

Sparse-grid sampling recovery and deep ReLU neural networks in high-dimensional approximation

Dinh Dũng^a and Van Kien Nguyen^b

^aVietnam National University, Hanoi, Information Technology Institute

144 Xuan Thuy, Cau Giay, Hanoi, Vietnam

Email: dinhzung@gmail.com

^bDepartment of Mathematics, University of Transport and Communications

No.3 Cau Giay Street, Lang Thuong Ward, Dong Da District, Hanoi, Vietnam

Email: kiennv@utc.edu.vn

July 17, 2020

Abstract

We investigate approximations of functions from the Hölder-Zygmund space of mixed smoothness $H_\infty^\alpha(\mathbb{I}^d)$ defined on the d -dimensional unit cube $\mathbb{I}^d := [0, 1]^d$, by linear algorithms of sparse-grid sampling recovery and by deep ReLU (Rectified Linear Unit) neural networks when the dimension d may be very large. The approximation error is measured in the norm of the isotropic Sobolev space $W_0^{1,p}(\mathbb{I}^d)$. The optimality of this sampling recovery is studied in terms of sampling n -widths. Optimal linear sampling algorithms are constructed on sparse grids using the piece-wise linear B-spline interpolation representation. We prove some tight dimension-dependent bounds of the sampling n -widths explicit in d and n . Based on the results on sampling recovery, we investigate the expressive power of deep ReLU neural networks to approximate functions in Hölder-Zygmund space. Namely, for any function $f \in H_\infty^\alpha(\mathbb{I}^d)$, we explicitly construct a deep ReLU neural network having an output that approximates f in the $W_0^{1,p}(\mathbb{I}^d)$ -norm with a prescribed accuracy ε , and prove tight dimension-dependent bounds of the computation complexity of this approximation, characterized as the number of weights and the depth of this deep ReLU neural network, explicitly in d and ε . Moreover, we show that under a certain restriction the curse of dimensionality can be avoided in the approximations by sparse-grid sampling recovery and deep ReLU neural networks.

Keywords and Phrases: High-dimensional approximation; Hölder-Zygmund spaces of mixed smoothness; Sampling recovery; Sparse grids; Deep ReLU neural networks.

Mathematics Subject Classifications (2000): 41A25; 41A30; 68T99; 82C32; 92B20.

1 Introduction

The purpose of the present paper is to investigate sampling recovery of functions on the d -dimensional unit cube $\mathbb{I}^d := [0, 1]^d$ with Hölder-Zygmund mixed smoothness by truncated on sparse grids tensor product Faber series when the dimension d may be very large, and then apply this sparse-grid sampling recovery to the high-dimensional approximation by deep ReLU (Rectified Linear Unit) neural networks.

The approximation error is measured in the norm of the isotropic Sobolev space $W_0^{1,p}$. We also study the optimality for these approximations.

In recent decades, the high-dimensional approximation by sparse-grid sampling recovery of functions or signals has been of great interest since they can be applied in a striking number of fields such as Information Technology, Mathematical Finance, Chemistry, Quantum Mechanics, Meteorology, and, in particular, in Uncertainty Quantification and Deep Machine Learning. A numerical method for such problems may require a computational cost increasing exponentially in dimension when the accuracy increases. This phenomenon is called the curse of dimensionality, coined by Bellman [3]. Hence for an efficient computation in high dimensional approximation, one of the key prerequisites is that the curse of dimension can be avoided at least to some extent. In some cases this can be obtained, particularly when the functions to be approximated have appropriate mixed smoothness. With this restriction one can apply approximation methods and sampling algorithms constructed on hyperbolic crosses and sparse grids which give a surprising effect since hyperbolic crosses and sparse grids have the number of elements much less than those of standard domains and grids but give the same approximation error. This essentially reduces the computational cost, and therefore makes the problem tractable.

Sparse grids for approximate sampling recovery and integration were first considered by Smolyak [35]. In computational mathematics, the sparse grid approach was initiated by Zenger [45]. There is a large number of papers on sparse-grid sampling recovery and numerical applications. The reader can consult [4, 10] for surveys about results and bibliography. We also refer to the monographs [27, 28] for concepts and results on high dimensional problems and computation complexity.

Let us mention that high-dimensional functions having mixed smoothness play a fundamental role not only in approximation theory but also in mathematical physics and finance and other fields. For instance, in a recent work on regularity properties of solutions of the electronic Schrödinger equation, Yserentant [44] has shown that the eigenfunctions of the electronic Schrödinger operator have a certain mixed smoothness. Triebel [40, Chapter 6] has indicated a relation between Faber bases and sampling recovery in the context of spaces with mixed smoothness and solutions of Navier-Stokes equations. In particular, when initial data belonging to spaces with mixed smoothness, Navier-Stokes equations admit a unique solution. In mathematical finance, many problems are expressed as the expectation of some payoff function depending on quantities, such as stock prices, which are solutions of stochastic equations governed by Brownian motions. The payoff function normally has kinks and jumps and belongs to a very high dimensional space. To approximate the expected value one can apply preliminary integration method with respect to a single well chosen variable to obtain a function of $d-1$ variables which belongs to appropriate mixed Sobolev spaces in which Quasi-Monte Carlo can be applied efficiently, see [16] and references therein. For a survey on various aspects of high-dimensional approximation of functions having a mixed smoothness we refer the reader to the book [10].

The object of our interest in high-dimensional sampling recovery are functions on $\mathbb{I}^d$ having Hölder-Zygmund mixed smoothness $\alpha > 0$ when the dimension d may be very large. Let us first introduce the space $H_\infty^\alpha(\mathbb{I}^d)$ of all such functions. For a univariate function f on $\mathbb{I}$, the r th difference operator Δ_h^r is defined by

$$\Delta_h^r(f, x) := \sum_{j=0}^r (-1)^{r-j} \binom{r}{j} f(x + jh),$$

for all x and $h \geq 0$ such that $x, x + rh \in \mathbb{I}$. If u is any subset of $[d] := \{1, \dots, d\}$, for a multivariate function f on $\mathbb{I}^d$ the mixed (r, u) th difference operator $\Delta_{\mathbf{h}}^{r,u}$ is defined by

$$\Delta_{\mathbf{h}}^{r,u} := \prod_{i \in u} \Delta_{h_i}^r, \quad \Delta_{\mathbf{h}}^{r,\emptyset} = \text{Id},$$

where the univariate operator $\Delta_{h_i}^r$ is applied to the univariate function f by considering f as a function of variable x_i with the other variables held fixed, and Id is the identity operator. If $0 < \alpha \leq r$, the Hölder-Zygmund space $H_\infty^\alpha(\mathbb{I}^d)$ of mixed smoothness α is defined as the set of functions $f \in C(\mathbb{I}^d)$ for which the norm

$$\|f\|_{H_\infty^\alpha(\mathbb{I}^d)} := \max_{u \subset [d]} \left\{ \sup_{\mathbf{h}} \prod_{i \in u} h_i^{-\alpha} \|\Delta_{\mathbf{h}}^{r,u}(f)\|_{C(\mathbb{I}^d(r\mathbf{h},u))} \right\} \quad (1.1)$$

is finite, where $\mathbb{I}^d(r\mathbf{h}, u) := \{\mathbf{x} \in \mathbb{I}^d : x_i + rh_i \in \mathbb{I}, i \in u\}$. Note, that when $u = \emptyset$ the term in brackets of (1.1) is $\|f\|_{C(\mathbb{I}^d)}$. By the definition we have the inclusions $H_\infty^\alpha(\mathbb{I}^d) \subset C(\mathbb{I}^d) \subset L_p(\mathbb{I}^d)$ for $0 < p \leq \infty$. Moreover, it is well-known that for $\alpha = r$, the space $H_\infty^r(\mathbb{I})$ coincides with the Sobolev space $W_\infty^r(\mathbb{I})$. Hence, by a tensor product argument one can deduce the equality $H_\infty^\alpha(\mathbb{I}^d) = W_\infty^r(\mathbb{I}^d)$ where $W_\infty^r(\mathbb{I}^d)$ is the Sobolev space of functions having bounded mixed derivatives of order r .

Denote by $\mathring{U}_\infty^\alpha$ the set of all functions f in the unit ball of $H_\infty^\alpha(\mathbb{I}^d)$ such that $\|f\|_{H_\infty^\alpha(\mathbb{I}^d)} \leq 1$ and f vanishes on the boundary $\partial\mathbb{I}^d$ of $\mathbb{I}^d$, i.e., $f(\mathbf{x}) = 0$ if $x_j = 0$ or $x_j = 1$ for some index $j \in [d]$. Let $X_n = \{\mathbf{x}^j\}_{j=1}^n$ be a set of n sample points in $\mathbb{I}^d$ and $\Phi_n = \{\phi_j\}_{j=1}^n$ a family of functions on $\mathbb{I}^d$. We define the linear sampling algorithm $L(\Phi_n, X_n, \cdot)$ for an approximate recovery of functions $f \in \mathring{U}_\infty^\alpha$ from the sampling values $f(\mathbf{x}^1), \dots, f(\mathbf{x}^n)$ by

$$L(\Phi_n, X_n, f) := \sum_{j=1}^n f(\mathbf{x}^j) \phi_j. \quad (1.2)$$

The error of approximation will be measured in the isotropic Sobolev space $W_0^{1,p} := W_0^{1,p}(\mathbb{I}^d)$, $1 \leq p \leq \infty$, consisting of all functions $f \in L_p(\mathbb{I}^d)$ vanishing on the boundary $\partial\mathbb{I}^d$ of $\mathbb{I}^d$ in the sense of trace such that the “energy” norm

$$\|f\|_{W_0^{1,p}} := \begin{cases} \left(\sum_{i=1}^d \int_{\mathbb{I}^d} \left| \frac{\partial}{\partial x_i} f(\mathbf{x}) \right|^p d\mathbf{x} \right)^{1/p}, & 1 \leq p < \infty, \\ \max_{1 \leq i \leq d} \text{ess sup}_{\mathbf{x} \in \mathbb{I}^d} \left| \frac{\partial}{\partial x_i} f(\mathbf{x}) \right|, & p = \infty, \end{cases}$$

is finite (this is a norm due to Poincaré inequality). It is known that the norm of spaces $W_0^{1,p}$, especially, the energy norm of the space $H_0^1 := W_0^{1,2}$ is of great interest in approximation and numerical methods of PDEs. To characterize the error of best recovery algorithm from n sampling values of $f \in \mathring{U}_\infty^\alpha$, we consider the quantity

$$r_n := r_n(\mathring{U}_\infty^\alpha, W_0^{1,p}) := \inf_{X_n, \Phi_n} \sup_{f \in \mathring{U}_\infty^\alpha} \|f - L(\Phi_n, X_n, f)\|_{W_0^{1,p}},$$

where the infimum is taken over all sets of n sample points $X_n = \{\mathbf{x}^j\}_{j=1}^n$ and all the family $\Phi_n = \{\phi_j\}_{j=1}^n \subset W_0^{1,p}$. This quantity is called sampling n -width. For its properties and relations to other approximation quantities we refer the reader to [27, 10].

Our present investigation on sampling recovery can be considered as a continuation of the recent study in [11] where the explicit dimension-dependent estimates of the approximation error for linear algorithms of sampling recovery from sampling values on Smolyak sparse grids by using piecewise linear B-splines are obtained. In this paper, however, the optimal linear sampling algorithm is constructed on energy-based sparse grids which leads to the tight dimension-dependent approximation rate of optimal recovery for the class $\mathring{U}_\infty^\alpha$ with $1 < \alpha \leq 2$

$$C_1(p) B_*^{-d} n^{-(\alpha-1)} \leq r_n \leq C_2(\alpha, p) B^{*-d} n^{-(\alpha-1)},$$

where $B_* = B_*(p) > 1$ and $B^* = B^*(d, \alpha) > 0$. Moreover, we show that under a very light additional restriction the constant $B^* > 1$ whenever $d > d_0(\alpha, p)$. We can also explicitly construct an optimal linear sampling operator on a sparse grid of the form (1.2) by using truncated tensor product Faber series. This is one of the main result of our paper.

It has been shown that there is a close relation between the approximation by sampling recovery based on B-spline interpolation and quasi-interpolation representation, and the approximation by deep ReLU neural networks and [42, 43], [25], [12], [36], [34]. Neural networks have been studied and used for almost 70 years, dating back to the foundational work of Hebb [19] and of Rosenblatt [33]. They are composed of layers of computational nodes interconnected by activation functions. The theory of approximating functions using shallow networks goes back to [6, 20] where they showed that that any continuous functions can be approximated by shallow networks which have one hidden layer. These approximation results, however, do not provide any information on the required neurons of a network to achieve a given accuracy. In [24, 32] it has shown that a shallow network with $\mathcal{O}(\varepsilon^{-d/s})$ neurons can approximate functions of d -variables in differential space C^s with prescribed error ε . Hence, standard convergence results for shallow networks suffer from the curse of dimensionality when the dimension d may be very large.

Neural networks used recently in machine learning are distinguished from those popular in the 1980's and 1990's by emphasis on the depth of networks. In applications, it is observed that deep networks with many hidden layers appear to perform better than shallow ones of comparable size. In recent years, deep neural networks have been successfully applied to a striking variety of Machine Learning problems, including computer vision [21], natural language processing [41], speech recognition and image classification [22]. The main advantage of deep neural networks is that they can output compositions of functions cheaply. Since their application range is getting wider, theoretical analysis to reveal the reason why deep neural networks could lead to significant practical improvements attracts substantial attention [2, 23, 26, 37, 38]. In the last several years, there has been a number of interesting papers that address the role of depth and architecture of deep neural networks in approximating sets of functions which have a very special regularity properties such as analytic functions [14, 24], differentiable functions [31, 42], oscillatory functions [17], and functions in isotropic Sobolev spaces [18]. Most of them use deep ReLU neural networks for approximation since the rectified linear unit is a simple and preferable activation function in many applications. The output of such a network is a continuous piece-wise linear function which is easily and cheaply computed. In the context of function spaces with mixed smoothness, using decomposition by Faber series, it has been proven in [25] that for functions in Koborov space the number of parameters in a deep ReLU neural network needed to achieve an error tolerance of ε is $\mathcal{O}(\varepsilon^{-1/2} \log(\varepsilon^{-1})^d)$. In another study [36], the author has employed results in [7, 8] of the first author of the present paper on B-spline interpolation or quasi-interpolation representation to approximate functions in Besov spaces with mixed smoothness by deep ReLU neural networks. However, these results did not give explicit dimension-dependent estimates for number of weights of the neural network needed.

From our results on sparse-grid sampling recovery we deduce approximation by deep ReLU neural networks for d -variate functions in Hölder-Zygmund classes of mixed smoothness. Namely, we investigate the high-dimensional approximation by deep ReLU neural networks of functions from the classes U_∞^α . We focus our attention on estimation of the computation complexity of a deep ReLU neural network characterized by the number of its weights and its depth required to achieve a given accuracy for the approximation (cf. [1, 12, 42]), emphasizing the dimension dependence of the computation complexity of the network.

Let us briefly describe our main results on the approximation by deep ReLU neural networks. For

any $f \in \mathring{U}_\infty^\alpha$ we explicitly construct a deep ReLU neural networks Φ_f having the output $\mathcal{N}(\Phi_f, \cdot)$ that approximates f in the $W_0^{1,p}$ -norm with a prescribed accuracy ε and having computation complexity expressing the dimension-dependent number of weights

$$W(\Phi_f) \leq C_3(\alpha, p) B^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}),$$

and the dimension-dependent depth

$$L(\Phi_f) \leq C_4(\alpha, p) \log d \log(\varepsilon^{-1}),$$

where $B > 0$. If a further restriction is given, in particular when $1 \leq p \leq 2$, we show that $B > 1$ when $d > d_0(\alpha, p)$ and therefore overcoming the curse of dimensionality. We also prove that these bounds cannot improved up to logarithm terms. Moreover, the continuous piece-wise linear function $\mathcal{N}(\Phi_f, \cdot)$ can be designed also as an output $\mathcal{N}(\Phi_f^*, \cdot)$ of another “very” deep ReLU neural network Φ_f^* having computation complexity expressing the ε -independent width $N_w(\Phi_f^*) \leq C_5 d$ and the dimension-dependent depth

$$L(\Phi_f^*) \leq C_6(\alpha, p) B^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}).$$

The outline of this paper is as follows. Section 2 is devoted to recalling decomposition of continuous functions on the unit cube $\mathbb{I}^d$ by Faber system. There an estimate for its coefficients in the Sobolev spaces is given. Our results on sampling recovery for functions in Hölder-Zygmund classes are formulated in Section 3. In Section 4 we construct a deep ReLU neural network that approximates functions in $\mathring{U}_\infty^\alpha$ and prove upper and lower estimates for number of weights and the depth required. We conclude our main results in Section 5.

