

# Besov Function Approximation and Binary Classification on Low-Dimensional Manifolds Using Convolutional Residual Networks

Hao Liu<sup>1</sup> Minshuo Chen<sup>2</sup> Tuo Zhao<sup>2</sup> Wenjing Liao<sup>3</sup>

## Abstract

Most of existing statistical theories on deep neural networks have sample complexities cursed by the data dimension and therefore cannot well explain the empirical success of deep learning on high-dimensional data. To bridge this gap, we propose to exploit low-dimensional geometric structures of the real world data sets. We establish theoretical guarantees of convolutional residual networks (ConvResNet) in terms of function approximation and statistical estimation for binary classification. Specifically, given the data lying on a  $d$ -dimensional manifold isometrically embedded in  $\mathbb{R}^D$ , we prove that if the network architecture is properly chosen, ConvResNets can (1) approximate *Besov functions* on manifolds with arbitrary accuracy, and (2) learn a classifier by minimizing the empirical logistic risk, which gives an *excess risk* in the order of  $n^{-\frac{s}{2s+2(s\vee d)}}$ , where  $s$  is a smoothness parameter. This implies that the sample complexity depends on the intrinsic dimension  $d$ , instead of the data dimension  $D$ . Our results demonstrate that ConvResNets are adaptive to low-dimensional structures of data sets.

## 1. Introduction

Deep learning has achieved significant success in various practical applications with high-dimensional data set, such as computer vision (Krizhevsky et al., 2012), natural language processing (Graves et al., 2013; Young et al., 2018; Wu et al., 2016), health care (Miotto et al., 2018; Jiang et al., 2017) and bioinformatics (Alipanahi et al., 2015; Zhou & Troyanskaya, 2015).

<sup>1</sup>Department of Mathematics, Hong Kong Baptist University, Kowloon Tong, Hong Kong. <sup>2</sup>School of Industrial and Systems Engineering, Georgia Institute of Technology, Atlanta, GA 30332 USA. <sup>3</sup>School of Mathematics, Georgia Institute of Technology, Atlanta, GA 30332 USA.. Correspondence to: Wenjing Liao <wliao60@gatech.edu>.

The success of deep learning clearly demonstrates the great power of neural networks in representing complex data. In the past decades, the representation power of neural networks has been extensively studied. The most commonly studied architecture is the feedforward neural network (FNN), as it has a simple composition form. The representation theory of FNNs has been developed with smooth activation functions (e.g., sigmoid) in Cybenko (1989); Barron (1993); McCaffrey & Gallant (1994); Hamers & Kohler (2006); Kohler & Krzyżak (2005); Kohler & Mehnert (2011) or nonsmooth activations (e.g., ReLU) in Lu et al. (2017); Yarotsky (2017); Lee et al. (2017); Suzuki (2019). These works show that if the network architecture is properly chosen, FNNs can approximate uniformly smooth functions (e.g., Hölder or Sobolev) with arbitrary accuracy.

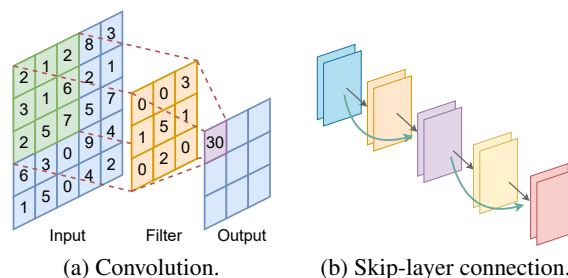


Figure 1. Illustration of (a) convolution and (b) skip-layer connection.

In real-world applications, convolutional neural networks (CNNs) are more popular than FNNs (LeCun et al., 1989; Krizhevsky et al., 2012; Sermanet et al., 2013; He et al., 2016; Chen et al., 2017; Long et al., 2015; Simonyan & Zisserman, 2014; Girshick, 2015). In a CNN, each layer consists of several filters (channels) which are convolved with the input, as demonstrated in Figure 1(a). Due to such complexity in the CNN architecture, there are limited works on the representation theory of CNNs (Zhou, 2020b;a; Fang et al., 2020; Petersen & Voigtlaender, 2020). The constructed CNNs in these works become extremely wide (in terms of the size of each layer’s output) as the approximation error goes to 0. In most real-life applications, the network width does not exceed 2048 (Zagoruyko & Komodakis, 2016; Zhang et al., 2020).

Convolutional residual networks (ConvResNet) is a special

CNN architecture with skip-layer connections, as shown in Figure 1(b). Specifically, in addition to CNNs, ConvResNets have identity connections between inconsecutive layers. In many applications, ConvResNets outperform CNNs in terms of generalization performance and computational efficiency, and alleviate the vanishing gradient issue. Using this architecture, He et al. (2016) won the 1st place on the ImageNet classification task with a 3.57% top 5 error in 2015.

Recently, Oono & Suzuki (2019) develops the only representation and statistical estimation theory of ConvResNets. Oono & Suzuki (2019) proves that if the network architecture is properly set, ConvResNets with a fixed filter size and a fixed number of channels can universally approximate Hölder functions with arbitrary accuracy. However, the sample complexity in Oono & Suzuki (2019) grows exponentially with respect to the data dimension and therefore cannot well explain the empirical success of ConvResNets for high dimensional data. In order to estimate a  $C^s$  function in  $\mathbb{R}^D$  with accuracy  $\varepsilon$ , the sample size required by Oono & Suzuki (2019) scales as  $\varepsilon^{-\frac{2s+D}{s}}$ , which is far beyond the sample size used in practical applications. For example, the ImageNet data set consists of 1.2 million labeled images of size  $224 \times 224 \times 3$ . According to this theory, to achieve a 0.1 error, the sample size is required to be in the order of  $10^{224 \times 224 \times 3}$  which greatly exceeds 1.2 million. Due to the curse of dimensionality, there is a huge gap between theory and practice.

