman and Jack Kiefer, Vol. II (Berkeley, Calif., 1983)*. Wadsworth Statist./Probab. Ser. 539–555. Wadsworth, Belmont, CA. MR0822052
- [33] GROENEBOOM, P. and JONGBLOED, G. (2014). *Nonparametric Estimation Under Shape Constraints: Estimators, Algorithms and Asymptotics*. Cambridge Series in Statistical and Probabilistic Mathematics **38**. Cambridge Univ. Press, New York. MR3445293 <https://doi.org/10.1017/CBO9781139020893>
- [34] GROENEBOOM, P. and JONGBLOED, G. (2015). Nonparametric confidence intervals for monotone functions. *Ann. Statist.* **43** 2019–2054. MR3375875 <https://doi.org/10.1214/15-AOS1335>
- [35] HALL, P. and HECKMAN, N. E. (2000). Testing for monotonicity of a regression mean by calibrating for linear functions. *Ann. Statist.* **28** 20–39. MR1762902 <https://doi.org/10.1214/aos/1016120363>
- [36] HAN, Q. (2021). Set structured global empirical risk minimizers are rate optimal in general dimensions. *Ann. Statist.* **49** 2642–2671. MR4338378 <https://doi.org/10.1214/21-aos2049>
- [37] HAN, Q., WANG, T., CHATTERJEE, S. and SAMWORTH, R. J. (2019). Isotonic regression in general dimensions. *Ann. Statist.* **47** 2440–2471. MR3988762 <https://doi.org/10.1214/18-AOS1753>
- [38] HAN, Q. and ZHANG, C.-H. (2020). Limit distribution theory for block estimators in multiple isotonic regression. *Ann. Statist.* **48** 3251–3282. MR4185808 <https://doi.org/10.1214/19-AOS1928>
- [39] HUANG, J. and WELLNER, J. A. (1995). Estimation of a monotone density or monotone hazard under random censoring. *Scand. J. Stat.* **22** 3–33. MR1334065
- [40] HUANG, Y. and ZHANG, C.-H. (1994). Estimating a monotone density from censored observations. *Ann. Statist.* **22** 1256–1274. MR1311975 <https://doi.org/10.1214/aos/1176325628>

- [41] KOSOROK, M. R. (2008). Bootstrapping in Grenander estimator. In *Beyond Parametrics in Interdisciplinary Research: Festschrift in Honor of Professor Pranab K. Sen. Inst. Math. Stat. (IMS) Collect.* **1** 282–292. IMS, Beachwood, OH. [MR2462212](#) <https://doi.org/10.1214/1939403070000000202>
- [42] KULIKOV, V. N. and LOPUHAĀ, H. P. (2005). Asymptotic normality of the L_k -error of the Grenander estimator. *Ann. Statist.* **33** 2228–2255. [MR2211085](#) <https://doi.org/10.1214/009053605000000462>
- [43] LIN, L. and DUNSON, D. B. (2014). Bayesian monotone regression using Gaussian process projection. *Biometrika* **101** 303–317. [MR3215349](#) <https://doi.org/10.1093/biomet/ast063>
- [44] NEELON, B. and DUNSON, D. B. (2004). Bayesian isotonic regression and trend analysis. *Biometrics* **60** 398–406. [MR2066274](#) <https://doi.org/10.1111/j.0006-341X.2004.00184.x>
- [45] PRAKASA RAO, B. L. S. (1969). Estimation of a unimodal density. *Sankhyā Ser. A* **31** 23–36. [MR0267677](#)
- [46] ROBERTSON, T., WRIGHT, F. T. and DYKSTRA, R. L. (1988). *Order Restricted Statistical Inference*. Wiley Series in Probability and Mathematical Statistics: Probability and Mathematical Statistics. Wiley, Chichester. [MR0961262](#)
- [47] SAARELA, O. and ARJAS, E. (2011). A method for Bayesian monotonic multiple regression. *Scand. J. Stat.* **38** 499–513. [MR2833843](#) <https://doi.org/10.1111/j.1467-9469.2010.00716.x>
- [48] SALOMOND, J.-B. (2014). Concentration rate and consistency of the posterior distribution for selected priors under monotonicity constraints. *Electron. J. Stat.* **8** 1380–1404. [MR3263126](#) <https://doi.org/10.1214/14-EJS929>
- [49] SALOMOND, J.-B. (2014). Adaptive Bayes test for monotonicity. In *The Contribution of Young Researchers to Bayesian Statistics*. Springer Proc. Math. Stat. **63** 29–33. Springer, Cham. [MR3133254](#) https://doi.org/10.1007/978-3-319-02084-6_7
- [50] SCHMIDT-HIEBER, J., MUNK, A. and DÜMBGEN, L. (2013). Multiscale methods for shape constraints in deconvolution: Confidence statements for qualitative features. *Ann. Statist.* **41** 1299–1328. [MR3113812](#) <https://doi.org/10.1214/13-AOS1089>
- [51] SEN, B., BANERJEE, M. and WOODROOFE, M. (2010). Inconsistency of bootstrap: The Grenander estimator. *Ann. Statist.* **38** 1953–1977. [MR2676880](#) <https://doi.org/10.1214/09-AOS777>
- [52] SHIVELY, T. S., SAGER, T. W. and WALKER, S. G. (2009). A Bayesian approach to non-parametric monotone function estimation. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **71** 159–175. [MR2655528](#) <https://doi.org/10.1111/j.1467-9868.2008.00677.x>
- [53] VAN DER VAART, A. W. (1998). *Asymptotic Statistics*. Cambridge Series in Statistical and Probabilistic Mathematics **3**. Cambridge Univ. Press, Cambridge. [MR1652247](#) <https://doi.org/10.1017/CBO9780511802256>
- [54] WANG, K. and GHOSAL, S. (2023). Supplement to “Coverage of credible intervals in Bayesian multivariate isotonic regression.” <https://doi.org/10.1214/23-AOS2298SUPP>
- [55] WANG, X. (2008). Bayesian free-knot monotone cubic spline regression. *J. Comput. Graph. Statist.* **17** 373–387. [MR2421700](#) <https://doi.org/10.1198/106186008X321077>
- [56] WRIGHT, F. T. (1981). The asymptotic behavior of monotone regression estimates. *Ann. Statist.* **9** 443–448. [MR0606630](#)
- [57] ZHANG, C.-H. (2002). Risk bounds in isotonic regression. *Ann. Statist.* **30** 528–555. [MR1902898](#) <https://doi.org/10.1214/aos/1021379864>