Notation. As usual, $\mathbb{N}$ denotes the natural numbers, $\mathbb{Z}$ denotes the integers, $\mathbb{R}$ the real numbers and $\mathbb{N}_0 := \{s \in \mathbb{Z} : s \geq 0\}$; $\mathbb{N}_{-1} = \mathbb{N}_0 \cup \{-1\}$. The letter d is always reserved for the underlying dimension of $\mathbb{R}^d$, $\mathbb{N}^d$, etc., and $[d]$ denotes the set of all natural numbers from 1 to d . Vectorial quantities are denoted by boldface letters and x_i denotes the i th coordinate of $\mathbf{x} \in \mathbb{R}^d$, i.e., $\mathbf{x} := (x_1, \dots, x_d)$. We use the notation $\mathbf{x}\mathbf{y}$ for the usual Euclidean inner product in $\mathbb{R}^d$ and $2^{\mathbf{x}} := (2^{x_1}, \dots, 2^{x_d})$. For $\mathbf{k}, \mathbf{s} \in \mathbb{N}_0^d$, we denote $2^{-\mathbf{k}}\mathbf{s} := (2^{-k_1}s_1, \dots, 2^{-k_d}s_d)$. For $\mathbf{x} \in \mathbb{R}^d$ we write $|\mathbf{x}|_0 = |\{x_j \neq 0, j = 1, \dots, d\}|$ and if $0 < p \leq \infty$ we denote $|\mathbf{x}|_p := (\sum_{i=1}^d |x_i|^p)^{1/p}$ with the usual modification when $p = \infty$. The notations $|\cdot|_0$ and $|\cdot|_p$ are extended to matrices $\mathbf{W} \in \mathbb{R}^{m \times n}$. For the function $f(p)$ of variable p , $f(\infty)$ is understood as $f(\infty) = \lim_{p \rightarrow \infty} f(p)$ when the limit exists.

2 Faber series

In this section we recall a decomposition of continuous functions on the unit cube $\mathbb{I}^d$ into tensor product Faber series and give an estimate in the $W_0^{1,p}$ -norm of the components of functions from the Hölder-Zygmund space of mixed smoothness $H_\infty^\alpha(\mathbb{I}^d)$. This decomposition plays a fundamental role in construction of linear algorithms of sparse-grid sampling recovery and of deep neural networks for approximation in the $W_0^{1,p}$ -norm of functions from the Hölder-Zygmund space of mixed smoothness $H_\infty^\alpha(\mathbb{I}^d)$.

We start by introducing the tensorized Faber basis. Let $\varphi(x) = (1 - |x - 1|)_+$, $x \in \mathbb{R}$, be the hat function (the piece-wise linear B-spline with knots at 0, 1, 2), where $x_+ := \max(x, 0)$ for $x \in \mathbb{R}$. For $k \in \mathbb{N}_{-1}$ we define the functions $\varphi_{k,s}$ on $\mathbb{I}$ by

$$\varphi_{k,s}(x) := \varphi(2^{k+1}x - 2s), \quad x \in \mathbb{I}, \quad k \geq 0, \quad s \in Z(k) := \{0, 1, \dots, 2^k - 1\}, \quad (2.1)$$

and

$$\varphi_{-1,s}(x) := \varphi(x - s + 1), \quad x \in \mathbb{I}, \quad s \in Z(-1) := \{0, 1\}. \quad (2.2)$$

Put $Z^d(\mathbf{k}) := \times_{i=1}^d Z(k_i)$. For $\mathbf{k} \in \mathbb{N}_{-1}^d$, $\mathbf{s} \in Z^d(\mathbf{k})$, define the d -variate tensor product hat functions

$$\varphi_{\mathbf{k},\mathbf{s}}(\mathbf{x}) := \prod_{i=1}^d \varphi_{k_i,s_i}(x_i), \quad \mathbf{x} \in \mathbb{I}^d. \quad (2.3)$$

For a univariate function f on $\mathbb{I}$, $k \in \mathbb{N}_{-1}$, and $s \in Z(k)$ we define

$$\lambda_{k,s}(f) := -\frac{1}{2} \Delta_{2^{-k-1}}^2 f(2^{-k}s), \quad k \geq 0, \quad \lambda_{-1,s}(f) := f(s).$$

We also define the linear functionals $\lambda_{\mathbf{k},\mathbf{s}}$ for multivariate function f on $\mathbb{I}^d$, $\mathbf{k} \in \mathbb{N}_{-1}^d$, and $\mathbf{s} \in Z^d(\mathbf{k})$ by

$$\lambda_{\mathbf{k},\mathbf{s}}(f) := \prod_{i=1}^d \lambda_{k_i,s_i}(f),$$

where the univariate functional λ_{k_i,s_i} is applied to the univariate function f by considering f as a function of variable x_i with the other variables held fixed.

We have the following decomposition.

Lemma 2.1 *The Faber system $\{\varphi_{\mathbf{k},\mathbf{s}} : \mathbf{k} \in \mathbb{N}_{-1}^d, \mathbf{s} \in Z^d(\mathbf{k})\}$ is a basis in $C(\mathbb{I}^d)$. Moreover, any function $f \in C(\mathbb{I}^d)$ can be represented by the Faber series*

$$f = \sum_{\mathbf{k} \in \mathbb{N}_{-1}^d} q_{\mathbf{k}}(f) := \sum_{\mathbf{k} \in \mathbb{N}_{-1}^d} \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k},\mathbf{s}}(f) \varphi_{\mathbf{k},\mathbf{s}}, \quad (2.4)$$

converging in the norm of $C(\mathbb{I}^d)$.

The system (2.1)-(2.2) and above decomposition for continuous functions on the interval $\mathbb{I}$ goes back to Faber [15]. The extension for tensorized Faber bases in higher dimensions was obtained in [39, Theorem 3.10] for $d = 2$ and in [7, 9] for the case $d \geq 2$. In [39, 7], the authors also established the decomposition (2.4) with equivalent discrete quasi-norm for function spaces of mixed smoothness. A generalization to B-spline interpolation and quasi-interpolation representation was obtained by the first author of this paper in [7, 9]. Note that the sequence in (2.4) converges conditionally in $C(\mathbb{I}^d)$ if $f \in C(\mathbb{I}^d)$, see comment in [39, Section 3.3.2]. Moreover, when $f \in \dot{U}_{\infty}^{\alpha}$ we can write

$$f = \sum_{\mathbf{k} \in \mathbb{N}_0^d} q_{\mathbf{k}}(f)$$

with unconditional convergence in $C(\mathbb{I}^d)$, see [39, Theorem 3.13]. In this case we have

$$\lambda_{\mathbf{k},\mathbf{s}}(f) := \prod_{i=1}^d \left(-\frac{1}{2} \Delta_{2^{-k_i-1}}^2 f(2^{-k_i}s_i) \right)$$

and the following estimate holds

$$|\lambda_{\mathbf{k},\mathbf{s}}(f)| \leq 2^{-(\alpha+1)d} 2^{-\alpha|\mathbf{k}|_1}, \quad \mathbf{k} \in \mathbb{N}_0^d, \quad \mathbf{s} \in Z^d(\mathbf{k}). \quad (2.5)$$

We now estimate the $W_0^{1,p}$ -norm of $q_{\mathbf{k}}(f)$.

Lemma 2.2 Let $d \in \mathbb{N}$, $1 \leq \alpha \leq 2$, and $1 \leq p \leq \infty$. Then for a function $f \in \dot{U}_\infty^\alpha$ and $\mathbf{k} \in \mathbb{N}_0^d$ we have

$$\|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} \leq \frac{2^{-\alpha|\mathbf{k}|_1+1}}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d}} |2^{\mathbf{k}}|_p.$$

Proof. Let us prove the lemma for the case $1 \leq p < \infty$. The case $p = \infty$ can be proven similarly with a slight modification. For $\mathbf{k} \in \mathbb{N}_0^d$, by disjoint supports of $\varphi_{\mathbf{k},\mathbf{s}}$, $\mathbf{s} \in Z^d(\mathbf{k})$, we have

$$\begin{aligned} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}}^p &= \sum_{i=1}^d \int_{\mathbb{I}^d} \left| \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k},\mathbf{s}}(f) \frac{\partial}{\partial x_i} \varphi_{\mathbf{k},\mathbf{s}}(\mathbf{x}) \right|^p d\mathbf{x} \\ &\leq \sup_{\mathbf{s} \in Z^d(\mathbf{k})} |\lambda_{\mathbf{k},\mathbf{s}}(f)|^p \sum_{i=1}^d |Z^d(\mathbf{k})| \left(\prod_{j \neq i} 2 \int_0^{2^{-k_j-1}} |2^{k_j+1} x_j|^p dx_j \right) \left(2 \int_0^{2^{-k_i-1}} 2^{pk_i+p} dx_i \right) \\ &\leq \left(\frac{2^{-\alpha|\mathbf{k}|_1}}{2^{(\alpha+1)d}} \right)^p \sum_{i=1}^d 2^{|\mathbf{k}|_1} \left(\prod_{j \neq i} \frac{2}{p+1} 2^{-k_j-1} \right) 2 \cdot 2^{(k_i+1)(p-1)} \\ &= \left(\frac{2^{-\alpha|\mathbf{k}|_1}}{2^{(\alpha+1)d}} \right)^p \frac{1}{(p+1)^{d-1}} \sum_{i=1}^d 2^{p(k_i+1)}. \end{aligned}$$

This proves the claim. $\square$

3 Sparse-grid sampling recovery

In this section we consider the approximation by explicitly constructed linear sampling operator on sparse grids with the error measured in the norm of the space $W_0^{1,p}$. The optimality is investigated in terms of sampling n -widths. We prove some tight dimension-dependent error estimates of the sampling n -widths explicit in d and n .

We start with constructing sparse grids and sampling operators on them for approximately recovering functions in $\dot{U}_\infty^\alpha$ from their values on these grids. For $\beta \geq 1$ and $m \in \mathbb{N}$, we define the sets of multi-indices

$$\Delta_\beta^d(m) := \{\mathbf{k} \in \mathbb{N}_0^d : |\mathbf{k}|_1 = m - j, |\mathbf{k}|_\infty \geq m - \lfloor \beta j \rfloor \text{ for } j = 0, \dots, m\},$$

and

$$D_\beta^d(m) := \{(\mathbf{k}, \mathbf{s}) : \mathbf{k} \in \Delta_\beta^d(m), \mathbf{s} \in Z^d(\mathbf{k})\}.$$

The definition of $\Delta_\beta^d(m)$ is similar to $\Delta(\alpha, \beta; \xi) = \{\mathbf{k} \in \mathbb{N}_0^d : \alpha|\mathbf{k}|_1 - \beta|\mathbf{k}|_\infty \leq \xi\}$ introduced in [5] but simpler. We also put

$$\Delta^d(m) := \{\mathbf{k} \in \mathbb{N}_0^d : |\mathbf{k}|_1 \leq m\}.$$

It is obvious that $\Delta_\beta^d(m)$ is a subset of $\Delta^d(m)$ for all $\beta \geq 1$.

Consider the operator

$$R_\beta(m, f) := \sum_{\mathbf{k} \in \Delta_\beta^d(m)} q_{\mathbf{k}}(f) = \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k},\mathbf{s}}(f) \varphi_{\mathbf{k},\mathbf{s}}, \quad (3.1)$$

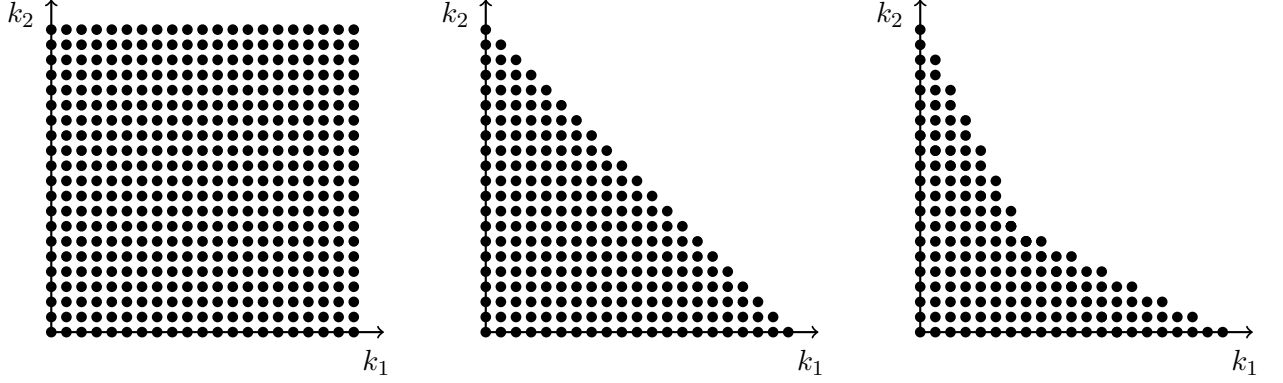


Figure 1: Illustration of different sets of multi-indices in $\mathbb{N}_0^2$. The left graph is $\{\mathbf{k} \in \mathbb{N}_0^2 : |\mathbf{k}|_\infty \leq 20\}$, the middle is $\Delta^2(20)$ and the right is $\Delta_2^2(20)$.

and the energy-norm-based grid

$$G_\beta^d(m) := \{2^{-\mathbf{k}} \mathbf{s} : \mathbf{k} \in \Delta_\beta^d(m), \mathbf{s} \in Z^d(\mathbf{k} + \mathbf{1})\}, \quad \mathbf{1} := (1, \dots, 1) \in \mathbb{R}^d.$$

It is easy to see that

$$G_\beta^d(m) = \{2^{-\mathbf{k}} \mathbf{s} : \mathbf{k} \in \bar{\Delta}_\beta^d(m), \mathbf{s} \in Z^d(\mathbf{k} + \mathbf{1})\},$$

where

$$\bar{\Delta}_\beta^d(m) := \{\mathbf{k} \in \Delta_\beta^d(m) : \mathbf{k} + \mathbf{e}^j \notin \Delta_\beta^d(m) \text{ for all } j \in [d]\},$$

and $\{\mathbf{e}^j\}_{j \in [d]}$ is the standard basis in $\mathbb{R}^d$.

We notice some important properties of the operator $R_\beta(m, \cdot)$ and the grid $G_\beta^d(m)$. The function $R_\beta(m, f)$ is completely determined by the values of f on the grid $G_\beta^d(m)$. Moreover, $R_\beta(m, f)$ interpolates f at the points of $G_\beta^d(m)$. Thus, $R_\beta(m, \cdot)$ is an interpolation sampling operator on the grid $G_\beta^d(m)$. As shown in what follows, with an appropriate choice of parameter β , the function $R_\beta(m, f)$ is suitable to approximately recovering the function f in $\dot{U}_\infty^\alpha$ from the sample values on the grid $G_\beta^d(m)$.

The following lemma gives an upper estimate of the cardinality of $D_\beta^d(m)$ and $G_\beta^d(m)$ showing their sparsity.

Lemma 3.1 *Let $d \in \mathbb{N}$. We have for every $\beta > 1$ and $m \in \mathbb{N}$,*

$$|D_\beta^d(m)| \leq \frac{\beta}{\beta - 1} d \left(1 - 2^{-\frac{1}{\beta-1}}\right)^{-d} 2^m,$$

and hence,

$$|G_\beta^d(m)| \leq \frac{\beta}{\beta - 1} d 2^d \left(1 - 2^{-\frac{1}{\beta-1}}\right)^{-d} 2^m.$$

Proof. Since $|G_\beta^d(m, f)| \leq |D_\beta^d(m + d)|$, it is sufficient to prove the first estimate in the lemma. We have

$$|D_\beta^d(m)| = \sum_{\mathbf{k} \in \Delta_\beta^d(m)} 2^{|\mathbf{k}|_1} = 2^m \sum_{j=0}^m 2^{-j} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty \geq m - \lfloor \beta j \rfloor}} 1.$$

Note, that for $i \in \mathbb{N}_0$ and $\ell \in [d]$ there are $\binom{d-2+\lfloor(\beta-1)j\rfloor-i}{d-2}$ multi-indices $\mathbf{k} \in \mathbb{N}_0^d$ satisfying

$$k_\ell = m - \lfloor \beta j \rfloor + i, \quad \text{and} \quad \sum_{r \neq \ell} k_r = \lfloor (\beta - 1)j \rfloor - i.$$

From this we can estimate

$$\sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty \geq m - \lfloor \beta j \rfloor}} 1 \leq d \sum_{i=0}^{\lfloor (\beta-1)j \rfloor} \binom{d-2+\lfloor(\beta-1)j\rfloor-i}{d-2} = d \binom{d-1+\lfloor(\beta-1)j\rfloor}{d-1}.$$

Hence,

$$|D_\beta^d(m)| \leq d 2^m \sum_{j=0}^m 2^{-j} \binom{d-1+\lfloor(\beta-1)j\rfloor}{d-1}.$$

Putting $\lfloor (\beta - 1)j \rfloor = k$ we obtain $j \geq \frac{k}{\beta-1}$ and

$$|\{j \in \mathbb{N}_0 : \lfloor (\beta - 1)j \rfloor = k\}| < \frac{1}{\beta-1} + 1$$

which leads to

$$|D_\beta^d(m)| \leq d 2^m \left(\frac{1}{\beta-1} + 1 \right) \sum_{j=0}^{\infty} 2^{-\frac{1}{\beta-1}j} \binom{d-1+j}{d-1} = d 2^m \frac{\beta}{\beta-1} (1 - 2^{-\frac{1}{\beta-1}})^{-d}.$$

The last equality is due to

$$\sum_{j=0}^{\infty} x^j \binom{k+j}{k} = (1-x)^{-k-1} \quad (3.2)$$

for $k \in \mathbb{N}_0$ and $x \in (0, 1)$ which is obtained by taking k th derivative both sides of $(1-x)^{-1} = \sum_{j=0}^{\infty} x^j$. The proof is completed. $\square$

We continue by proving a supplementary result.