We bridge such a gap by taking low-dimensional geometric structures of data sets into consideration. It is commonly believed that real world data sets exhibit low-dimensional structures due to rich local regularities, global symmetries, or repetitive patterns (Hinton & Salakhutdinov, 2006; Osher et al., 2017; Tenenbaum et al., 2000). For example, the ImageNet data set contains many images of the same object with certain transformations, such as rotation, translation, projection and skeletonization. As a result, the degree of freedom of the ImageNet data set is significantly smaller than the data dimension (Gong et al., 2019).

The function space considered in Oono & Suzuki (2019) is the Hölder space in which functions are required to be differentiable everywhere up to certain order. In practice, the target function may not have high order derivatives. Function spaces with less restrictive conditions are more desirable. In this paper, we consider the Besov space  $B_{p,q}^s$ , which is more general than the Hölder space. In particular, the Hölder and Sobolev spaces are special cases of the Besov space:

$$W^{s+\alpha,\infty} = \mathcal{H}^{s,\alpha} \subseteq B_{\infty,\infty}^{s+\alpha} \subseteq B_{p,q}^{s+\alpha}$$

for any  $0 < p, q \leq \infty, s \in \mathbb{N}$  and  $\alpha \in (0, 1]$ . For practical applications, it has been demonstrated in image processing that Besov norms can capture important features, such

as edges (Jaffard et al., 2001). Due to the generality of the Besov space, it is shown in Suzuki & Nitanda (2019) that kernel ridge estimators have a sub-optimal rate when estimating Besov functions.

In this paper, we establish theoretical guarantees of ConvResNets for the approximation of Besov functions on a low-dimensional manifold, and a statistical theory on binary classification. Let  $\mathcal{M}$  be a  $d$ -dimensional compact Riemannian manifold isometrically embedded in  $\mathbb{R}^D$ . Denote the Besov space on  $\mathcal{M}$  as  $B_{p,q}^s(\mathcal{M})$  for  $0 < p, q \leq \infty$  and  $0 < s < \infty$ . Our function approximation theory is established for  $B_{p,q}^s(\mathcal{M})$ . For binary classification, we are given  $n$  i.i.d. samples  $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$  where  $\mathbf{x}_i \in \mathcal{M}$  and  $y_i \in \{-1, 1\}$  is the label. The label  $y$  follows the Bernoulli-type distribution

$$\mathbb{P}(y = 1|\mathbf{x}) = \eta(\mathbf{x}), \mathbb{P}(y = -1|\mathbf{x}) = 1 - \eta(\mathbf{x})$$

for some  $\eta : \mathcal{M} \rightarrow [0, 1]$ . Our results (Theorem 1 and 2) are summarized as follows:

**Theorem** (informal). Assume  $s \geq d/p + 1$ .

1. Given  $\varepsilon \in (0, 1)$ , we construct a ConvResNet architecture such that, for any  $f^* \in B_{p,q}^s(\mathcal{M})$ , if the weight parameters of this ConvResNet are properly chosen, it gives rises to  $\bar{f}$  satisfying

$$\|\bar{f} - f^*\|_{L^\infty} \leq \varepsilon.$$

2. Assume  $\eta \in B_{p,q}^s(\mathcal{M})$ . Let  $f_\phi^*$  be the minimizer of the population logistic risk. If the ConvResNet architecture is properly chosen, minimizing the empirical logistic risk gives rise to  $\hat{f}_{\phi,n}$  with the following excess risk bound

$$\mathbb{E}(\mathcal{E}_\phi(\hat{f}_{\phi,n}, f_\phi^*)) \leq Cn^{-\frac{s}{2s+2(s \vee d)}} \log^4 n,$$

where  $\mathcal{E}_\phi(\hat{f}_{\phi,n}, f_\phi^*)$  denotes the excess logistic risk of  $\hat{f}_{\phi,n}$  against  $f_\phi^*$  and  $C$  is a constant independent of  $n$ .

We remark that the first part of the theorem above requires the network size to depend on the intrinsic dimension  $d$  and only weakly depend on  $D$ . The second part is built upon the first part and shows a fast convergence rate of the excess risk in terms of  $n$  where the exponent depends on  $d$  instead of  $D$ . Our results demonstrate that ConvResNets are adaptive to low-dimensional structures of data sets.

**Related work.** Approximation theories of FNNs with the ReLU activation have been established for Sobolev (Yarotsky, 2017), Hölder (Schmidt-Hieber, 2017) and Besov (Suzuki, 2019) spaces. The networks in these works have certain cardinality constraint, i.e., the number of nonzero parameters is bounded by certain constant, which requires a lot of efforts for training.

Approximation theories of CNNs are developed in Zhou (2020b); Petersen & Voigtlaender (2020); Oono & Suzuki (2019). Among these works, Zhou (2020b) shows that

Table 1. Comparison of our approximation theory and existing theoretical results.

|                                | Network type | Function class | Low dim. structure | Fixed width  | Training                                             |
|--------------------------------|--------------|----------------|--------------------|--------------|------------------------------------------------------|
| Yarotsky (2017)                | FNN          | Sobolev        | $\times$           | $\times$     | difficult to train due to the cardinality constraint |
| Suzuki (2019)                  | FNN          | Besov          | $\times$           | $\times$     |                                                      |
| Chen et al. (2019a)            | FNN          | Hölder         | $\checkmark$       | $\times$     |                                                      |
| Petersen & Voigtlaender (2020) | CNN          | FNN            | $\times$           | $\times$     |                                                      |
| Zhou (2020b)                   | CNN          | Sobolev        | $\times$           | $\times$     | can be trained without the cardinality constraint    |
| Oono & Suzuki (2019)           | ConvResNet   | Hölder         | $\times$           | $\checkmark$ |                                                      |
| Ours                           | ConvResNet   | Besov          | $\checkmark$       | $\checkmark$ |                                                      |

CNNs can approximate Sobolev functions in  $W^{s,2}$  for  $s \geq D/2 + 2$  with an arbitrary accuracy  $\varepsilon \in (0, 1)$ . The network in Zhou (2020b) has width increasing linearly with respect to depth and has depth growing in the order of  $\varepsilon^{-2}$  as  $\varepsilon$  decreases to 0. It is shown in Petersen & Voigtlaender (2020); Zhou (2020a) that any approximation error achieved by FNNs can be achieved by CNNs. Combining Zhou (2020a) and Yarotsky (2017), we can show that CNNs can approximate  $W^{s,\infty}$  functions in  $\mathbb{R}^D$  with arbitrary accuracy  $\varepsilon$ . Such CNNs have the number of channels in the order of  $\varepsilon^{-D/s}$  and a cardinality constraint. The only theory on ConvResNet can be found in Oono & Suzuki (2019), where an approximation theory for Hölder functions is proved for ConvResNets with fixed width.

Statistical theories for binary classification by FNNs are established with the hinge loss (Ohn & Kim, 2019; Hu et al., 2020) and the logistic loss (Kim et al., 2018). Among these works, Hu et al. (2020) uses a parametric model given by a teacher-student network. The nonparametric results in Ohn & Kim (2019); Kim et al. (2018) are cursed by the data dimension, and therefore require a large number of samples for high-dimensional data.

Binary classification by CNNs has been studied in Kohler et al. (2020); Kohler & Langer (2020); Nitanda & Suzuki (2018); Huang et al. (2018). Image binary classification is studied in Kohler et al. (2020); Kohler & Langer (2020) in which the probability function is assumed to be in a hierarchical max-pooling model class. ResNet type classifiers are considered in Nitanda & Suzuki (2018); Huang et al. (2018) while the generalization error is not given explicitly.

Low-dimensional structures of data sets are explored for neural networks in Chui & Mhaskar (2018); Shaham et al. (2018); Chen et al. (2019a,b); Schmidt-Hieber (2019); Nakada & Imaizumi (2019); Cloninger & Klock (2020); Chen et al. (2020); Montanelli & Yang (2020). These works show that, if data are near a low-dimensional manifold, the performance of FNNs depends on the intrinsic dimension of the manifold and only weakly depends on the data dimension. Our work focuses on ConvResNets for practical applications.

The networks in many aforementioned works have a cardinality constraint. From the computational perspective,

training such networks requires substantial efforts (Han et al., 2016; 2015; Blalock et al., 2020). In comparison, the ConvResNet in Oono & Suzuki (2019) and this paper does not require any cardinality constraint. Additionally, our constructed network has a fixed filter size and a fixed number of channels, which is desirable for practical applications.

As a summary, we compare our approximation theory and existing results in Table 1.

The rest of this paper is organized as follows: In Section 2, we briefly introduce manifolds, Besov functions on manifolds and convolution. Our main results are presented in Section 3. We give a proof sketch in Section 4 and conclude this paper in Section 5.

## 2. Preliminaries

**Notations:** We use bold lower-case letters to denote vectors, upper-case letters to denote matrices, calligraphic letters to denote tensors, sets and manifolds. For any  $x > 0$ , we use  $\lceil x \rceil$  to denote the smallest integer that is no less than  $x$  and use  $\lfloor x \rfloor$  to denote the largest integer that is no larger than  $x$ . For any  $a, b \in \mathbb{R}$ , we denote  $a \vee b = \max(a, b)$ . For a function  $f : \mathbb{R}^d \rightarrow \mathbb{R}$  and a set  $\Omega \subset \mathbb{R}^d$ , we denote the restriction of  $f$  to  $\Omega$  by  $f|_{\Omega}$ . We use  $\|f\|_{L^p}$  to denote the  $L^p$  norm of  $f$ . We denote the Euclidean ball centered at  $\mathbf{c}$  with radius  $\omega$  by  $B_{\omega}(\mathbf{c})$ .

### 2.1. Low-dimensional manifolds

We first introduce some concepts on manifolds. We refer the readers to Tu (2010); Lee (2006) for details. Throughout this paper, we let  $\mathcal{M}$  be a  $d$ -dimensional Riemannian manifold  $\mathcal{M}$  isometrically embedded in  $\mathbb{R}^D$  with  $d \leq D$ . We first introduce charts, an atlas and the partition of unity.

**Definition 1** (Chart). *A chart on  $\mathcal{M}$  is a pair  $(U, \phi)$  where  $U \subset \mathcal{M}$  is open and  $\phi : U \rightarrow \mathbb{R}^d$ , is a homeomorphism (i.e., bijective,  $\phi$  and  $\phi^{-1}$  are both continuous).*

In a chart  $(U, \phi)$ ,  $U$  is called a coordinate neighborhood and  $\phi$  is a coordinate system on  $U$ . A collection of charts which covers  $\mathcal{M}$  is called an atlas of  $\mathcal{M}$ .

**Definition 2** ( $C^k$  Atlas). *A  $C^k$  atlas for  $\mathcal{M}$  is a collection of charts  $\{(U_{\alpha}, \phi_{\alpha})\}_{\alpha \in \mathcal{A}}$  which satisfies  $\bigcup_{\alpha \in \mathcal{A}} U_{\alpha} = \mathcal{M}$ ,*

and are pairwise  $C^k$  compatible:

$$\phi_\alpha \circ \phi_\beta^{-1} : \phi_\beta(U_\alpha \cap U_\beta) \rightarrow \phi_\alpha(U_\alpha \cap U_\beta) \quad \text{and}$$

$$\phi_\beta \circ \phi_\alpha^{-1} : \phi_\alpha(U_\alpha \cap U_\beta) \rightarrow \phi_\beta(U_\alpha \cap U_\beta)$$

are both  $C^k$  for any  $\alpha, \beta \in \mathcal{A}$ . An atlas is called finite if it contains finitely many charts.