Lemma 3.2 *Let $d \in \mathbb{N}$, $\ell \in \mathbb{N}_0$, and $1 \leq p \leq \infty$. Then it holds*

$$\sum_{\mathbf{k} \in \mathbb{N}_0^d, |\mathbf{k}|_1 = \ell} |2^{\mathbf{k}}|_p \leq d 2^{\ell+d-1}. \quad (3.3)$$

Proof. By monotonicity in p of ℓ_p norms, it is enough to prove the lemma for $p = 1$. We use induction argument with respect to d . It is obvious that the inequality holds for $d = 1$ and $\ell \in \mathbb{N}_0$ or $d \in \mathbb{N}$ and $\ell = 0$. Assume that

$$\sum_{\mathbf{k} \in \mathbb{N}_0^d, |\mathbf{k}|_1 = j} \sum_{i=1}^d 2^{k_i} \leq d 2^{j+d-1}$$

for $j = 0, \dots, \ell$. We show that the inequality (3.3) holds for $d+1$ instead of d . Indeed, we have

$$\begin{aligned} \sum_{\mathbf{k} \in \mathbb{N}_0^{d+1}, |\mathbf{k}|_1 = \ell} \sum_{i=1}^{d+1} 2^{k_i} &= \sum_{j=0}^{\ell} \sum_{k_1+\dots+k_d=\ell-j} 2^j + \sum_{i=1}^d 2^{k_i} \\ &\leq \sum_{j=0}^{\ell} d 2^{\ell-j+d-1} + \sum_{j=0}^{\ell} 2^j \binom{\ell-j+d-1}{d-1} \\ &= d 2^{d-1} (2^{\ell+1} - 1) + 2^{\ell} \sum_{j=0}^{\ell} 2^{j-\ell} \binom{\ell-j+d-1}{d-1}. \end{aligned}$$

Using (3.2) we finally obtain

$$\sum_{\mathbf{k} \in \mathbb{N}_0^{d+1}, |\mathbf{k}|_1 = \ell} \sum_{i=1}^{d+1} 2^{k_i} \leq d 2^{d-1} (2^{\ell+1} - 1) + 2^{\ell+d} \leq (d+1) 2^{\ell+d}.$$

The proof is completed. $\square$

We give a preliminary dimension-dependent error estimate in terms of parameter m of the approximation of a $f \in \dot{U}_\infty^\alpha$ by the sampling operator $R_\beta(m, \cdot)$.

Theorem 3.1 *Let $d \geq 2$, $1 < \alpha \leq 2$, $\beta > \alpha$, and $1 \leq p \leq \infty$. Then for every $f \in \dot{U}_\infty^\alpha$ we have*

$$\|f - R_\beta(m, f)\|_{W_0^{1,p}} \leq K_1 \frac{d^2 2^{-m(\alpha-1)}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d}, \quad (3.4)$$

where $K_1 = K_1(\alpha, \beta, p) = 2(p+1)^{\frac{1}{p}} \max \left\{ \frac{2\beta}{\beta-1}, \frac{1}{2^{\alpha-1}-1} \right\}$.

Proof. Let us prove the theorem for the case $1 \leq p < \infty$. The case $p = \infty$ can be proven similarly with a slight modification. Recall that $\Delta_\beta^d(m)$ with $\beta > 1$ is a subset of $\Delta^d(m)$. Hence, for every $f \in \dot{U}_\infty^\alpha$ we have

$$\begin{aligned} \|f - R_\beta(m, f)\|_{W_0^{1,p}} &\leq \sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \Delta_\beta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} \\ &= \sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \Delta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} + \sum_{\mathbf{k} \in \Delta^d(m) \setminus \Delta_\beta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}}. \end{aligned} \quad (3.5)$$

For the first term on the right side, from Lemma 2.2 we have

$$\begin{aligned} \sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \Delta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} &\leq \frac{1}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d}} \sum_{\mathbf{k} \in \mathbb{N}_0^d, |\mathbf{k}|_1 > m} 2^{-\alpha|\mathbf{k}|_1+1} |2^{\mathbf{k}}|_p \\ &= \frac{2}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d}} \sum_{\ell=m+1}^{\infty} 2^{-\alpha\ell} \sum_{\mathbf{k} \in \mathbb{N}_0^d, |\mathbf{k}|_1 = \ell} |2^{\mathbf{k}}|_p. \end{aligned}$$

In view of Lemma 3.2 we get

$$\sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \Delta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} \leq \frac{d}{(p+1)^{\frac{d-1}{p}} 2^{\alpha d}} \sum_{\ell=m+1}^{\infty} 2^{-(\alpha-1)\ell} = \frac{d 2^{-(\alpha-1)m}}{(p+1)^{\frac{d-1}{p}} 2^{\alpha d} (2^{\alpha-1} - 1)}.$$

We now consider the second term on the right side of (3.5). Denote j^* the maximum value of j such that the set

$$\{\mathbf{k} \in \mathbb{N}_0^d : |\mathbf{k}|_1 = m - j, |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor\}$$

is not empty. Following the argument in the proof of [4, Theorem 3.10] and using Lemma 2.2 we get

$$\begin{aligned}
\sum_{\mathbf{k} \in \Delta^d(m) \setminus \Delta_\beta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} &\leq \sum_{j=0}^{j^*} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} \\
&\leq \frac{2}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d}} \sum_{j=0}^{j^*} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} 2^{-\alpha|\mathbf{k}|_1} |2^{\mathbf{k}}|_p \\
&= \frac{2^{-\alpha m+1}}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d}} \sum_{j=0}^{j^*} 2^{\alpha j} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} |2^{\mathbf{k}}|_p.
\end{aligned} \tag{3.6}$$

The sum on the right side can be estimated as

$$\begin{aligned}
\sum_{j=0}^{j^*} 2^{\alpha j} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} |2^{\mathbf{k}}|_p &\leq \sum_{j=0}^{j^*} 2^{\alpha j} \sum_{i=1}^{m-1-\lfloor \beta j \rfloor} d \binom{m+d-2-i-j}{d-2} 2^i \\
&= d 2^{m-1} \sum_{j=0}^{j^*} 2^{-\lfloor (\beta-\alpha)j \rfloor} \sum_{\ell=0}^{m-2-\lfloor \beta j \rfloor} \binom{d-1+\lfloor (\beta-1)j \rfloor + \ell}{d-2} 2^{-\ell},
\end{aligned}$$

where in the equality we put $i = m-1-\lfloor \beta j \rfloor - \ell$. As in the proof of Lemma 3.1 we put $\lfloor (\beta-1)j \rfloor = k$ and get $\lfloor (\beta-\alpha)j \rfloor \geq \frac{k(\beta-\alpha)}{\beta-1} - 1$. Since $\frac{\beta-\alpha}{\beta-1} < 1$, we have

$$\begin{aligned}
\sum_{j=0}^{j^*} 2^{\alpha j} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} |2^{\mathbf{k}}|_p &\leq \frac{d 2^m \beta}{\beta-1} \sum_{k=0}^{\infty} \sum_{\ell=0}^{\infty} \binom{d-1+k+\ell}{d-2} 2^{-\ell} 2^{-\frac{\beta-\alpha}{\beta-1}k} \\
&\leq \frac{d 2^m \beta}{\beta-1} \sum_{k=0}^{\infty} \sum_{\ell=0}^{\infty} \binom{d-1+k+\ell}{d-2} 2^{-\frac{\beta-\alpha}{\beta-1}(k+\ell)} \\
&= \frac{d 2^m \beta}{\beta-1} \sum_{j=0}^{\infty} \binom{d-1+j}{d-2} 2^{-\frac{\beta-\alpha}{\beta-1}j} (j+1) \\
&= (d-1) \frac{d 2^m \beta}{\beta-1} \sum_{j=0}^{\infty} \binom{d-1+j}{d-1} 2^{-\frac{\beta-\alpha}{\beta-1}j}.
\end{aligned}$$

From the assumption $\beta > \alpha > 1$ and (3.2) we arrive at

$$\sum_{j=0}^{j^*} 2^{\alpha j} \sum_{\substack{|\mathbf{k}|_1 = m-j \\ |\mathbf{k}|_\infty < m - \lfloor \beta j \rfloor}} |2^{\mathbf{k}}|_p \leq \frac{d^2 2^m \beta}{\beta-1} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^{-d}.$$

Inserting this into (3.6) we get

$$\sum_{\mathbf{k} \in \Delta^d(m) \setminus \Delta_\beta^d(m)} \|q_{\mathbf{k}}(f)\|_{W_0^{1,p}} \leq \frac{2\beta}{\beta-1} \frac{d^2 2^{-m(\alpha-1)}}{(p+1)^{\frac{d-1}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d}.$$

Since $2^{-d}(1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^{-d} > 1$ which is due to $\alpha > 1$, we finally obtain the desired estimate. $\square$

Remark 3.2 For approximation of functions from $\dot{U}_\infty^\alpha$, we could take the sampling operator $R_\square(m, f) := \sum_{|\mathbf{k}|_\infty \leq m} q_{\mathbf{k}}(f)$ on the corresponding traditional standard grid

$$G_\square^d(m) := \{2^{-\mathbf{k}} \mathbf{s} : |\mathbf{k}|_\infty = m+1, \mathbf{s} \in Z^d(\mathbf{k})\},$$

and the sampling operator $R_\Delta(m, f) := \sum_{|\mathbf{k}|_1 \leq m} q_{\mathbf{k}}(f)$ on the corresponding Smolyak grid

$$G_\Delta^d(m) := \{2^{-\mathbf{k}} \mathbf{s} : |\mathbf{k}|_1 = m, \mathbf{s} \in Z^d(\mathbf{k} + \mathbf{1})\}.$$

It is easy to see that the error of approximation in the $W_0^{1,p}$ norm of $f \in \dot{U}_\infty^\alpha$ by $R_\square(m, f)$ or $R_\Delta(m, f)$ is the same as the error of approximation in the $W_0^{1,p}$ norm of $f \in \dot{U}_\infty^\alpha$ by $R_\beta(m, f)$. On the other hand, we notice that the sparsity of the grid $G_\beta^d(m)$ in the operator $R_\beta(m, f)$, is much higher than the sparsity of the grids $G_\square^d(m)$ and $G_\Delta^d(m)$, see the estimate of $|G_\beta^d(m)|$ in Lemma 3.1 in comparing with $|G_\square^d(m)| \approx 2^{dm}$ and $|G_\Delta^d(m)| \approx 2^{m+d} \binom{m+d-1}{d-1}$.

Some results similar to (3.4) were obtained in [4, Theorem 3.8] for the approximation in energy norm of functions with 2nd mixed derivatives bounded in L_2 - or L_∞ -norm.

In the following theorem, from Theorem 3.1 we deduce a dimension-dependent estimate of the computation complexity characterized by the number of sample values in the operator $R_\beta(m, \cdot)$ necessary for approximating functions from $\dot{U}_\infty^\alpha$ with an accuracy ε . This estimate will be used as an auxiliary result in the next section on approximation of functions from $\dot{U}_\infty^\alpha$ by deep ReLU neural networks.

Theorem 3.3 Let $d \geq 2$, $1 < \alpha \leq 2$, $\beta > \alpha$, and $1 \leq p \leq \infty$. Then there is $\varepsilon_0 \in (0, 1]$ such that for every $\varepsilon \in (0, \varepsilon_0)$ and every $f \in \dot{U}_\infty^\alpha$ we have

$$\|f - R_\beta(m, f)\|_{W_0^{1,p}} \leq \varepsilon$$

and

$$|D_\beta^d(m)| \leq K_2 B_1^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}}, \quad |G_\beta^d(m)| \leq K_2 (B_1/2)^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}},$$

with

$$m := \left\lceil \frac{1}{\alpha-1} \log \left(\frac{K_1 d^2 \varepsilon^{-1}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right) \right\rceil \quad (3.7)$$

and

$$B_1 = B_1(d, \alpha, \beta, p) := (1 - 2^{-\frac{1}{\beta-1}}) \left(\frac{(p+1)^{\frac{1}{p}} 2^{(\alpha+1)} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})}{d^{\frac{\alpha+1}{d}}} \right)^{\frac{1}{\alpha-1}}, \quad (3.8)$$

where K_1 is the constant in Theorem 3.1 and $K_2 = K_2(\alpha, \beta, p) := \frac{2\beta}{\beta-1} K_1^{\frac{1}{\alpha-1}}$. Moreover, if α and p satisfy

$$2^{\frac{1}{\alpha}} - (p+1)^{-\frac{1}{p\alpha}}/2 > 1 \quad (3.9)$$

and

$$\frac{\alpha + \log \left(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}} \right)}{1 + \log \left(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}} \right)} < \beta < 1 - \frac{1}{\log \left(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}} \right)} \quad (3.10)$$

then there exist constants $d(\alpha, \beta, p) \in \mathbb{N}$ and $B(\alpha, \beta, p) > 1$ such that $B_1 \geq B(\alpha, \beta, p) > 1$ for all $d \geq d(\alpha, \beta, p)$.

Proof. Choose

$$\varepsilon_0 = \min \left\{ 1, \frac{K_1 d^2}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right\}.$$

For $\varepsilon \in (0, \varepsilon_0)$ we take m as in (3.7). This implies that

$$m - 1 < \frac{1}{\alpha - 1} \log \left(\frac{K_1 d^2 \varepsilon^{-1}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right) \leq m.$$

Hence from Theorem 3.1 and Lemma 3.1 we obtain

$$\|f - R_\beta(m, f)\|_{W_0^{1,p}} \leq \varepsilon$$

and

$$\begin{aligned} |D_\beta^d(m)| &\leq d 2^m \frac{\beta}{\beta-1} (1 - 2^{-\frac{1}{\beta-1}})^{-d} \\ &\leq 2d \frac{\beta}{\beta-1} (1 - 2^{-\frac{1}{\beta-1}})^{-d} \left(\frac{K_1 d^2 \varepsilon^{-1}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right)^{\frac{1}{\alpha-1}} \\ &\leq \frac{2\beta}{\beta-1} K_1^{\frac{1}{\alpha-1}} (1 - 2^{-\frac{1}{\beta-1}})^{-d} \left(\frac{d^{(\alpha+1)/d}}{(p+1)^{\frac{1}{p}} 2^{(\alpha+1)} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})} \right)^{\frac{d}{\alpha-1}} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \\ &= K_2 B_1^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \end{aligned} \tag{3.11}$$

which is the first statement. To prove the second one, we show that under the conditions (3.9) and (3.10) it holds

$$(1 - 2^{-\frac{1}{\beta-1}})^{(\alpha-1)} (p+1)^{\frac{1}{p}} 2^{\alpha+1} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}}) > 1.$$

Indeed, the last inequality is equivalent to

$$\frac{1}{(p+1)^{\frac{1}{p}} 2^{\alpha+1}} < (1 - 2^{-\frac{1}{\beta-1}})^{(\alpha-1)} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}}). \tag{3.12}$$

If $\beta - \alpha \geq 1$ then $(1 - 2^{-\frac{1}{\beta-1}}) \leq (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})$. Hence, the last inequality is fulfilled if

$$\frac{1}{(p+1)^{\frac{1}{p}} 2^{\alpha+1}} < (1 - 2^{-\frac{1}{\beta-1}})^\alpha \iff \beta < 1 - \frac{1}{\log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})}.$$

If $\beta - \alpha < 1$ then $(1 - 2^{-\frac{1}{\beta-1}}) > (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})$. Hence the inequality (3.12) is fulfilled if

$$\frac{1}{(p+1)^{\frac{1}{p}} 2^{\alpha+1}} < (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^\alpha \iff \beta > \frac{\alpha + \log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})}{1 + \log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})}.$$

Assigning

$$1 - \frac{1}{\log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})} > \frac{\alpha + \log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})}{1 + \log(1 - (p+1)^{-\frac{1}{p\alpha}} 2^{-1-\frac{1}{\alpha}})}$$

we find $2^{\frac{1}{\alpha}} - (p+1)^{-\frac{1}{p\alpha}}/2 > 1$. Since $d^{\frac{\alpha+1}{d}}$ tends to 1 when $d \rightarrow \infty$, there are $d(\alpha, \beta, p) \in \mathbb{N}$ and $B(\alpha, \beta, p) > 1$ such that $B_1 \geq B(\alpha, \beta, p) > 1$ for all $d \geq d(\alpha, \beta, p)$. The proof is completed. $\square$

Remark 3.4 For $1 \leq p \leq 2$ the condition (3.9) is satisfied for all $\alpha \in (1, 2]$, and, therefore, we can always find $\beta = \beta(\alpha) > \alpha$, $d(\alpha) \in \mathbb{N}$, and $B(\alpha) > 1$ such that $B_1 \geq B(\alpha) > 1$ for all $d \geq d(\alpha)$. In this case, the approximation error is going to 0 as fast as the exponent B_1^{-d} when d tending to ∞ .

We show the optimality in terms of sampling n -widths of the sampling operator $R_\beta(m, \cdot)$ on the grid $G_\beta^d(m)$ with the number of sample points $|G_\beta^d(m)| \leq n$, by giving dimension-dependent tight upper and lower bounds of the sampling n -widths r_n .

Theorem 3.5 Let $d \geq 2$, $1 < \alpha \leq 2$, $\beta > \alpha$, and $1 \leq p \leq \infty$. Then there is $n_0 \in \mathbb{N}$ such that for $n \geq n_0$ we have

$$\frac{3}{M4^{\alpha-1}} \left(\frac{3}{2p+2} \right)^{1/p} B_*^{-d} n^{-(\alpha-1)} \leq r_n \leq K_2^{\alpha-1} B^{*-d} n^{-(\alpha-1)} \quad (3.13)$$

where $B_* := 18 \cdot 3^{1/p} M^{-1}$, $M = \|M_3\|_{L_p(\mathbb{R})}$, $B^* := (B_1/2)^{\alpha-1}$, K_2 and B_1 are given in Theorem 3.3. The lower bound holds for $n \geq 2$. Moreover, we can explicitly define $m(n) \in \mathbb{N}$ so that $|G_\beta^d(m(n))| \leq n$ and for the sampling operator $R_\beta(m(n), \cdot)$ on the grid $G_\beta^d(m(n))$, we have that

$$r_n \leq \sup_{f \in \dot{U}_\infty^\alpha} \|f - R_\beta(m(n), f)\|_{W_0^{1,p}} \leq K_2^{\alpha-1} B^{*-d} n^{-(\alpha-1)}.$$

Proof.

Upper bound. We take n_0 the smallest positive integer such that $(K_2(B_1/2)^{-d})^{\alpha-1} n_0^{-(\alpha-1)} < \varepsilon_0$. For $n \geq n_0$, we define

$$\varepsilon = (K_2(B_1/2)^{-d})^{\alpha-1} n^{-(\alpha-1)}.$$

By Theorem 3.3 there is $m(n) \in \mathbb{N}$ such that

$$\|f - R_\beta(m(n), f)\|_{W_0^{1,p}} \leq (K_2(B_1/2)^{-d})^{\alpha-1} n^{-(\alpha-1)},$$

and $|G_\beta^d(m(n))| \leq n$. If we define the set of points $X_n = G_\beta^d(m(n))$ then by the definition of $R_\beta(m(n), f)$ we obtain

$$r_n \leq \sup_{f \in \dot{U}_\infty^\alpha} \|f - R_\beta(m(n), f)\|_{W_0^{1,p}} \leq (K_2(B_1/2)^{-d})^{\alpha-1} n^{-(\alpha-1)} = K_2^{\alpha-1} B^{*-d} n^{-(\alpha-1)}.$$

Lower bound. Let us prove the lower bound for the case $1 \leq p < \infty$. The case $p = \infty$ can be proven similarly with a slight modification. Our strategy to prove the lower bound is as follows. Based on the obvious inequality

$$r_n \geq \inf_{X_n \subset \mathbb{I}^d} \sup_{f \in \dot{U}_\infty^\alpha, f(\mathbf{x}^j)=0, \mathbf{x}^j \in X_n} \|f\|_{W_0^{1,p}}, \quad (3.14)$$

where the infimum is taken over all sets X_n of n points. For any point set $X_n = \{\mathbf{x}^j\}_{j=1}^n$ in $\mathbb{I}^d$, we will construct a test function $f \in \dot{U}_\infty^\alpha$ vanishing at the points $\mathbf{x}^j$, $j = 1, \dots, n$, and prove that the norm $\|f\|_{W_0^{1,p}}$ is bounded from below by the quantity in the left side of (3.13).

Let M_3 be the B-spline with knots at the points 0, 1, 2, 3, i.e.,

$$M_3(x) = \frac{1}{2} \begin{cases} x^2 & \text{if } 0 \leq x < 1 \\ -2x^2 + 6x - 3 & \text{if } 1 \leq x < 2 \\ (3-x)^2 & \text{if } 2 \leq x < 3 \end{cases}$$

and $M_3(x) = 0$ otherwise. Let $\psi(x) := M_3(3x)$. We define the univariate non-negative functions $\psi_{k,s}$ by

$$\psi_{k,s}(x) := \psi(2^k x - s + 1), \quad k \in \mathbb{N}_0, \quad s = 1, \dots, 2^k.$$

One can also verify that

$$\text{supp } \psi_{k,s} = I_{k,s} =: [2^{-k}(s-1), 2^{-k}s], \quad \text{int } I_{k,s} \cap \text{int } I_{k,s'} = \emptyset, \quad s \neq s'.$$

For every $n \in \mathbb{N}$, $n \geq 2$, we choose $m \in \mathbb{N}$ such that $2^{m-2} < n \leq 2^{m-1}$. Let J_m be the set of s such that $X_n \cap \text{int}(I_{m,s} \times [0, 1]^{d-1}) = \emptyset$. Then we have $|J_m| \geq 2^{m-1}$. We define the function

$$f_m(\mathbf{x}) := 18^{-d} 2^{-\alpha m} \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \prod_{j=2}^d \psi_{0,1}(x_j), \quad \mathbf{x} \in \mathbb{I}^d$$

and show that

$$\Delta_{\mathbf{h}}^{2,u}(f_m, \mathbf{x}) \leq \prod_{j \in u} |h_j|^\alpha, \quad \mathbf{x} \in \mathbb{I}^d, \quad \mathbf{h} \in [-1, 1]^d, \quad u \subset [d], \quad (3.15)$$

following partly in [11]. Let us prove this inequality for $u = [d]$ and $\mathbf{h} \in \mathbb{I}^d$, the general case of u can be proven in a similar way with a slight modification. We have

$$\Delta_{\mathbf{h}}^{2,[d]}(f_m, \mathbf{x}) = 18^{-d} 2^{-\alpha m} \Delta_{h_1}^2 \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \prod_{j=2}^d \Delta_{h_j}^2 \psi_{0,1}(x_j).$$

By using the formula

$$\Delta_h^2(f, x) = h^2 \int_{\mathbb{R}} f^{(2)}(x+y) [h^{-1} M_2(h^{-1}y)] dy, \quad h \in \mathbb{R}$$

for univariate function f having locally absolutely continuous f' , see e.g. [13, page 45], we get

$$\begin{aligned} \Delta_{h_1}^2 \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) &= h_1^2 \int_{\mathbb{R}} \left(\sum_{s \in J_m} \psi_{m,s}(t) \right)'' h_1^{-1} M_2((h_1^{-1}(t-x_1))) dt \\ &= 9h_1^2 2^{2m} \int_{\mathbb{R}} \left(\sum_{s \in J_m} \chi_{m,s}^0(t) \right) h_1^{-1} M_2((h_1^{-1}(t-x_1))) dt, \end{aligned}$$

where $\chi_{m,s}^0(t) = \chi_{I_{m,s}^1} - 2\chi_{I_{m,s}^2} + \chi_{I_{m,s}^3}$ and $\chi_{I_{m,s}^j}$ are characteristic functions of the intervals

$$I_{m,s}^j := \left[2^{-m}(s-1) + \frac{2^{-m}(j-1)}{3}, 2^{-m}(s-1) + \frac{2^{-m}j}{3} \right], \quad j = 1, 2, 3.$$

If $(2^m h_1) \leq 1$ we have from $|\chi_{m,s}^0| \leq 2$

$$\begin{aligned} 2^{-\alpha m} \left| \Delta_{h_1}^2 \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \right| &= 9h_1^\alpha (2^m h_1)^{(2-\alpha)} \left| \int_{\mathbb{R}} \left(\sum_{s \in J_m} \chi_{m,s}^0(t) \right) h_1^{-1} M_2((h_1^{-1}(t-x_1))) dt \right| \\ &\leq 18h_1^\alpha \int_{\mathbb{R}} h_1^{-1} M_2((h_1^{-1}(t-x_1))) dt = 18h_1^\alpha, \end{aligned} \quad (3.16)$$

where in the last equality we used $\int_{\mathbb{R}} M_2(t) dt = 1$. If $(2^m h_1) > 1$ we have by changing variable

$$2^{-\alpha m} \left| \Delta_{h_1}^2 \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \right| = 9h_1^2 2^{m(2-\alpha)} \int_{\mathbb{R}} \left(\sum_{s \in J_m} \chi_{m,s}^0(x_1 + h_1 y) \right) M_2(y) dy.$$

Denote $K_{m,s} = \text{supp } \chi_{m,s}(x_1 + h_1 \cdot) = [\frac{2^{-m}(s-1)-x_1}{h_1}, \frac{2^{-m}s-x_1}{h_1}]$. If $K_{m,s} \subset [0, 1]$ or $K_{m,s} \subset [1, 2]$ we have

$$\int_{\mathbb{R}} \chi_{m,s}^0(x_1 + h_1 y) M_2(y) dy = \int_{\mathbb{R}} (\chi_{I_{m,s}^1} - 2\chi_{I_{m,s}^2} + \chi_{I_{m,s}^3})(x_1 + h_1 y) M_2(y) dy = 0$$

and there are at most three $s^j \in J_m$, $j = 0, 1, 2$ such that $j \in \text{int}(K_{m,s^j})$, $j = 0, 1, 2$. It is not difficult to verify that

$$\left| \int_{\mathbb{R}} \chi_{m,s^1}^0(x_1 + h_1 y) M_2(y) dy \right| \leq \frac{3}{2} \left(\frac{1}{3 \cdot 2^m h_1} \right)^2$$

and

$$\left| \int_{\mathbb{R}} \chi_{m,s^j}^0(x_1 + h_1 y) M_2(y) dy \right| \leq \frac{1}{2} \left(\frac{1}{3 \cdot 2^m h_1} \right)^2$$

if $j = 0, 2$. From this we obtain

$$\begin{aligned} 2^{-\alpha m} \left| \Delta_{h_1}^2 \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \right| &= 9h_1^2 2^{m(2-\alpha)} \left| \int_{\mathbb{R}} \left(\sum_{j=0,1,2} \chi_{m,s^j}^0(x_1 + h_1 y) \right) M_2(y) dy \right| \\ &\leq 9h_1^2 2^{m(2-\alpha)} \frac{5}{2} \left(\frac{1}{3 \cdot 2^m h_1} \right)^2 = \frac{5}{2} h^\alpha \frac{1}{(2^m h_1)^\alpha} \leq \frac{5}{2} h_1^\alpha. \end{aligned}$$

Since $h_j \in [0, 1]$, $j = 2, \dots, d$, similar to (3.16) we can show that

$$\left| \prod_{j=2}^d \Delta_{h_j}^2 \psi_{0,1}(x_j) \right| \leq 18^{d-1} \prod_{j=2}^d h_j^\alpha.$$

Consequently, the inequality (3.15) is proven. This means that $f_m \in \dot{U}_\infty^\alpha$. Due to the disjoint supports of $\psi_{m,s}$ we have

$$\begin{aligned} \|f_m\|_{W_0^{1,p}}^p &= 18^{-pd} \sum_{j=1}^d \int_{\mathbb{I}^d} \left| 2^{-\alpha m} \frac{\partial}{\partial x_j} \left(\sum_{s \in J_m} \psi_{m,s}(x_1) \right) \prod_{j=2}^d \psi_{0,1}(x_j) \right|^p dx \\ &\geq 18^{-pd} \int_{\mathbb{I}^d} \left| 2^{-\alpha m} \sum_{s \in J_m} \frac{\partial}{\partial x_1} \psi_{m,s}(x_1) \prod_{j=2}^d \psi_{0,1}(x_j) \right|^p dx \\ &= 18^{-pd} 2^{-p\alpha m} |J_m| \int_{\mathbb{I}} |\psi'_{m,1}(x_1)|^p dx_1 \left(\int_{\mathbb{I}} |\psi_{0,1}(t)|^p dt \right)^{d-1} \\ &\geq 18^{-pd} 2^{-p\alpha m} 2^{m-1} 2^{pm} 3^p \int_{\mathbb{I}} |M'_3(3 \cdot 2^m x_1)|^p dx_1 \left(\int_{\mathbb{I}} |M_3(3t)|^p dt \right)^{d-1} \\ &= \frac{3^p}{2} 18^{-pd} 3^{-d} 2^{-pm(\alpha-1)} \int_{\mathbb{R}} |M'_3(t)|^p dt \left(\int_{\mathbb{R}} |M_3(t)|^p dt \right)^{d-1}, \end{aligned}$$

where in the last equality we changed variables. Since $f_m(\xi) = 0$, $\xi \in X_n$, and

$$\int_{\mathbb{R}} |M'_3(t)|^p dt = \frac{3}{p+1}, \quad M := \left(\int_{\mathbb{R}} M_3^p(t) dt \right)^{1/p},$$

by (3.14) we get

$$\begin{aligned} r_n \geq \|f_m\|_{W_0^{1,p}} &\geq 18^{-d} \left(\frac{3^{p+1-d}}{2p+2} \right)^{1/p} M^{d-1} 2^{-m(\alpha-1)} \geq 18^{-d} \left(\frac{3^{p+1-d}}{2p+2} \right)^{1/p} M^{d-1} \left(\frac{n^{-1}}{4} \right)^{(\alpha-1)} \\ &\geq \frac{1}{M 4^{\alpha-1}} \left(\frac{3^{p+1}}{2p+2} \right)^{1/p} (18 \cdot 3^{1/p} M^{-1})^{-d} n^{-(\alpha-1)} = \frac{1}{M 4^{\alpha-1}} \left(\frac{3^{p+1}}{2p+2} \right)^{1/p} B_*^{-d} n^{-(\alpha-1)}. \end{aligned}$$

The proof is completed. $\square$

Corollary 3.6 *Let $d \geq 2$, $1 < \alpha \leq 2$, $\beta > \alpha$, and $1 \leq p \leq \infty$. If $2^{\frac{2}{\alpha}} - (p+1)^{-\frac{1}{p\alpha}} > 2$, then there is $n_0 \in \mathbb{N}$ such that for $n \geq n_0$ we have*

$$C_1 B_*^{-d} n^{-(\alpha-1)} \leq r_n \leq C_2 B^* n^{-(\alpha-1)}$$

with $B_ > 1$ and $B^* > 1$, where the constants $B_* = B_*(p)$, $B^* = B^*(\alpha, p)$, $C_1 = C_1(p)$, $C_2 = C_2(\alpha, p)$ are independent of n and d .*

Proof. The lower bound in the corollary for the case $1 \leq p \leq \infty$ follows from the case $p = 1$ which is obvious due to the lower bound in Theorem 3.5. By the definitions of B_1 in (3.8) and of B^* in Theorem 3.5, to prove the upper bound it is sufficient to show that under the condition $2^{\frac{2}{\alpha}} - (p+1)^{-\frac{1}{p\alpha}} > 2$ it holds

$$2^{1-\alpha} (1 - 2^{-\frac{1}{\beta-1}})^{(\alpha-1)} (p+1)^{\frac{1}{p}} 2^{\alpha+1} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}}) > 1$$

or

$$(1 - 2^{-\frac{1}{\beta-1}})^{(\alpha-1)} (p+1)^{\frac{1}{p}} 2^2 (1 - 2^{-\frac{\beta-\alpha}{\beta-1}}) > 1.$$

The further argument is similar to the proof of Theorem 3.3. $\square$

4 Approximation by deep ReLU neural networks

In this section, based on the result on sparse-grid sampling recovery in Theorem 3.3, for any function $f \in \dot{U}_\infty^\alpha$, we explicitly construct a deep ReLU neural network having an output that approximates f in the $W_0^{1,p}$ -norm with a prescribed accuracy ε and prove dimension-dependent error estimates of the number of its weights and its depth. We also prove some lower bounds for the number of weights of this deep ReLU neural network necessary for this approximation.