**Definition 3** (Smooth Manifold). A smooth manifold is a manifold  $\mathcal{M}$  together with a  $C^\infty$  atlas.

The Euclidean space, the torus and the unit sphere are examples of smooth manifolds.  $C^s$  functions on a smooth manifold  $\mathcal{M}$  are defined as follows:

**Definition 4** ( $C^s$  functions on  $\mathcal{M}$ ). Let  $\mathcal{M}$  be a smooth manifold and  $f : \mathcal{M} \rightarrow \mathbb{R}$  be a function on  $\mathcal{M}$ . We say  $f$  is a  $C^s$  function on  $\mathcal{M}$ , if for every chart  $(U, \phi)$  on  $\mathcal{M}$ , the function  $f \circ \phi^{-1} : \phi(U) \rightarrow \mathbb{R}$  is a  $C^s$  function.

We next define the  $C^\infty$  partition of unity which is an important tool for the study of functions on manifolds.

**Definition 5** (Partition of Unity). A  $C^\infty$  partition of unity on a manifold  $\mathcal{M}$  is a collection of  $C^\infty$  functions  $\{\rho_\alpha\}_{\alpha \in \mathcal{A}}$  with  $\rho_\alpha : \mathcal{M} \rightarrow [0, 1]$  such that for any  $\mathbf{x} \in \mathcal{M}$ ,

1. there is a neighbourhood of  $\mathbf{x}$  where only a finite number of the functions in  $\{\rho_\alpha\}_{\alpha \in \mathcal{A}}$  are nonzero, and
2.  $\sum_{\alpha \in \mathcal{A}} \rho_\alpha(\mathbf{x}) = 1$ .

An open cover of a manifold  $\mathcal{M}$  is called locally finite if every  $\mathbf{x} \in \mathcal{M}$  has a neighbourhood which intersects with a finite number of sets in the cover. The following proposition shows that a  $C^\infty$  partition of unity for a smooth manifold always exists (Spivak, 1970, Chapter 2, Theorem 15).

**Proposition 1** (Existence of a  $C^\infty$  partition of unity). Let  $\{U_\alpha\}_{\alpha \in \mathcal{A}}$  be a locally finite cover of a smooth manifold  $\mathcal{M}$ . There is a  $C^\infty$  partition of unity  $\{\rho_\alpha\}_{\alpha \in \mathcal{A}}$  such that  $\text{supp}(\rho_\alpha) \subset U_\alpha$ .

Let  $\{(U_\alpha, \phi_\alpha)\}_{\alpha \in \mathcal{A}}$  be a  $C^\infty$  atlas of  $\mathcal{M}$ . Proposition 1 guarantees the existence of a partition of unity  $\{\rho_\alpha\}_{\alpha \in \mathcal{A}}$  such that  $\rho_\alpha$  is supported on  $U_\alpha$ .

The reach of  $\mathcal{M}$  introduced by Federer (Federer, 1959) is an important quantity defined below. Let  $d(\mathbf{x}, \mathcal{M}) = \inf_{\mathbf{y} \in \mathcal{M}} \|\mathbf{x} - \mathbf{y}\|_2$  be the distance from  $\mathbf{x}$  to  $\mathcal{M}$ .

**Definition 6** (Reach (Federer, 1959; Niyogi et al., 2008)). Define the set

$$G = \{\mathbf{x} \in \mathbb{R}^D : \exists \text{ distinct } \mathbf{p}, \mathbf{q} \in \mathcal{M} \text{ such that } d(\mathbf{x}, \mathcal{M}) = \|\mathbf{x} - \mathbf{p}\|_2 = \|\mathbf{x} - \mathbf{q}\|_2\}.$$

The closure of  $G$  is called the medial axis of  $\mathcal{M}$ . The reach of  $\mathcal{M}$  is defined as

$$\tau = \inf_{\mathbf{x} \in \mathcal{M}} \inf_{\mathbf{y} \in G} \|\mathbf{x} - \mathbf{y}\|_2.$$

We illustrate large and small reach in Figure 2.

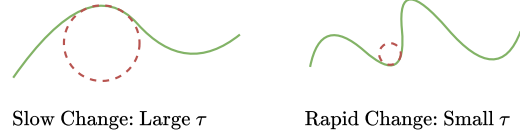


Figure 2. Illustration of manifolds with large and small reach.

## 2.2. Besov functions on a smooth manifold

We next define Besov function spaces on  $\mathcal{M}$ , which generalizes more elementary function spaces such as the Sobolev and Hölder spaces. To define Besov functions, we first introduce the modulus of smoothness.

**Definition 7** (Modulus of Smoothness ( DeVore & Lorentz, 1993; Suzuki, 2019)). Let  $\Omega \subset \mathbb{R}^D$ . For a function  $f : \mathbb{R}^D \rightarrow \mathbb{R}$  be in  $L^p(\Omega)$  for  $p > 0$ , the  $r$ -th modulus of smoothness of  $f$  is defined by

$$w_{r,p}(f, t) = \sup_{\|\mathbf{h}\|_2 \leq t} \|\Delta_{\mathbf{h}}^r(f)\|_{L^p}, \text{ where}$$

$$\Delta_{\mathbf{h}}^r(f)(\mathbf{x}) = \begin{cases} \sum_{j=0}^r \binom{r}{j} (-1)^{r-j} f(\mathbf{x} + j\mathbf{h}) & \text{if } \mathbf{x} \in \Omega, \mathbf{x} + r\mathbf{h} \in \Omega, \\ 0 & \text{otherwise.} \end{cases}$$

**Definition 8** (Besov Space  $B_{p,q}^s(\Omega)$ ). For  $0 < p, q \leq \infty$ ,  $s > 0$ ,  $r = \lfloor s \rfloor + 1$ , define the seminorm  $|\cdot|_{B_{p,q}^s}$  as

$$|f|_{B_{p,q}^s(\Omega)} := \begin{cases} \left( \int_0^\infty (t^{-s} w_{r,p}(f, t))^q \frac{dt}{t} \right)^{\frac{1}{q}} & \text{if } q < \infty, \\ \sup_{t>0} t^{-s} w_{r,p}(f, t) & \text{if } q = \infty. \end{cases}$$

The norm of the Besov space  $B_{p,q}^s(\Omega)$  is defined as  $\|f\|_{B_{p,q}^s(\Omega)} := \|f\|_{L^p(\Omega)} + |f|_{B_{p,q}^s(\Omega)}$ . The Besov space is  $B_{p,q}^s(\Omega) = \{f \in L^p(\Omega) \mid \|f\|_{B_{p,q}^s} < \infty\}$ .

We next define  $B_{p,q}^s$  functions on  $\mathcal{M}$  (Geller & Pesenson, 2011; Triebel, 1983; 1992).

**Definition 9** ( $B_{p,q}^s$  Functions on  $\mathcal{M}$ ). Let  $\mathcal{M}$  be a compact smooth manifold of dimension  $d$ . Let  $\{(U_i, \phi_i)\}_{i=1}^{C_{\mathcal{M}}}$  be a finite atlas on  $\mathcal{M}$  and  $\{\rho_i\}_{i=1}^{C_{\mathcal{M}}}$  be a partition of unity on  $\mathcal{M}$  such that  $\text{supp}(\rho_i) \subset U_i$ . A function  $f : \mathcal{M} \rightarrow \mathbb{R}$  is in  $B_{p,q}^s(\mathcal{M})$  if

$$\|f\|_{B_{p,q}^s(\mathcal{M})} := \sum_{i=1}^{C_{\mathcal{M}}} \|(f\rho_i) \circ \phi_i^{-1}\|_{B_{p,q}^s(\mathbb{R}^d)} < \infty. \quad (1)$$

Since  $\rho_i$  is supported on  $U_i$ , the function  $(f\rho_i) \circ \phi_i^{-1}$  is supported on  $\phi(U_i)$ . We can extend  $(f\rho_i) \circ \phi_i^{-1}$  from  $\phi(U_i)$  to  $\mathbb{R}^d$  by setting the function to be 0 on  $\mathbb{R}^d \setminus \phi(U_i)$ . The extended function lies in the Besov space  $B_{p,q}^s(\mathbb{R}^d)$  (Triebel, 1992, Chapter 7).

## 2.3. Convolution and residual block

In this paper, we consider one-sided stride-one convolution in our network. Let  $\mathcal{W} = \{\mathcal{W}_{j,k,l}\} \in \mathbb{R}^{C' \times K \times C}$  be a filter

where  $C'$  is the output channel size,  $K$  is the filter size and  $C$  is the input channel size. For  $z \in \mathbb{R}^{D \times C}$ , the convolution of  $\mathcal{W}$  with  $z$  gives  $y \in \mathbb{R}^{D \times C'}$  such that

$$y = \mathcal{W} * z, \quad y_{i,j} = \sum_{k=1}^K \sum_{l=1}^C \mathcal{W}_{j,k,l} z_{i+k-1,l}, \quad (2)$$

where  $1 \leq i \leq D$ ,  $1 \leq j \leq C'$  and we set  $z_{i+k-1,l} = 0$  for  $i+k-1 > D$ , as demonstrated in Figure 3(a).

The building blocks of ConvResNets are residual blocks. For an input  $\mathbf{x}$ , each residual block computes

$$\mathbf{x} + F(\mathbf{x})$$

where  $F$  is a subnetwork consisting of convolutional layers (see more details in Section 3.1). A residual block is demonstrated in Figure 3(b).

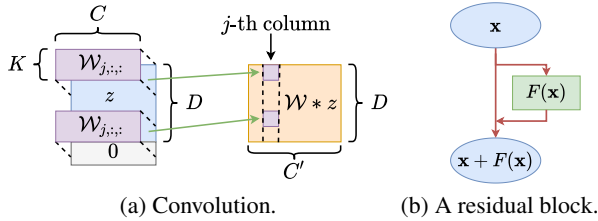


Figure 3. (a) Demonstration of  $\mathcal{W} * z$ , where the input is  $z \in \mathbb{R}^{D \times C}$ , and the output is  $\mathcal{W} * z \in \mathbb{R}^{D \times C'}$ . Here  $\mathcal{W} = \{\mathcal{W}_{j,k,l}\} \in \mathbb{R}^{C' \times K \times C}$  is a filter where  $C'$  is the output channel size,  $K$  is the filter size and  $C$  is the input channel size.  $\mathcal{W}_{j,:,:}$  is a  $D \times C$  matrix for the  $j$ -th output channel. (b) Demonstration of a residual block.

### 3. Theory

In this section, we first introduce the ConvResNet architecture, and then present our main results.

#### 3.1. Convolutional residual neural network

We study the ConvResNet with the rectified linear unit (ReLU) activation function:  $\text{ReLU}(z) = \max(z, 0)$ . The ConvResNet we consider consists of a padding layer and several residual blocks followed by a fully connected feed-forward layer.

We first define the padding layer. Given an input  $A \in \mathbb{R}^{D \times C_1}$ , the network first applies a padding operator  $P : \mathbb{R}^{D \times C_1} \rightarrow \mathbb{R}^{D \times C_2}$  for some integer  $C_2 \geq C_1$  such that

$$Z = P(A) = \begin{bmatrix} A & \mathbf{0} & \cdots & \mathbf{0} \end{bmatrix} \in \mathbb{R}^{D \times C_2}.$$

Then the matrix  $Z$  is passed through  $M$  residual blocks.