There is a wide variety of neural network architectures and each of them is adapted to specific tasks. For our purpose of approximation functions from Höder-Zygmund spaces, in this section we introduce feed-forward deep ReLU neural networks with one-dimension output. We are interested in deep neural networks where only connections between neighboring layers are allowed. Let us introduce necessary definitions and elementary facts on deep ReLU neural networks.

Definition 4.1 *Let $d, L \in \mathbb{N}$, $L \geq 2$.*

- *A deep neural network Φ with input dimension d and L layers is a sequence of matrix-vector tuples*

$$\Phi = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L))$$

where $\mathbf{W}^\ell = (w_{i,j}^\ell)$ is an $N_\ell \times N_{\ell-1}$ matrix, and $\mathbf{b}^\ell = (b_j^\ell) \in \mathbb{R}^{N_\ell}$ with $N_0 = d$, $N_L = 1$, and $N_1, \dots, N_{L-1} \in \mathbb{N}$. We call the number of layers $L(\Phi) = L$ the depth and $\mathbf{N}(\Phi) = (N_0, N_1, \dots, N_L)$ the dimension of the network. The real numbers $w_{i,j}^\ell$ and b_j^ℓ are called edge and node weights of the network Φ , respectively. The number of nonzero weights $w_{i,j}^\ell$ and b_j^ℓ is called the number of weights of the network Φ and denoted by $W(\Phi)$, i.e., $W(\Phi) := \sum_{\ell=1}^L |\mathbf{W}^\ell|_0 + \sum_{\ell=1}^L |\mathbf{b}^\ell|_0$. We call $N_w(\Phi) = \max_{\ell=0, \dots, L} \{N_\ell\}$ the width of the network Φ .

- *A neural network architecture $\mathbb{A}$ with input dimension d and L layers is a neural network*

$$\mathbb{A} = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L)),$$

where elements of $\mathbf{W}^\ell$ and $\mathbf{b}^\ell$, $\ell = 1, \dots, L$, are in $\{0, 1\}$.

Since we only are interested in neural network with scalar output, $\mathbf{b}^L$ is a constant. However for consistent notation, we still use bold letter.

A graph associated to a deep neural network Φ defined in Definition 4.1 is a graph consisting of $|\mathcal{N}(\Phi)|_1$ nodes and $\sum_{\ell=1}^L |\mathbf{W}^\ell|_0$ edges. $|\mathcal{N}(\Phi)|_1$ nodes are placed in $L + 1$ layers which are numbered from 0 to L . The ℓ th layer has N_ℓ nodes which are numbered from 1 to N_ℓ . If $w_{i,j}^\ell \neq 0$, then there is an edge connecting the node j in the layer $\ell - 1$ to the node i in the layer ℓ . See Figure 2 for an illustration of a graph associated to a deep neural network.

Definition 4.2 Given $L \in \mathbb{N}$, $L \geq 2$, and a deep neural network architecture $\mathbb{A} = ((\overline{\mathbf{W}}^1, \overline{\mathbf{b}}^1), \dots, (\overline{\mathbf{W}}^L, \overline{\mathbf{b}}^L))$. We say that a neural network $\Phi = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L))$ has architecture $\mathbb{A}$ if

- $\mathcal{N}(\Phi) = \mathcal{N}(\mathbb{A})$
- $\overline{w}_{i,j}^\ell = 0$ implies $w_{i,j}^\ell = 0$, $\overline{b}_i^\ell = 0$ implies $b_i^\ell = 0$ for all $i = 1, \dots, N_\ell$, $j = 1, \dots, N_{\ell-1}$, and $\ell = 1, \dots, L$. Here $\overline{w}_{i,j}^\ell$ are entries of $\overline{\mathbf{W}}^\ell$ and $\overline{b}_i^\ell$ are elements of $\overline{\mathbf{b}}^\ell$, $\ell = 1, \dots, L$.

For a given deep neural network $\Phi = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L))$, there exists a unique deep neural network architecture $\mathbb{A} = ((\overline{\mathbf{W}}^1, \overline{\mathbf{b}}^1), \dots, (\overline{\mathbf{W}}^L, \overline{\mathbf{b}}^L))$ such that

- $\mathcal{N}(\Phi) = \mathcal{N}(\mathbb{A})$
- $\overline{w}_{i,j}^\ell = 0 \iff w_{i,j}^\ell = 0$, $\overline{b}_i^\ell = 0 \iff b_i^\ell = 0$ for all $i = 1, \dots, N_\ell$, $j = 1, \dots, N_{\ell-1}$, and $\ell = 1, \dots, L$.

We call this architecture $\mathbb{A}$ the minimal architecture of Φ (this definition is proper in the sense that any architecture of Φ is also an architecture of $\mathbb{A}$.)

A deep neural network associated with an activation function. In this paper we focus on ReLU (Rectified Linear Unit) activation function defined by $\sigma(t) := \max\{t, 0\}$, $t \in \mathbb{R}$. We will use the notation $\sigma(\mathbf{x}) := (\sigma(x_1), \dots, \sigma(x_d))$ for $\mathbf{x} \in \mathbb{R}^d$.

Definition 4.3 A deep ReLU neural network with input dimension d and L layers is a neural network

$$\Phi = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L))$$

in which the following computation scheme is implemented

$$\begin{aligned} \mathbf{z}^0 &:= \mathbf{x} \in \mathbb{R}^d, \\ \mathbf{z}^\ell &:= \sigma(\mathbf{W}^\ell \mathbf{z}^{\ell-1} + \mathbf{b}^\ell), \quad \ell = 1, \dots, L-1, \\ \mathbf{z}^L &:= \mathbf{W}^L \mathbf{z}^{L-1} + \mathbf{b}^L. \end{aligned}$$

We call $\mathbf{z}^0$ the input and $\mathcal{N}(\Phi, \mathbf{x}) := \mathbf{z}^L$ the output of Φ .

Several deep ReLU neural networks can be combined into a larger deep ReLU neural network whose output is a linear combination of outputs of sub-networks as in the following lemma. This combination is called parallelization. For other combinations, such as concatenation, we refer to [18, Section 2] or [12]. Note that our parallelization construction differs slightly from that in [18] since we only consider deep ReLU neural networks with scalar output.

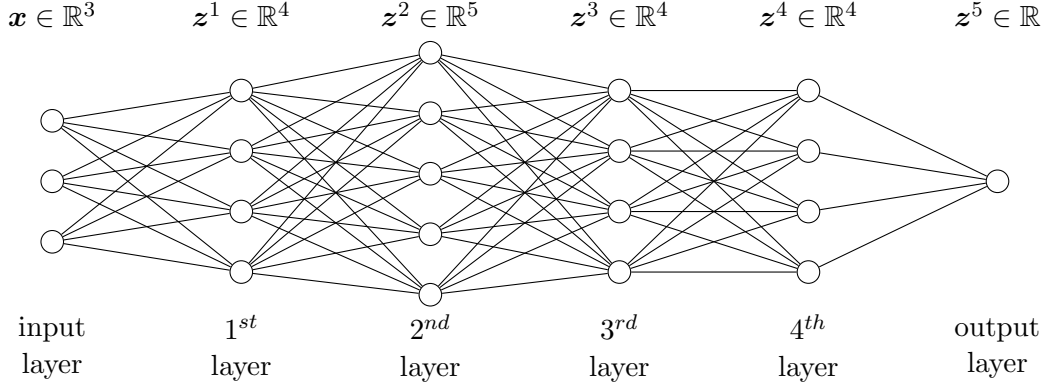


Figure 2: The graph associated to a deep neural network with input dimension 3 and 5 layers

Lemma 4.1 *Let $N \in \mathbb{N}$, $\Omega \subset \mathbb{R}^d$ be a bounded set, $\lambda_j \in \mathbb{R}$, $j = 1, \dots, N$. Let Φ_j , $j = 1, \dots, N$ be deep neural networks with input dimension d , L_j layers, and W_j number of weights respectively. Then there is a deep neural network denoted by Φ such that*

$$\mathcal{N}(\Phi, \mathbf{x}) = \sum_{j=1}^N \lambda_j \mathcal{N}(\Phi_j, \mathbf{x}), \quad \mathbf{x} \in \Omega,$$

with $L(\Phi) = \max_{j=1, \dots, N} \{L_j\}$ and $W(\Phi) = \sum_{j=1}^N W_j + \sum_{j: L_j < L} (L - L_j + 2)$.

Proof. We prove first for $N = 2$. Without loss of generality we assume that $L_1 \leq L_2$ and

$$\Phi_1 = ((\mathbf{W}_1^1, \mathbf{b}_1^1), \dots, (\mathbf{W}_1^{L_1}, \mathbf{b}_1^{L_1})); \quad \Phi_2 = ((\mathbf{W}_2^1, \mathbf{b}_2^1), \dots, (\mathbf{W}_2^{L_2}, \mathbf{b}_2^{L_2})).$$

If $L = L_1 = L_2$, then we can choose

$$\Phi = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L))$$

where

$$\mathbf{W}^1 = \begin{bmatrix} \mathbf{W}_1^1 \\ \mathbf{W}_2^1 \end{bmatrix}, \quad \mathbf{W}^\ell = \begin{bmatrix} \mathbf{W}_1^\ell & 0 \\ 0 & \mathbf{W}_2^\ell \end{bmatrix}, \quad \ell = 2, \dots, L-1, \quad \mathbf{W}^L = [\lambda_1 \mathbf{W}_1^L \quad \lambda_2 \mathbf{W}_2^L]$$

and

$$\mathbf{b}^1 = \begin{bmatrix} \mathbf{b}_1^1 \\ \mathbf{b}_2^1 \end{bmatrix}, \quad \ell = 1, \dots, L_1, \quad \mathbf{b}^L = \lambda_1 \mathbf{b}_1^L + \lambda_2 \mathbf{b}_2^L.$$

In this case we have $W(\Phi) \leq W_1 + W_2$. If $L_1 < L_2$ we construct a network $\tilde{\Phi}_1$ with output $\mathcal{N}(\tilde{\Phi}_1, \mathbf{x}) = \mathcal{N}(\Phi_1, \mathbf{x})$ and having L_2 layers. The strategy here is to modify the network Φ_1 by making its output layer satisfying $\mathcal{N}(\tilde{\Phi}_1, \mathbf{x}) + C \geq 0$, for some constant C , so that this value does not change when we apply function σ for layers from $L_1 + 1$ to L_2 . For this we put $M_1 = \sup_{\mathbf{x} \in \Omega} |\mathcal{N}(\Phi_1, \mathbf{x})|$. Note that $\mathcal{N}(\Phi_1, \cdot)$ is a continuous function on Ω hence $M_1 < \infty$. The network $\tilde{\Phi}_1$ is

$$((\mathbf{W}_1^1, \mathbf{b}_1^1), \dots, (\mathbf{W}_1^{L_1-1}, \mathbf{b}_1^{L_1-1}), (\mathbf{W}_1^{L_1}, \mathbf{b}_1^{L_1} + M_1), (1, 0), \dots, (1, 0), (1, -M_1)).$$

Hence $W(\tilde{\Phi}_1) \leq W_1 + L_2 - L_1 + 2$. Now following procedure as the case $L_1 = L_2$ with Φ_1 replaced by $\tilde{\Phi}_1$ we obtain the assertion when $N = 2$. The case $N > 2$ is extended in a similar manner. $\square$

An illustration of parallelization of neural networks is given in Figure 3.

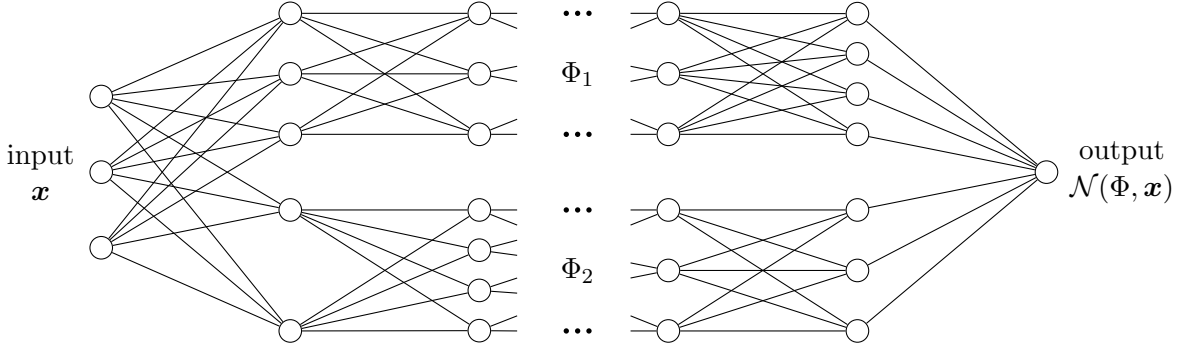


Figure 3: The graph associated to parallelization of two neural networks

Let us introduce a concept of special deep neural network borrowed from [12]. It is quite useful in construction deep ReLU neural networks with a fixed width, whose outputs are able to approximate multivariate functions.

Definition 4.4 *A special deep neural network with input dimension d and depth L (and a given activation function) can be defined as follows. In each hidden layer a special role is reserved for d first (top) nodes and the last (bottom) node. The top d nodes and the bottom node are free of the activation function, other nodes in each hidden layer have the activation function. The top d nodes are used to simply copy the input $\mathbf{x}$. The d parallel concatenations of all these top nodes can be viewed as special channels that skip computation altogether and just carry $\mathbf{x}$ forward. They are called the source channels. The bottom node in each hidden layer is used to collect intermediate outputs by addition. The concatenation of all these nodes is called collation channel. This channel never feeds forward into subsequent calculation, it only accepts previous calculations.*

An illustration of special deep neural network is given in Figure 4.

Lemma 4.2 *Let Φ be a special deep ReLU neural network with input dimension d depth L . Then there is a deep ReLU neural network Φ' such that, $N_w(\Phi') = N_w(\Phi)$, $L(\Phi') = L$, and $\mathcal{N}(\Phi', \mathbf{x}) = \mathcal{N}(\Phi, \mathbf{x})$, $\mathbf{x} \in \mathbb{I}^d$.*

Proof. The proof follows from [12, Remark 3.1]. First note, that the input $\mathbf{x}$ belongs to $\mathbb{I}^d = [0, 1]^d$, we have $\mathbf{x} = \sigma(\mathbf{x})$. For $\ell = 1, \dots, L$, the bottom node in the ℓ -th layer collects a continuous piece-wise linear function $g_\ell(\mathbf{x})$ on $\mathbb{I}^d$, and the output is $\mathcal{N}(\Phi, \mathbf{x}) = \sum_{\ell=1}^{L-1} g_\ell(\mathbf{x})$. Thus, there is a constant c_ℓ such that $g_\ell(\mathbf{x}) + c_\ell \geq 0$ for all $\mathbf{x} \in \mathbb{I}^d$. Hence we take Φ' having the same graph as Φ but the computation at ℓ -th nodes in the collation channel is replaced by $\sigma(g_\ell(\mathbf{x}) + c_\ell) = g_\ell(\mathbf{x}) + c_\ell$, and the output node is

$$\mathcal{N}(\Phi', \mathbf{x}) = \sum_{\ell=1}^{L-1} \sigma(g_\ell(\mathbf{x}) + c_\ell) - \sum_{\ell=1}^{L-1} c_\ell = \mathcal{N}(\Phi, \mathbf{x}).$$

□

We now consider the problem of approximation of functions from $\mathring{U}_\infty^\alpha$ by deep ReLU neural networks. For every $\varepsilon > 0$ and every $f \in \mathring{U}_\infty^\alpha$, we will construct a deep ReLU neural network Φ_f having an architecture $\mathbb{A}_\varepsilon$ independent of f , and the output $\mathcal{N}(\Phi_f, \cdot)$ which approximates f with accuracy ε , and give dimension-dependent upper bounds for the number of weights and the depth of Φ_f . Our strategy is as follows. Based on the result in Theorem 3.3 on approximation of f by the sparse-grid

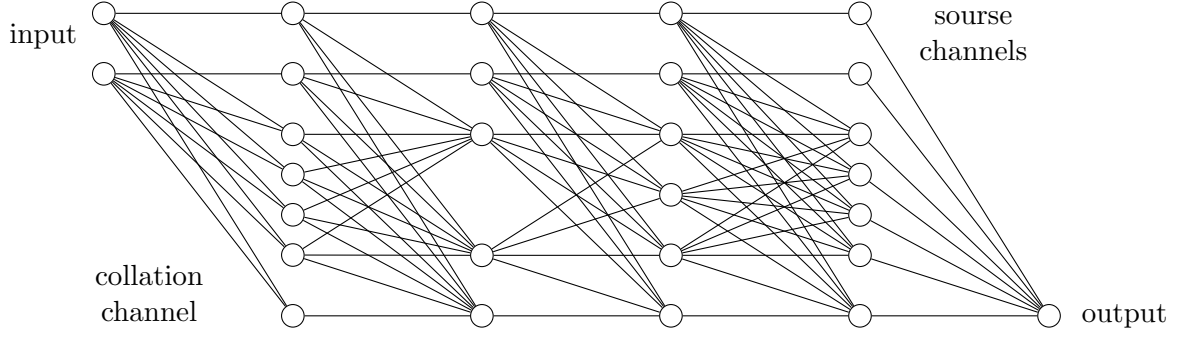


Figure 4: The graph associated to a special neural network with two source channels and 5 layers

sampling sum $R_\beta(m, f)$, we use $R_\beta(m, f)$ as a mediate approximation, and then construct a deep ReLU neural network Φ_f for approximating this sum by the output $\mathcal{N}(\Phi_f, \cdot)$. Since $R_\beta(m, f)$ is a sum of tensor products of hat functions, first of all we will process the approximation such tensor products by deep ReLU neural networks which can be done based on a result in the following lemma which has been proven in [30, Proposition 2.6] on approximating by deep ReLU neural networks the product of d numbers.