In the  $m$ -th block, let  $\mathcal{W}_m = \{\mathcal{W}_m^{(1)}, \dots, \mathcal{W}_m^{(L_m)}\}$  and  $\mathcal{B}_m = \{\mathcal{B}_m^{(1)}, \dots, \mathcal{B}_m^{(L_m)}\}$  be a collection of filters and biases. The  $m$ -th residual block maps a matrix from  $\mathbb{R}^{D \times C}$  to  $\mathbb{R}^{D \times C}$  by

$$\begin{aligned} & \text{Conv}_{\mathcal{W}_m, \mathcal{B}_m} + \text{id}, \\ & \text{where id is the identity operator and} \\ & \text{Conv}_{\mathcal{W}_m, \mathcal{B}_m}(Z) = \text{ReLU}\left(\mathcal{W}_m^{(L_m)} * \dots \right. \\ & \left. \dots * \text{ReLU}\left(\mathcal{W}_m^{(1)} * Z + \mathcal{B}_m^{(1)}\right) \dots + \mathcal{B}_m^{(L_m)}\right), \end{aligned} \quad (3)$$

with ReLU applied entrywise. Denote

$$Q(\mathbf{x}) = (\text{Conv}_{\mathcal{W}_M, \mathcal{B}_M} + \text{id}) \circ \dots \circ (\text{Conv}_{\mathcal{W}_1, \mathcal{B}_1} + \text{id}) \circ P(\mathbf{x}). \quad (4)$$

For networks only consisting of residual blocks, we define the network class as

$$\begin{aligned} \mathcal{C}^{\text{Conv}}(M, L, J, K, \kappa) = \\ \{Q | Q(\mathbf{x}) \text{ is in the form of (4) with } M \text{ residual blocks.} \\ \text{Each block has filter size bounded by } K, \text{ number of} \\ \text{channels bounded by } J, \max_m L_m \leq L, \\ \max_{m,l} \|\mathcal{W}_m^{(l)}\|_\infty \vee \|\mathcal{B}_m^{(l)}\|_\infty \leq \kappa\}. \end{aligned} \quad (5)$$

where  $\|\cdot\|_\infty$  denotes  $\ell^\infty$  norm of a vector, and for a tensor  $\mathcal{W}$ ,  $\|\mathcal{W}\|_\infty = \max_{j,k,l} |\mathcal{W}_{j,k,l}|$ .

Based on the network  $Q$  in (4), a ConvResNet has an additional fully connected layer and can be expressed as

$$f(\mathbf{x}) = WQ(\mathbf{x}) + \mathbf{b} \quad (6)$$

where  $W$  and  $\mathbf{b}$  are the weight matrix and the bias in the fully connected layer. The class of ConvResNets is defined as

$$\begin{aligned} \mathcal{C}(M, L, J, K, \kappa_1, \kappa_2, R) = \\ \{f | f(\mathbf{x}) = WQ(\mathbf{x}) + \mathbf{b} \text{ with } Q \in \mathcal{C}^{\text{Conv}}(M, L, J, K, \kappa_1), \\ \|W\|_\infty \vee \|\mathbf{b}\| \leq \kappa_2, \|f\|_{L^\infty} \leq R\}. \end{aligned} \quad (7)$$

Sometimes we do not have restriction on the output, we omit the parameter  $R$  and denote the network class by  $\mathcal{C}(M, L, J, K, \kappa_1, \kappa_2)$ .

#### 3.2. Approximation theory

Our approximation theory is based on the following assumptions of  $\mathcal{M}$  and the object function  $f^* : \mathcal{M} \rightarrow \mathbb{R}$ .

**Assumption 1.**  $\mathcal{M}$  is a  $d$ -dimensional compact smooth Riemannian manifold isometrically embedded in  $\mathbb{R}^D$ . There is a constant  $B$  such that for any  $\mathbf{x} \in \mathcal{M}$ ,  $\|\mathbf{x}\|_\infty \leq B$ .

**Assumption 2.** The reach of  $\mathcal{M}$  is  $\tau > 0$ .

**Assumption 3.** Let  $0 < p, q \leq \infty$ ,  $d/p + 1 \leq s < \infty$ . Assume  $f^* \in B_{p,q}^s(\mathcal{M})$  and  $\|f^*\|_{B_{p,q}^s(\mathcal{M})} \leq c_0$  for a constant  $c_0 > 0$ . Additionally, we assume  $\|f^*\|_{L^\infty} \leq R$  for a constant  $R > 0$ .

Assumption 3 implies that  $f^*$  is Lipschitz continuous (Triebe, 1983, Section 2.7.1 Remark 2 and Section 3.3.1).

Our first result is the following universal approximation error of ConvResNets for Besov functions on  $\mathcal{M}$ .

**Theorem 1.** Assume Assumption 1-3. For any  $\varepsilon \in (0, 1)$  and positive integer  $K \in [2, D]$ , there is a ConvResNet architecture  $\mathcal{C}(M, L, J, K, \kappa_1, \kappa_2)$  such that, for any

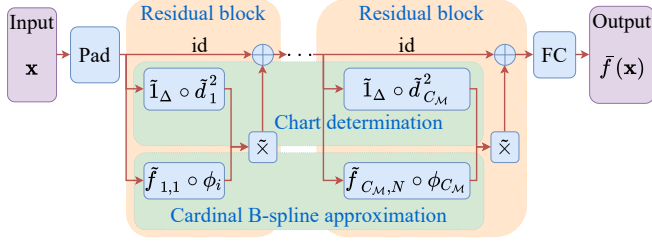


Figure 4. The ConvResNet in Theorem 1 contains a padding layer,  $M$  residual blocks, and a fully connected (FC) layer.

$f^* \in B_{p,q}^s(\mathcal{M})$ , if the weight parameters of this ConvResNet are properly chosen, the network yields a function  $\tilde{f} \in \mathcal{C}(M, L, J, K, \kappa_1, \kappa_2)$  satisfying

$$\|\tilde{f} - f^*\|_{L^\infty} \leq \varepsilon. \quad (8)$$

Such a network architecture has

$$M = O\left(\varepsilon^{-d/s}\right), \quad L = O(\log(1/\varepsilon) + D + \log D),$$

$$J = O(D), \quad \kappa_1 = O(1), \quad \log \kappa_2 = O(\log^2(1/\varepsilon)). \quad (9)$$

The constant hidden in  $O(\cdot)$  depend on  $d, s, \frac{2d}{sp-d}, p, q, c_0, \tau$  and the surface area of  $\mathcal{M}$ .

The architecture of the ConvResNet in Theorem 1 is illustrated in Figure 4. It has the following properties:

- The network has a fixed filter size and a fixed number of channels.
- There is no cardinality constraint.
- The network size depends on the intrinsic dimension  $d$ , and only weakly depends on  $D$ .

Theorem 1 can be compared with Suzuki (2019) on the approximation theory for Besov functions in  $\mathbb{R}^D$  by FNNs as follows: (1) To universally approximate Besov functions in  $\mathbb{R}^D$  with  $\varepsilon$  error, the FNN constructed in Suzuki (2019) requires  $O(\log(1/\varepsilon))$  depth,  $O(\varepsilon^{-D/s})$  width and  $O(\varepsilon^{-D/s} \log(1/\varepsilon))$  nonzero parameters. By exploiting the manifold model, our network size depends on the intrinsic dimension  $d$  and weakly depends on  $D$ . (2) The ConvResNet in Theorem 1 does not require any cardinality constraint, while such a constraint is needed in Suzuki (2019).

### 3.3. Statistical theory

We next consider binary classification on  $\mathcal{M}$ . For any  $\mathbf{x} \in \mathcal{M}$ , denote its label by  $y \in \{-1, 1\}$ . The label  $y$  follows the following Bernoulli-type distribution

$$\mathbb{P}(y = 1|\mathbf{x}) = \eta(\mathbf{x}), \quad \mathbb{P}(y = -1|\mathbf{x}) = 1 - \eta(\mathbf{x}) \quad (10)$$

for some  $\eta : \mathcal{M} \rightarrow [0, 1]$ .

We assume the following data model:

**Assumption 4.** We are given i.i.d. sample  $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ , where  $\mathbf{x}_i \in \mathcal{M}$ , and the  $y_i$ 's are sampled according to (10).

In binary classification, a classifier  $f$  predicts the label of  $\mathbf{x}$  as  $\text{sign}(f(\mathbf{x}))$ . To learn the optimal classifier, we consider

the logistic loss  $\phi(z) = \log(1 + \exp(-z))$ . The logistic risk  $\mathcal{E}_\phi(f)$  of a classifier  $f$  is defined as

$$\mathcal{E}_\phi(f) = \mathbb{E}(\phi(yf(\mathbf{x}))). \quad (11)$$

The minimizer of  $\mathcal{E}_\phi(f)$  is denoted by  $f_\phi^*$ , which satisfies

$$f_\phi^*(\mathbf{x}) = \log \frac{\eta(\mathbf{x})}{1 - \eta(\mathbf{x})}. \quad (12)$$

For any classifier  $f$ , we define its logistic excess risk as

$$\mathcal{E}_\phi(f, f_\phi^*) = \mathcal{E}_\phi(f) - \mathcal{E}_\phi(f_\phi^*). \quad (13)$$

In this paper, we consider ConvResNets with the following architecture:

$$\mathcal{C}^{(n)} = \{f | f = \bar{g}_2 \circ \bar{h} \circ \bar{g}_1 \circ \bar{\eta} \text{ where}$$

$$\bar{\eta} \in \mathcal{C}^{\text{Conv}}(M_1, L_1, J_1, K, \kappa_1), \quad \bar{g}_1 \in \mathcal{C}^{\text{Conv}}(1, 4, 8, 1, \kappa_2),$$

$$\bar{h} \in \mathcal{C}^{\text{Conv}}(M_2, L_2, J_2, 1, \kappa_1), \quad \bar{g}_2 \in \mathcal{C}(1, 3, 8, 1, \kappa_3, 1, R)\} \quad (14)$$

where  $M_1, M_2, L, J, K, \kappa_1, \kappa_2, \kappa_3$  are some parameters to be determined.

The empirical classifier is learned by minimizing the empirical logistic risk:

$$\hat{f}_{\phi,n} = \underset{f \in \mathcal{C}^{(n)}}{\text{argmin}} \frac{1}{n} \sum_{i=1}^n \phi(y_i f(\mathbf{x}_i)). \quad (15)$$

We establish an upper bound on the excess risk of  $\hat{f}_{\phi,n}$ :

**Theorem 2.** Assume Assumption 1, 2 and 4. Assume  $0 < p, q \leq \infty, 0 < s < \infty, s \geq d/p + 1$  and  $\eta \in B_{p,q}^s(\mathcal{M})$  with  $\|\eta\|_{B_{p,q}^s} \leq c_0$  for some constant  $c_0$ . For any  $2 \leq K \leq D$ , we set

$$M_1 = O\left(n^{\frac{2d}{s+2(s\vee d)}}\right), \quad M_2 = O\left(n^{\frac{2s}{s+2(s\vee d)}}\right),$$

$$L_1 = O(\log(1/\varepsilon) + D + \log D), \quad L_2 = O(\log(1/\varepsilon)),$$

$$J_1 = O(D), \quad J_2 = O(1), \quad \kappa_1 = O(1),$$

$$\log \kappa_2 = O(\log^2 n), \quad \kappa_3 = O(\log n), \quad R = O(\log n)$$

for  $\mathcal{C}^{(n)}$ . Then

$$\mathbb{E}(\mathcal{E}_\phi(\hat{f}_{\phi,n}, f_\phi^*)) \leq C n^{-\frac{s}{2s+2(s\vee d)}} \log^4 n \quad (16)$$

for some constant  $C$ . Here  $C$  is linear in  $D \log D$  and additionally depends on  $d, s, \frac{2d}{sp-d}, p, q, c_0, \tau$  and the surface area of  $\mathcal{M}$ . The constant hidden in  $O(\cdot)$  depends on  $d, s, \frac{2d}{sp-d}, p, q, c_0, \tau$  and the surface area of  $\mathcal{M}$ .