Lemma 4.3 *For any $\delta \in (0, 1)$, $d \in \mathbb{N}$, $d \geq 2$, there exists a deep ReLU neural network Φ_P such that*

$$\sup_{\mathbf{x} \in [-1, 1]^d} \left| \prod_{i=1}^d x_i - \mathcal{N}(\Phi_P, \mathbf{x}) \right| \leq \delta,$$

and

$$\text{ess sup}_{\mathbf{x} \in [-1, 1]^d} \sup_{j=1, \dots, d} \left| \frac{\partial}{\partial x_j} \prod_{i=1}^d x_i - \frac{\partial}{\partial x_j} \mathcal{N}(\Phi_P, \mathbf{x}) \right| \leq \delta,$$

where $\frac{\partial}{\partial x_j}$ denotes a weak partial derivative. Furthermore, there exists a constant C independent of $\delta \in (0, 1)$ and $d \in \mathbb{N}$ such that

$$W(\Phi_P) \leq Cd \log(d\delta^{-1}) \quad \text{and} \quad L(\Phi_P) \leq C \log d \log(d\delta^{-1}).$$

Moreover, if $x_j = 0$ for some $j \in [d]$, then $\mathcal{N}(\Phi_P, \mathbf{x}) = 0$.

The last statement $\mathcal{N}(\Phi_P, \mathbf{x}) = 0$ ($d = 2$) when $x_1 \cdot x_2 = 0$ was proved in [34, Proposition 3.1], see also [18, Proposition C.2] and [29, Proposition 4.1]. But this implies that the statement also holds for general d since the network Φ_P is constructed as an binary tree of the network Φ_P when $d = 2$. Inspecting the proof of Lemma 4.3, see [30, Proposition 2.6] and [34, Proposition 3.3] we also find that when $\mathbf{x} \in \mathbb{I}^d$ it holds $N_w(\Phi_P) \leq 12d$. The above lemma immediately yields the following lemma on approximating by deep ReLU neural networks tensor products of d hat functions.

Lemma 4.4 *For any dimension $d \geq 2$, $0 < \delta < 1$ and for the d -variate hat functions $\varphi_{\mathbf{k}, \mathbf{s}}$, $\mathbf{k} \in \mathbb{N}_0^d$, $\mathbf{s} \in Z^d(\mathbf{k})$, defined as in (2.3), there is a deep ReLU neural network $\Phi_{\mathbf{k}, \mathbf{s}}$ such that $\mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \cdot)$ approximates $\varphi_{\mathbf{k}, \mathbf{s}}$ with accuracy δ , and*

$$N_w(\Phi_{\mathbf{k}, \mathbf{s}}) \leq Cd, \quad W(\Phi_{\mathbf{k}, \mathbf{s}}) \leq Cd \log(d\delta^{-1}), \quad \text{and} \quad L(\Phi_{\mathbf{k}, \mathbf{s}}) \leq C \log d \log(d\delta^{-1}).$$

Moreover, $\text{supp}(\mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \cdot)) \subset \text{supp} \varphi_{\mathbf{k}, \mathbf{s}}$ and

$$\text{ess sup}_{\mathbf{x} \in \mathbb{I}^d} \left| \frac{\partial}{\partial x_j} \varphi_{\mathbf{k}, \mathbf{s}}(\mathbf{x}) - \frac{\partial}{\partial x_j} \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}) \right| \leq 2^{k_j+1} \delta.$$

Proof. Indeed, we write

$$y_i := \varphi_{k_i, s_i}(x_i) = \sigma(1 - \sigma(2^{k_i+1}x_i - 2(s_i + 1))) - \sigma(2(s_i + 1) - 2^{k_i+1}x_i).$$

Let Φ_P be the deep ReLU neural network in Lemma 4.3 and $\mathbf{y} = (y_1, \dots, y_d)$ be the inputs of Φ_P . Then we obtain a deep ReLU neural network denoted by $\Phi_{\mathbf{k}, \mathbf{s}}$. We have

$$\sup_{\mathbf{x} \in \mathbb{I}^d} \left| \prod_{i=1}^d \varphi_{k_i, s_i}(x_i) - \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}) \right| = \sup_{\mathbf{y} \in \mathbb{I}^d} \left| \prod_{i=1}^d y_i - \mathcal{N}(\Phi_P, \mathbf{y}) \right| \leq \delta$$

and

$$\text{ess sup}_{\mathbf{x} \in \mathbb{I}^d} \left| \frac{\partial}{\partial x_j} \varphi_{\mathbf{k}, \mathbf{s}}(\mathbf{x}) - \frac{\partial}{\partial x_j} \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}) \right| = \text{ess sup}_{\mathbf{y} \in \mathbb{I}^d} \left| \left(\frac{\partial}{\partial y_j} \prod_{i=1}^d y_i - \frac{\partial}{\partial y_j} \mathcal{N}(\Phi_P, \mathbf{y}) \right) \frac{dy_j}{dx_j} \right| \leq 2^{k_j+1} \delta.$$

Moreover, we have

$$L(\Phi_{\mathbf{k}, \mathbf{s}}) = L(\Phi_P) + 2 \quad \text{and} \quad W(\Phi_{\mathbf{k}, \mathbf{s}}) \leq W(\Phi_P) + 7d.$$

Now, from Lemma 4.3 we obtain the desired result. $\square$

We now in position to formulate and prove the following theorem.

Theorem 4.5 *Let $d \geq 2$, $1 < \alpha \leq 2$, $\beta \geq \alpha$ and $1 \leq p \leq \infty$. Let*

$$\varepsilon_0 = \min \left\{ 1, \frac{d}{2^{\alpha d}(1 - 2^{1-\alpha})}, \frac{K_1 d^2}{(p+1)^{d/p} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right\},$$

where K_1 is the constant given in Theorem 3.3.

Then for every $\varepsilon \in (0, \varepsilon_0)$ there exists a deep neural network architecture $\mathbb{A}_\varepsilon$ with the following property. For every $f \in \dot{U}_\infty^\alpha$ there exist a deep ReLU neural network Φ_f having the architecture $\mathbb{A}_\varepsilon$, and positive constants $K_3 = K_3(\alpha)$, $K_4 = K_4(\alpha, \beta)$, such that

$$\|f - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} \leq \varepsilon, \quad (4.1)$$

and there hold the estimates

$$L(\mathbb{A}_\varepsilon) \leq K_3 \log d \log(\varepsilon^{-1}) \quad \text{and} \quad W(\mathbb{A}_\varepsilon) \leq K_4 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}),$$

where

$$B_2 = B_2(d, \alpha, \beta) := (1 - 2^{-\frac{1}{\beta-1}}) \left(\frac{(p+1)^{\frac{1}{p}} 2^{(\alpha+1)} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})}{d^{\frac{2\alpha}{d}}} \right)^{\frac{1}{\alpha-1}}. \quad (4.2)$$

Proof. Let $f \in \dot{U}_\infty^\alpha$. For $\varepsilon \in (0, \varepsilon_0)$ we take

$$m = \left\lceil \frac{1}{\alpha-1} \log \left(\frac{2K_1 d^2 \varepsilon^{-1}}{(p+1)^{d/p} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right) \right\rceil. \quad (4.3)$$

Let $(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)$ and $\Phi_{\mathbf{k}, \mathbf{s}}$ be the deep ReLU neural network obtained in Lemma 4.4 such that $\mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \cdot)$ approximates $\varphi_{\mathbf{k}, \mathbf{s}}$ with accuracy $\delta \in (0, 1)$ which is chosen later. Let Φ_f be the deep ReLU neural network obtained by parallelization as in Lemma 4.1 with the output

$$\mathcal{N}(\Phi_f, \mathbf{x}) = \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k}, \mathbf{s}}(f) \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}), \quad \mathbf{x} \in \mathbb{I}^d.$$

Then we can write

$$\|f - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} \leq \|f - R_\beta(m, f)\|_{W_0^{1,p}} + \|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}}, \quad (4.4)$$

where $R_\beta(m, f)$ is the operator given in (3.1). With the choice of m as in (4.3) we get from Theorem 3.1

$$\|f - R_\beta(m, f)\|_{W_0^{1,p}} \leq K_1 \frac{d^2 2^{-m(\alpha-1)}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \leq \varepsilon/2. \quad (4.5)$$

Let us estimate the norm $\|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}}$. Since $\text{supp } \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \cdot) \subset \text{supp } \varphi_{\mathbf{k}, \mathbf{s}}$ we have

$$\begin{aligned} \|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} &\leq \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \left\| \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k}, \mathbf{s}}(f) (\varphi_{\mathbf{k}, \mathbf{s}} - \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \cdot)) \right\|_{W_0^{1,p}} \\ &= \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \left(\sum_{j=1}^d \int_{\mathbb{I}^d} \left| \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \lambda_{\mathbf{k}, \mathbf{s}}(f) \left(\frac{\partial}{\partial x_j} \varphi_{\mathbf{k}, \mathbf{s}}(\mathbf{x}) - \frac{\partial}{\partial x_j} \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}) \right) \right|^p d\mathbf{x} \right)^{1/p} \\ &= \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \left(\sum_{j=1}^d \sum_{\mathbf{s} \in Z^d(\mathbf{k})} |\lambda_{\mathbf{k}, \mathbf{s}}(f)|^p \int_{\mathbb{I}^d} \left| \frac{\partial}{\partial x_j} \varphi_{\mathbf{k}, \mathbf{s}}(\mathbf{x}) - \frac{\partial}{\partial x_j} \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}) \right|^p d\mathbf{x} \right)^{1/p}. \end{aligned}$$

Using Lemma 4.4 and estimate (2.5) we get

$$\begin{aligned} \|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} &\leq \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \sup_{\mathbf{s} \in Z^d(\mathbf{k})} |\lambda_{\mathbf{k}, \mathbf{s}}(f)| \left(\sum_{j=1}^d \sum_{\mathbf{s} \in Z^d(\mathbf{k})} 2^{-|\mathbf{k}|_1} 2^{p(k_j+1)} \right)^{1/p} \delta \\ &\leq 2^{-(\alpha+1)d} \sum_{\mathbf{k} \in \Delta_\beta^d(m)} 2^{-|\mathbf{k}|_1 \alpha} \left(\sum_{j=1}^d 2^{p(k_j+1)} \right)^{1/p} \delta \\ &\leq 2^{-(\alpha+1)d} \sum_{\ell=0}^{\infty} 2^{-\ell \alpha} \sum_{|\mathbf{k}|_1=\ell} \left(\sum_{j=1}^d 2^{p(k_j+1)} \right)^{1/p} \delta. \end{aligned}$$

Now Lemma 3.2 leads to

$$\|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} \leq d 2^d 2^{-(\alpha+1)d} \delta \sum_{\ell=0}^{\infty} 2^{-\ell(\alpha-1)} \leq \frac{d 2^{-\alpha d}}{1 - 2^{1-\alpha}} \delta.$$

Define $\delta = \delta(\varepsilon) := \frac{1-2^{1-\alpha}}{d 2^{-\alpha d}} \frac{\varepsilon}{2}$. Since $\varepsilon < \frac{2d}{2^{\alpha d}(1-2^{1-\alpha})}$ we get $\delta < 1$. This choice of δ gives

$$\|R_\beta(m, f) - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,p}} \leq \varepsilon/2$$

which together with (4.4) and (4.5) proves (4.1).

We now prove the bounds for the depth and the number of weights of Φ_f . From Lemmata 4.1 and 3.2 we have

$$\begin{aligned} L(\Phi_f) &= \max_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} L(\Phi_{\mathbf{k}, \mathbf{s}}) \leq C \log d \log(d \delta^{-1}) \\ &\leq C \log d \log \left(\frac{2d^2 \varepsilon^{-1}}{2^{\alpha d} (1 - 2^{1-\alpha})} \right) \leq K_3 \log d \log(\varepsilon^{-1}) \end{aligned} \quad (4.6)$$

for some positive constant $K_3 = K_3(\alpha)$, and

$$\begin{aligned}
W(\Phi_f) &\leq \sum_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} W(\Phi_{\mathbf{k}, \mathbf{s}}) + \sum_{(\mathbf{k}, \mathbf{s}) : L(\Phi_{\mathbf{k}, \mathbf{s}}) < L(\Phi_f)} (L(\Phi_f) - L(\Phi_{\mathbf{k}, \mathbf{s}}) + 2) \\
&\leq \sum_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} W(\Phi_{\mathbf{k}, \mathbf{s}}) + \sum_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} \left(\max_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} W(\Phi_{\mathbf{k}, \mathbf{s}}) + 2 \right) \\
&\leq 2|D_\beta^d(m)| \max_{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)} (W(\Phi_{\mathbf{k}, \mathbf{s}}) + 1) \leq \frac{2\beta}{\beta - 1} \frac{d2^m}{(1 - 2^{-\frac{1}{\beta-1}})^d} (Cd \log(d\delta^{-1}) + 1).
\end{aligned} \tag{4.7}$$

Similar to (3.11) we have

$$\begin{aligned}
d^2 2^m (1 - 2^{-\frac{1}{\beta-1}})^{-d} &\leq 2d^2 (1 - 2^{-\frac{1}{\beta-1}})^{-d} \left(\frac{2K_1 d^2 \varepsilon^{-1}}{(p+1)^{\frac{d}{p}} 2^{(\alpha+1)d} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})^d} \right)^{\frac{1}{\alpha-1}} \\
&\leq 2(2K_1)^{\frac{1}{\alpha-1}} (1 - 2^{-\frac{1}{\beta-1}})^{-d} \left(\frac{d^{\frac{2\alpha}{d}}}{(p+1)^{\frac{1}{p}} 2^{(\alpha+1)} (1 - 2^{-\frac{\beta-\alpha}{\beta-1}})} \right)^{\frac{d}{\alpha-1}} (\varepsilon^{-1})^{\frac{1}{\alpha-1}}.
\end{aligned}$$

Inserting this into (4.7) we find

$$W(\Phi_f) \leq K_4 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} (\log \varepsilon^{-1})$$

with B_2 given in (4.2) and some positive constant K_4 depending on α and β .

To complete the proof it is sufficient to notice that Φ_f has the architecture $\mathbb{A}_\varepsilon$ (independent of f) which is defined as the minimal architecture of the deep ReLU neural network Φ obtained by parallelization as in Lemma 4.1 with the output

$$\mathcal{N}(\Phi, \mathbf{x}) = \sum_{\mathbf{k} \in \Delta_\beta^d(m)} \sum_{\mathbf{s} \in Z^d(\mathbf{k})} \mathcal{N}(\Phi_{\mathbf{k}, \mathbf{s}}, \mathbf{x}), \quad \mathbf{x} \in \mathbb{I}^d.$$

□

Remark 4.6 Since $d^{2\alpha/d}$ tends to 1 when $d \rightarrow \infty$, if α and p satisfy (3.9), β satisfies (3.10), then there are $d(\alpha, \beta, p) \in \mathbb{N}$ and $B(\alpha, \beta, p) > 1$ such that $B_2 \geq B(\alpha, \beta, p) > 1$ for all $d \geq d(\alpha, \beta, p)$. In this case, the approximation error is going to 0 as fast as the exponent B_2^{-d} when d tending to ∞ . See also Remark 3.4.

For $f \in \mathring{U}_\infty^\alpha$ and $\varepsilon \in (0, \varepsilon_0)$, we observe that the width of the deep ReLU neural network Φ_f constructed in Theorem 4.5 may depend on ε . To construct a deep neural network Φ'_f with the same output that has width independent of ε we can concatenate the deep ReLU networks $\Phi_{\mathbf{k}, \mathbf{s}}$, $(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)$, with the help of special deep ReLU networks.

Theorem 4.7 Under the assumptions and notations of Theorem 4.5, for every $\varepsilon \in (0, \varepsilon_0)$ there exists a deep neural network architecture $\mathbb{A}_\varepsilon^*$ with the following property. For every $f \in \mathring{U}_\infty^\alpha$ there exist a deep ReLU neural network Φ_f^* having the architecture $\mathbb{A}_\varepsilon^*$, and positive constants K_5 and $K_6 = K_6(\alpha, \beta)$ such that

$$\|f - \mathcal{N}(\Phi_f^*, \cdot)\|_{W_0^{1,p}} \leq \varepsilon, \tag{4.8}$$

and there hold the estimates

$$N_w(\mathbb{A}_\varepsilon^*) \leq K_5 d \quad \text{and} \quad L(\mathbb{A}_\varepsilon^*) \leq K_6 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}). \tag{4.9}$$

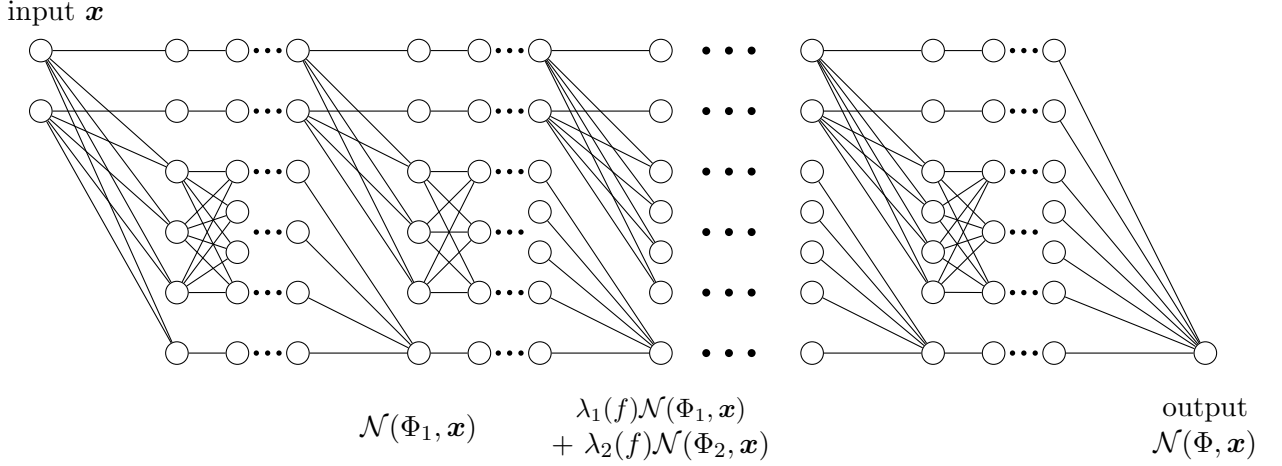


Figure 5: The graph associated to concatenation of neural networks $\Phi_{\mathbf{k}, \mathbf{s}}, (\mathbf{k}, \mathbf{s}) \in D_\beta^d(m), (d = 2)$

Proof. Consider the special deep ReLU neural network Φ'_f as in Figure 5 ($d = 2$). We number the set $\{(\mathbf{k}, \mathbf{s}) \in D_\beta^d(m)\}$ from 1 to J , where $J = |D_\beta^d(m)|$. The d source channels carry $\mathbf{x} \in \mathbb{I}^d$ forward so that it is the input of all networks $\Phi_\ell, \ell = 1, \dots, J$. The $(\sum_{j=1}^\ell L(\Phi_j))$ th node in the collation channel stores the partial sum $\sum_{j=1}^\ell \lambda_j(f) \mathcal{N}(\Phi_j, \mathbf{x})$ of the outputs of $\Phi_\ell, \ell = 1, \dots, J$. Hence,

$$\mathcal{N}(\Phi'_f, \mathbf{x}) = \mathcal{N}(\Phi_f, \mathbf{x}) = \sum_{j=1}^J \lambda_j(f) \mathcal{N}(\Phi_j, \mathbf{x}), \quad (4.10)$$

where Φ_f is the deep ReLU neural network in Theorem 4.5. From Lemma 4.4 we can find an absolute positive constant K_5 so that $N_w(\Phi'_f) \leq K_5 d$ and a positive constant $K_6 = K_6(\alpha, \beta)$ such that

$$L(\Phi'_f) \leq \sum_{j=1}^J L(\Phi_j) \leq C |D_\beta^d(m)| \log d \log(d \delta^{-1}) \leq K_6 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}),$$

where the last inequality follows from (4.6) and (4.7). Hence, by Lemma 4.2 and (4.10), Φ'_f generates a deep ReLU neural network Φ_f^* such that $\mathcal{N}(\Phi_f^*, \mathbf{x}) = \mathcal{N}(\Phi_f, \mathbf{x})$ and, consequently, there hold (4.8) and

$$N_w(\Phi_f^*) \leq K_5 d \quad \text{and} \quad L(\Phi_f^*) \leq K_6 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1}).$$

The proof of existence of an architecture $\mathbb{A}_\varepsilon^*$ of Φ_f^* satisfying (4.9) is similar to the proof of existence of $\mathbb{A}_\varepsilon$ at the end of the proof of Theorem 4.5. $\square$

In the case when the error is measured in the norm of the space $W_0^{1,\infty}$, we are able to give dimension-dependent lower bounds for the numbers of weights of deep ReLU networks whose outputs approximate functions from $\dot{U}_\infty^\alpha$ with a given accuracy. More precisely, we have the following results.

Theorem 4.8 *Let $d \in \mathbb{N}, d \geq 2$, and $1 < \alpha \leq 2$. Let $\varepsilon \in (0, 24^{-d})$ and $\mathbb{A}$ be a neural network architecture such that for any $f \in \dot{U}_\infty^\alpha$, there is a deep ReLU neural network Φ_f having the architecture $\mathbb{A}$ and*

$$\|f - \mathcal{N}(\Phi_f, \cdot)\|_{W_0^{1,\infty}} \leq \varepsilon.$$

Then there is a positive constant $K_7 = K_7(\alpha)$ such that

$$W(\mathbb{A}) \geq K_7 24^{-\frac{d}{2(\alpha-1)}} \varepsilon^{-\frac{1}{2(\alpha-1)}}.$$

If assume in addition that

$$L(\mathbb{A}) \leq C(\log \varepsilon^{-1})^\lambda$$

for some constant $C > 0$ and $\lambda \geq 0$, then there exists a constant $K_8 = K_8(\alpha) > 0$ such that

$$W(\mathbb{A}) \geq K_8 24^{-\frac{d}{\alpha-1}} \varepsilon^{-\frac{1}{\alpha-1}} (\log \varepsilon^{-1})^{-2\lambda-1}.$$

Some non-dimension-dependent lower bounds have been obtained in [42, 18] for isotropic Sobolev spaces. In order to prove the above theorem we develop some techniques in [42, 18] which are relied on upper bounds for VC-dimension of ReLU networks with Boolean outputs [1]. We start with recalling a definition of VC-dimension.

Definition 4.9 Let H be a set of functions $h : X \rightarrow \{0, 1\}$ for some set X . Then the VC-dimension of H , denoted by $\text{VCdim}(H)$, is defined as the supremum of all number m that there exist $x^1, \dots, x^m \in X$ such that for any sequence $(y_j)_{j=1}^m \in \{0, 1\}^m$ there is a function $h \in H$ with $h(x^j) = y_j$ for $j = 1, \dots, m$.

The relation between VC-dimension and number of weights and depth of a neural network architecture is given in the following lemma, see [1, Theorems 8.7 and 8.8].

Lemma 4.5 Let $\mathbb{A}$ be a deep neural network architecture. Let F be the class of functions-outputs $f = \mathcal{N}(\Phi, \cdot)$ of all deep ReLU neural networks Φ having the architecture $\mathbb{A}$. Let $a \in \mathbb{R}$ and H be the class of all functions $h_f : \mathbb{I}^d \rightarrow \{0, 1\}$, $f \in F$, defined by threshold:

$$h_f(\mathbf{x}) := \begin{cases} 1 & \text{if } f(\mathbf{x}) > a \\ 0 & \text{if } f(\mathbf{x}) \leq a. \end{cases}$$

Then

$$\text{VCdim}(H) \leq CW(\mathbb{A})^2$$

for some positive constant C . Moreover, there exists a $C' > 0$ such that

$$\text{VCdim}(H) \leq C' L(\mathbb{A})^2 W(\mathbb{A}) \log(W(\mathbb{A})).$$

The following elementary property of deep ReLU neural networks has been proven in [18, Lemma D.1] which is based on the piece-wise linearity of ReLU activation function.

Lemma 4.6 Let Φ be a deep ReLU neural network, $\boldsymbol{\xi} \in (0, 1)^d$, and $\boldsymbol{\nu} \in \mathbb{R}^d$. Then there exists an open set $G = G(\boldsymbol{\xi}, \boldsymbol{\nu}) \subset (0, 1)^d$ and $\delta = \delta(\boldsymbol{\xi}, \boldsymbol{\nu}) > 0$ such that $\boldsymbol{\xi} + \lambda \delta \boldsymbol{\nu} \in \overline{G}$ for $\lambda \in [0, 1]$ and $\mathcal{N}(\Phi, \cdot)$ is affine on G .

We are now in position to prove Theorem 4.8.

Proof. Given $\varepsilon \in (0, 24^{-d})$ and $m = m(\varepsilon) \in \mathbb{N}$ which will be chosen later. Denote $L = L(\mathbb{A})$, $W = W(\mathbb{A})$ the depth and the number of weights of $\mathbb{A}$. We assume that

$$\mathbb{A} = ((\mathbf{W}^1, \mathbf{b}^1), \dots, (\mathbf{W}^L, \mathbf{b}^L)),$$

where $\mathbf{W}^\ell$ is an $N_\ell \times N_{\ell-1}$ matrix, and $\mathbf{b}^\ell \in \mathbb{R}^{N_\ell}$ with $N_0 = d$, $N_L = 1$, and $N_1, \dots, N_{L-1} \in \mathbb{N}$. Let Φ be a deep ReLU neural network of the architecture $\mathbb{A}$. For $\mathbf{x} = (x_1, \dots, x_d) \in \mathbb{I}^d$ and $\delta > 0$ we put $\bar{\mathbf{x}} = (x_1 + \frac{2^{-m-2}}{3}, x_2, \dots, x_d)$ and define the deep ReLU neural network Φ^δ by parallelization construction in Lemma 4.1 with output

$$\mathcal{N}(\Phi^\delta, \mathbf{x}) := \frac{\mathcal{N}(\Phi, \bar{\mathbf{x}}) - \mathcal{N}(\Phi, \bar{\mathbf{x}} - \delta \mathbf{e}^1)}{\delta}, \quad \mathbf{e}^1 = (1, 0, \dots, 0) \in \mathbb{R}^d.$$

Then Φ^δ is a deep ReLU neural network having the architecture

$$\tilde{\mathbb{A}} = \left(\left(\begin{bmatrix} \mathbf{I}_d \\ \mathbf{I}_d \end{bmatrix}, \begin{bmatrix} \mathbf{e}^1 \\ \mathbf{e}^1 \end{bmatrix} \right), \left(\begin{bmatrix} \mathbf{W}^1 & 0 \\ 0 & \mathbf{W}^1 \end{bmatrix}, \begin{bmatrix} \mathbf{b}^1 \\ \mathbf{b}^1 \end{bmatrix} \right), \dots, \left(\begin{bmatrix} \mathbf{W}^{L-1} & 0 \\ 0 & \mathbf{W}^{L-1} \end{bmatrix}, \begin{bmatrix} \mathbf{b}^{L-1} \\ \mathbf{b}^{L-1} \end{bmatrix} \right) ([\mathbf{W}^L, \mathbf{W}^L], \mathbf{b}^L) \right),$$

where $\mathbf{I}_d$ is the identity matrix of size d . It is clear that $\tilde{\mathbb{A}}$ has depth and number of weights

$$\tilde{L} := L(\tilde{\mathbb{A}}) = L + 1, \quad \text{and} \quad \tilde{W} := W(\tilde{\mathbb{A}}) = 2W + 2d + 1. \quad (4.11)$$

Let $a = a(m, d) \in \mathbb{R}$ be a constant which will be clarified later. For a deep ReLU neural network $\tilde{\Phi}$ having architecture $\tilde{\mathbb{A}}$, we define the function

$$h(\tilde{\Phi}, \mathbf{x}) = \begin{cases} 1 & \text{if } \mathcal{N}(\tilde{\Phi}, \mathbf{x}) > a/2, \\ 0 & \text{if } \mathcal{N}(\tilde{\Phi}, \mathbf{x}) \leq a/2, \end{cases}$$

and the set

$$H := H(\tilde{\mathbb{A}}) = \{h(\tilde{\Phi}, \mathbf{x}) : \tilde{\Phi} \text{ is a deep ReLU neural network having architecture } \tilde{\mathbb{A}}\}.$$

In the following we will show that $\text{VCdim}(H) \geq 2^m$. For this we put

$$\mathbf{x}^j = \left(2^{-m} \left(j - \frac{3}{4} \right), \frac{1}{2}, \dots, \frac{1}{2} \right) \in (0, 1)^d, \quad j = 1, \dots, 2^m.$$

It is clear that $\bar{\mathbf{x}}^j \in (0, 1)^d$ for $j = 1, \dots, 2^m$. For every $\mathbf{y} = (y_1, \dots, y_{2^m}) \in \{0, 1\}^{2^m}$, we define

$$f_{\mathbf{y}}(\mathbf{x}) = 18^{-d} 2^{-\alpha m} \left(\sum_{j=1}^{2^m} y_j \psi_{m,j}(x_1) \right) \prod_{\ell=2}^d \psi_{0,1}(x_\ell),$$

which is similar to $f_m(\mathbf{x})$ in the proof of Theorem 3.5. Then it holds $f_{\mathbf{y}} \in \mathring{U}_\infty^\alpha$. Moreover, we have

$$\frac{\partial f_{\mathbf{y}}}{\partial x_1}(\bar{\mathbf{x}}^j) = 18^{-d} 2^{-\alpha m} y_j \psi'_{m,j} \left(x_1^j + \frac{2^{-m-2}}{3} \right) \prod_{\ell=2}^d \psi(x_\ell^j) = y_j 18^{-d} 2^{-\alpha m} \psi \left(\frac{1}{2} \right)^{(d-1)} 2^m \psi' \left(\frac{1}{3} \right).$$

From $\psi(\frac{1}{2}) = M_3(\frac{3}{2}) = \frac{3}{4}$ and $\psi'(\frac{1}{3}) = 3M'_3(1) = 3$ we get

$$\frac{\partial f_{\mathbf{y}}}{\partial x_1}(\bar{\mathbf{x}}^j) = 4y_j 18^{-d} 2^{-(\alpha-1)m} \left(\frac{3}{4} \right)^d = 4y_j 2^{-(\alpha-1)m} 24^{-d}.$$

Since $f_{\mathbf{y}} \in \mathring{U}_\infty^\alpha$, by the assumption, there exists a neural network $\Phi_{\mathbf{y}}$ having architecture $\mathbb{A}$ such that

$$\|f_{\mathbf{y}} - \mathcal{N}(\Phi_{\mathbf{y}}, \cdot)\|_{W_0^{1,\infty}} \leq \varepsilon.$$

By Lemma 4.6, for $\bar{\mathbf{x}}^j \in (0, 1)^d$, there exists an open set $G_j \subset (0, 1)^d$ and $\delta_j > 0$ such that $\bar{\mathbf{x}}^j + \lambda \delta_j \mathbf{e}^1 \in \bar{G}_j$ for $\lambda \in [0, 1]$ and $\mathcal{N}(\Phi_{\mathbf{y}}, \cdot)$ is affine on $\bar{G}_j$. Let

$$\delta^* = \min_{1 \leq j \leq 2^m} \delta_j > 0.$$

and $B(\bar{\mathbf{x}}^j, \theta_j)$ be the open ball centered at $\bar{\mathbf{x}}^j$ with radius θ_j . Then we have

$$\mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) = \frac{\mathcal{N}(\Phi_{\mathbf{y}}, \bar{\mathbf{x}}^j) - \mathcal{N}(\Phi_{\mathbf{y}}, \bar{\mathbf{x}}^j - \delta^* \mathbf{e}^1)}{\delta^*} = \frac{\partial \mathcal{N}(\Phi_{\mathbf{y}}, \xi^j)}{\partial x_1},$$

for some $\xi^j \in B(\bar{\mathbf{x}}^j, \theta_j) \cap G_j$. In case $\bar{\mathbf{x}}^j \in G_j$ we can choose $\xi^j = \bar{\mathbf{x}}^j$. Define m the largest positive integer such that

$$\varepsilon \leq 2^{-(\alpha-1)m} 24^{-d}.$$

We then obtain for $y_j = 1$,

$$\begin{aligned} \mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) &= \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} + \frac{\partial \mathcal{N}(\Phi_{\mathbf{y}}, \xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} \\ &\geq \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} - \left| \frac{\partial \mathcal{N}(\Phi_{\mathbf{y}}, \xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} \right| - \left| \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} \right|. \end{aligned}$$

Since $\frac{\partial f_{\mathbf{y}}(\mathbf{x})}{\partial x_1}$ is a continuous function due to $f_{\mathbf{y}} \in \dot{U}_{\infty}^{\alpha}$, $\alpha > 1$. Then we can choose θ_j small enough such that

$$\left| \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} \right| \leq \frac{\varepsilon}{2}$$

which implies

$$\mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) \geq 4 \cdot 2^{-(\alpha-1)m} 24^{-d} - \frac{3}{2} \varepsilon \geq \frac{5}{2} \cdot 2^{-(\alpha-1)m} 24^{-d}.$$

For $y_j = 0$,

$$\begin{aligned} |\mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j)| &= \left| \frac{\partial \mathcal{N}(\Phi_{\mathbf{y}}, \xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} + \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} \right| \\ &\leq \left| \frac{\partial \mathcal{N}(\Phi_{\mathbf{y}}, \xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} \right| + \left| \frac{\partial f_{\mathbf{y}}(\xi^j)}{\partial x_1} - \frac{\partial f_{\mathbf{y}}(\bar{\mathbf{x}}^j)}{\partial x_1} \right| \leq \frac{3}{2} \varepsilon \leq \frac{3}{2} \cdot 2^{-(\alpha-1)m} 24^{-d}. \end{aligned}$$

Putting $a = 4 \cdot 2^{-(\alpha-1)m} 24^{-d}$, we get

$$\begin{cases} \mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) > a/2 & \text{if } y_j = 1, \\ \mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) \leq a/2 & \text{if } y_j = 0. \end{cases}$$

Then for $\mathbf{y} \in \{0, 1\}^{2^m}$ the function

$$h(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}) = \begin{cases} 1 & \text{if } \mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}) > a/2, \\ 0 & \text{if } \mathcal{N}(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}) \leq a/2, \end{cases}$$

belongs to H and satisfies $h(\Phi_{\mathbf{y}}^{\delta^*}, \mathbf{x}^j) = y_j$. By definition of VC-dimension, see Definition 4.9, we obtain $2^m \leq \text{VCdim}(H)$. Moreover, from (4.11) and Lemma 4.5 we have

$$2^m \leq \text{VCdim}(H) \leq C(2W + 2d + 1)^2.$$

Since $\mathbb{A}$ has input dimension d and depth $L \geq 2$, we find that $2d + 1 \leq 2W$. From this and $2^{-(\alpha-1)(m+1)}24^{-d} \leq \varepsilon$ we get

$$C4^2W^2 \geq 2^m \geq \frac{1}{2}24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}$$

or

$$W \geq 2^m \geq \frac{1}{4\sqrt{2}C}24^{-\frac{d}{2(\alpha-1)}}\varepsilon^{-\frac{1}{2(\alpha-1)}}$$

which is the first statement.

Concerning second one, we have

$$\begin{aligned} \frac{1}{2}24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}} &\leq C'\tilde{L}^2\tilde{W}\log\tilde{W} = C'(L+1)^2(2W+2d+1)\log(2W+2d+1) \\ &\leq C''(\log\varepsilon^{-1})^{2\lambda}W\log W \end{aligned}$$

which implies

$$W\log W \geq \frac{1}{2}24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}(\log\varepsilon^{-1})^{-2\lambda}.$$

Consider for $\kappa < 1$,

$$W_\kappa = \kappa 24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}\log(\varepsilon^{-1})^{-2\lambda-1}.$$

Then we have that

$$\log W_\kappa = \log \kappa - \frac{d\log 24}{\alpha-1} - (2\lambda+1)\log(\log\varepsilon^{-1}) + \frac{1}{\alpha-1}(\log\varepsilon^{-1}) \leq \frac{1}{\alpha-1}(\log\varepsilon^{-1}).$$

From this we obtain

$$W_\kappa \log W_\kappa \leq \frac{\kappa}{\alpha-1}24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}(\log\varepsilon^{-1})^{-2\lambda} \leq \frac{1}{2C''}24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}(\log\varepsilon^{-1})^{-2\lambda} \leq W\log(W)$$

if we choose $\frac{\kappa}{\alpha-1} \leq \min\{1, \frac{1}{2C''}\}$. Consequently, we get

$$\kappa 24^{-\frac{d}{\alpha-1}}\varepsilon^{-\frac{1}{\alpha-1}}(\log\varepsilon^{-1})^{-2\lambda-1} = W_\kappa \leq W$$

which is the second statement. The proof is completed. $\square$

5 Concluding remarks

We have constructed interpolation linear algorithms of sampling recovery on sparse grids of points which are tailored for approximation in the norm of the isotropic Sobolev space $W_0^{1,p}$ of functions from the Hölder-Zygmund space of mixed smoothness $H_\infty^\alpha(\mathbb{I}^d)$, $\alpha \in (1, 2]$. These grids have a definite advantage over the standard grids and classical Smolyak sparse grids. They are in some sense optimal in terms of sampling n -widths r_n which characterizes the best approximation error by linear algorithms of sampling recovery of functions f from the unit ball $\mathring{U}_\infty^\alpha$ based on n sampled values of f . Moreover, with some restrictions we proved the tight dimension-dependent estimates $C_1B_*^{-d}n^{-(\alpha-1)} \leq r_n \leq C_2B^{*-d}n^{-(\alpha-1)}$ with $B_* > 1$ and $B^* > 1$, which show that the sampling n -widths are decreasing as fast as the exponent B^{*-d} when the dimension d going to ∞ .

Based on the constructed linear algorithms of sparse-grid sampling recovery, we have explicitly constructed a deep ReLU neural network Φ_f having an output that approximates in $W_0^{1,p}$ -norm any function $f \in \mathring{U}_\infty^\alpha$ with a prescribed accuracy ε and established a dimension-dependent estimate for

the computation complexity characterized by number of weights $W(\Phi_f)$ and the depth $L(\Phi_f)$ of this deep ReLU neural network: $L(\Phi_f) \leq C_3 \log d \log(\varepsilon^{-1})$ and $W(\Phi_f) \leq C_4 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1})$ with $B_2 > 1$. Moreover, this output can be designed also as an output of another “very” deep ReLU neural network Φ_f^* having computation complexity expressing the ε -independent width $N_w(\Phi_f^*) \leq C_5 d$ and the dimension-dependent depth $L(\Phi_f^*) \leq C_6 B_2^{-d} (\varepsilon^{-1})^{\frac{1}{\alpha-1}} \log(\varepsilon^{-1})$. This shows in particular, that the computation complexity is decreasing as fast as the exponent B_2^{-d} when the dimension d going to ∞ . In the case when $p = \infty$, we gave dimension-dependent lower bounds for the numbers of weights of deep ReLU networks whose outputs approximate functions from $\dot{U}_\infty^\alpha$ with a given accuracy.

Thus, we have shown that under some reasonable restrictions the curse of dimensionality can be avoided in the both approximations by sparse-grid sampling recovery and by deep ReLU neural networks. The result also indicated that the decomposition of functions from the Hölder-Zygmund space of mixed smoothness $H_\infty^\alpha(\mathbb{I}^d)$ into tensor product Faber series plays a fundamental role in construction of linear algorithms of sparse-grid sampling recovery and of deep ReLU neural networks for approximation of functions from $H_\infty^\alpha(\mathbb{I}^d)$. In the present paper, our concerns are non-adaptive approximations by sparse-grid sampling recovery and by deep ReLU neural networks for which algorithms of sampling recovery are linear and the architecture of deep ReLU neural networks is the same for all functions. In a forthcoming paper, we will discuss a problem of adaptive nonlinear approximation by sparse-grid sampling recovery and by deep ReLU neural networks of multivariate functions having a mixed smoothness.

Acknowledgments. This work is funded by Vietnam National Foundation for Science and Technology Development (NAFOSTED) under Grant No. 102.01-2020.03. A part of this work was done when the authors were working at the Vietnam Institute for Advanced Study in Mathematics (VIASM). They would like to thank the VIASM for providing a fruitful research environment and working condition.

References

- [1] M. Anthony and P. Bartlett. *Neural Network Learning: Theoretical Foundations*. Cambridge University Press, Cambridge, 2009.
- [2] R. Arora, A. Basu, P. Mianjy, and A. Mukherjee. Understanding deep neural networks with rectified linear units. *Electronic Colloquium on Computational Complexity*, Report No. 98, 2017.
- [3] R. Bellmann. *Dynamic Programming*. Princeton University Press, Princeton, 1957.
- [4] H.-J. Bungartz and M. Griebel. Sparse grids. *Acta Numer.*, 13:147–269, 2004.
- [5] G. Byrenheid, D. Dũng, W. Sickel, and T. Ullrich. Sampling on energy-norm based sparse grids for the optimal recovery of Sobolev type functions in H_γ . *J. Approx. Theory*, 207:207–231, 2016.
- [6] G. Cybenko. Approximation by superpositions of a sigmoidal function. *Math. Control Signal*, 2:303–314, 1989.
- [7] D. Dũng. B-spline quasi-interpolant representations and sampling recovery of functions with mixed smoothness. *J. Complexity*, 27:541–567, 2011.
- [8] D. Dũng. Optimal adaptive sampling recovery. *Adv. Comput. Math*, 34:1–41, 2011.

- [9] D. Dũng. Sampling and cubature on sparse grids based on a B-spline quasi-interpolation. *Found. Comp. Math.*, 16:1193–1240, 2016.
- [10] D. Dũng, V. N. Temlyakov, and T. Ullrich. *Hyperbolic Cross Approximation*. Advanced Courses in Mathematics - CRM Barcelona, Birkhäuser/Springer, 2018.
- [11] D. Dũng and M. X. Thao. Dimension-dependent error estimates for sampling recovery on Smolyak grids based on B-spline quasi-interpolation. *J. Approx. Theory*, 250:185–205, 2020.
- [12] I. Daubechies, R. DeVore, S. Foucart, B. Hanin, and G. Petrova. Nonlinear approximation and (Deep) ReLU networks. *arXiv:1905.02199*, 2019.
- [13] R. DeVore and G. Lorentz. *Constructive Approximation*. Springer-Verlag, New York, 1993.
- [14] W. E and Q. Wang. Exponential convergence of the deep neural network approximation for analytic functions. *Sci. China Math.*, 61:1733–1740, 2018.
- [15] G. Faber. Über stetige Funktionen. *Math. Ann.*, 66:81–94, 1909.
- [16] A. Griewank, F. Y. Kuo, H. Leövey, and I. H. Sloan. High dimensional integration of kinks and jumps – smoothing by preintegration. *J. Comput. Appl. Math.*, 344:259–274, 2018.
- [17] P. Grohs, D. Perekrestenko, D. Elbrachter, and H. Bolcskei. Deep neural network approximation theory. *arXiv: 1901.02220*, 2019.
- [18] I. Gühring, G. Kutyniok, and P. Petersen. Error bounds for approximations with deep ReLU neural networks in $W^{s,p}$ norms. <https://doi.org/10.1142/S0219530519410021>, 00:00, 2020.
- [19] D. Hebb. *The Organization of Behavior: A Neuropsychological Theory*. Wiley, 1949.
- [20] K. Hornik, M. Stinchcombe, and H. White. Multilayer feedforward networks are universal approximators. *Neural Netw.*, 2:359–366, 1989.
- [21] A. Krizhevsky, I. Sutskever, and G. E. Hinton. ImageNet classification with deep convolutional neural networks. *NeurIPS*, pages 1106–1114, 2012.
- [22] Y. LeCun, Y. Bengio, and G. Hinton. Deep learning. *Nature*, 521:436–444, 2015.
- [23] H. Mhaskar and T. Poggio. Deep vs. shallow networks: An approximation theory perspective. *Anal. Appl.*, 14:829–848, 2016.
- [24] H. N. Mhaskar. Neural networks for optimal approximation of smooth and analytic functions. *Neural Comput.*, 8:164–177, 1996.
- [25] H. Montanelli and Q. Du. New error bounds for deep ReLU networks using sparse grids. *SIAM J. Math. Data Sci.*, 1:78–92, 2019.
- [26] G. Montúfar, R. Pascanu, K. Cho, and Y. Bengio. On the number of linear regions of deep neural networks. In *Advances in neural information processing systems*, pages 2924–2932, 2014.
- [27] E. Novak and H. Woźniakowski. *Tractability of Multivariate Problems, Volume I: Linear Information*. EMS Tracts in Mathematics, Vol. 6, Eur. Math. Soc. Publ. House, Zürich, 2008.
- [28] E. Novak and H. Woźniakowski. *Tractability of Multivariate Problems, Volume II: Standard Information for Functionals*. EMS Tracts in Mathematics, Vol. 12, Eur. Math. Soc. Publ. House, Zürich, 2010.

- [29] J. A. A. Opschoor, P. C. Petersen, and C. Schwab. Deep ReLU networks and high-order finite element methods. *Anal. Appl. (Singap.)*, <https://doi.org/10.1142/S0219530519410136>, 2020.
- [30] J. A. A. Opschoor, C. Schwab, and J. Zech. Exponential ReLU DNN expression of holomorphic maps in high dimension. *SAM, Research Report No. 2019-35*, 2019.
- [31] P. Petersen and F. Voigtlaender. Optimal approximation of piecewise smooth functions using deep ReLU neural networks. *Neural Netw.*, 108:296–330, 2018.
- [32] A. Pinkus. Approximation theory of the MLP model in neural networks. *Acta Numer.*, 8:143–195, 1999.
- [33] F. Rosenblatt. The perceptron: a probabilistic model for information storage and organization in the brain. *Psychol. Rev.*, 65:386–408, 1958.
- [34] C. Schwab and J. Zech. Deep learning in high dimension: Neural network expression rates for generalized polynomial chaos expansions in UQ. *Anal. Appl. (Singap.)*, 17:19–55, 2019.
- [35] S. Smolyak. Quadrature and interpolation formulas for tensor products of certain classes of functions. *Dokl. Akad. Nauk*, 148:1042–1045, 1963.
- [36] T. Suzuki. Adaptivity of deep ReLU network for learning in Besov and mixed smooth Besov spaces: optimal rate and curse of dimensionality. *International Conference on Learning Representations*, 2019.
- [37] M. Telgarsky. Representation benefits of deep feedforward networks. *arXiv:1509.08101*, 2015.
- [38] M. Telgrasky. Benefits of depth in neural nets. In *Proceedings of the JMLR: Workshop and Conference Proceedings, New York, NY, USA*, 49:1–23, 2016.
- [39] H. Triebel. *Bases in Function Spaces, Sampling, Discrepancy, Numerical Integration*. European Math. Soc. Publishing House, Zürich, 2010.
- [40] H. Triebel. *Hybrid Function Spaces, Heat and Navier-Stokes Equations*. European Mathematical Society, 2015.
- [41] Y. Wu, M. Schuster, Z. Chen, Q. V. Le, and M. Norouzi. Googles neural machine translation system: Bridging the gap between human and machine translation. *arXiv: 1609.08144*, 2016.
- [42] D. Yarotsky. Error bounds for approximations with deep ReLU networks. *Neural Netw.*, 94:103–114, 2017.
- [43] D. Yarotsky. Quantified advantage of discontinuous weight selection in approximations with deep neural networks. *arXiv: 1705.01365*, 2017.
- [44] H. Yserentant. *Regularity and Approximability of Electronic Wave Functions*. Lecture Notes in Mathematics, Springer, 2010.
- [45] C. Zenger. Sparse grids. In *Parallel Algorithms for Partial Differential Equations. Volume 31 of Notes on Numerical Fluid Mechanics (Vieweg, Wiesbaden 1991)*, pages 241–251.