Theorem 2 shows that a properly designed ConvResNet gives rise to an empirical classifier, of which the excess risk converges at a fast rate with an exponent depending on the intrinsic dimension  $d$ , instead of  $D$ .

Theorem 2 is proved in Appendix A. Each building block of  $\mathcal{C}^{(n)}$  is constructed for the following purpose:

- $\bar{g}_1 \circ \bar{\eta}$  is designed to approximate a truncated  $\eta$  on  $\mathcal{M}$ , which is realized by Theorem 1.
- $\bar{g}_2 \circ \bar{h}$  is designed to approximate a truncated univariate function  $\log \frac{z}{1-z}$ .

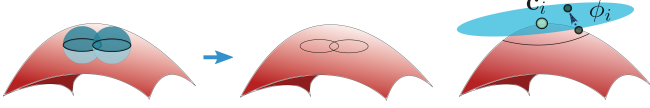


Figure 5. An atlas given by covering  $\mathcal{M}$  using Euclidean balls.

#### 4. Proof of Theorem 1

We provide a proof sketch of Theorem 1 in this section. More technical details are deferred to Appendix C.

We prove Theorem 1 in the following four steps:

1. Decompose  $f^* = \sum_i f_i$  as a sum of locally supported functions according to the manifold structure.
2. Locally approximate each  $f_i$  using cardinal B-splines.
3. Implement the cardinal B-splines using CNNs.
4. Implement the sum of all CNNs by a ConvResNet for approximating  $f^*$ .

##### Step 1: Decomposition of $f^*$ .

• **Construct an atlas on  $\mathcal{M}$ .** Since the manifold  $\mathcal{M}$  is compact, we can cover  $\mathcal{M}$  by a finite collection of open balls  $B_\omega(\mathbf{c}_i)$  for  $i = 1, \dots, C_{\mathcal{M}}$ , where  $\mathbf{c}_i$  is the center of the ball and  $\omega$  is the radius to be chosen later. Accordingly, the manifold is partitioned as  $\mathcal{M} = \bigcup_i U_i$  with  $U_i = B_\omega(\mathbf{c}_i) \cap \mathcal{M}$ . We choose  $\omega < \tau/2$  such that  $U_i$  is diffeomorphic to an open subset of  $\mathbb{R}^d$  (Niyogi et al., 2008, Lemma 5.4). The total number of partitions is then bounded by  $C_{\mathcal{M}} \leq \left\lceil \frac{\text{SA}(\mathcal{M})}{\omega^d} T_d \right\rceil$ , where  $\text{SA}(\mathcal{M})$  is the surface area of  $\mathcal{M}$  and  $T_d$  is the average number of  $U_i$ 's that contain a given point on  $\mathcal{M}$  (Conway et al., 1987, Chapter 2 Equation (1)).

On each partition, we define a projection-based transformation  $\phi_i$  as

$$\phi_i(\mathbf{x}) = a_i V_i^\top (\mathbf{x} - \mathbf{c}_i) + \mathbf{b}_i,$$

where the scaling factor  $a_i \in \mathbb{R}$  and the shifting vector  $\mathbf{b}_i \in \mathbb{R}^d$  ensure  $\phi_i(U_i) \subset [0, 1]^d$ , and the column vectors of  $V_i \in \mathbb{R}^{D \times d}$  form an orthonormal basis of the tangent space  $T_{\mathbf{c}_i}(\mathcal{M})$ . The atlas on  $\mathcal{M}$  is the collection  $(U_i, \phi_i)$  for  $i = 1, \dots, C_{\mathcal{M}}$ . See Figure 5 for a graphical illustration of the atlas.

• **Decompose  $f^*$  according to the atlas.** We decompose  $f^*$  as

$$f^* = \sum_{i=1}^{C_{\mathcal{M}}} f_i \quad \text{with} \quad f_i = f \rho_i, \quad (17)$$

where  $\{\rho_i\}_{i=1}^{C_{\mathcal{M}}}$  is a  $C^\infty$  partition of unity with  $\text{supp}(\phi_i) \subset U_i$ . The existence of such a  $\{\rho_i\}_{i=1}^{C_{\mathcal{M}}}$  is guaranteed by Proposition 1. As a result, each  $f_i$  is supported on a subset of  $U_i$ , and therefore, we can rewrite (17) as

$$f^* = \sum_{i=1}^{C_{\mathcal{M}}} (f_i \circ \phi_i^{-1}) \circ \phi_i \times \mathbb{1}_{U_i} \quad \text{with} \quad f_i = f \rho_i, \quad (18)$$

where  $\mathbb{1}_{U_i}$  is the indicator function of  $U_i$ . Since  $\phi_i$  is a bijection between  $U_i$  and  $\phi_i(U_i)$ ,  $f_i \circ \phi_i^{-1}$  is supported on  $\phi_i(U_i) \subset [0, 1]^d$ . We extend  $f_i \circ \phi_i^{-1}$  on  $[0, 1]^d \setminus \phi_i(U_i)$  by 0. The extended function is in  $B_{p,q}^s([0, 1]^d)$  (see Lemma 4 in Appendix C.1). This allows us to use cardinal B-splines to locally approximate each  $f_i \circ \phi_i^{-1}$  as detailed in Step 2.

**Step 2: Local cardinal B-spline approximation.** We approximate  $f_i \circ \phi_i^{-1}$  using cardinal B-splines  $\tilde{f}_i$  as

$$f_i \circ \phi_i^{-1} \approx \tilde{f}_i \equiv \sum_{j=1}^N \tilde{f}_{i,j} \quad \text{with} \quad \tilde{f}_{i,j} = \alpha_{k,\mathbf{j}}^{(i)} M_{k,\mathbf{j},m}^d, \quad (19)$$

where  $\alpha_{k,\mathbf{j}}^{(i)} \in \mathbb{R}$  is a coefficient and  $M_{k,\mathbf{j},m}^d : [0, 1]^d \rightarrow \mathbb{R}$  denotes a cardinal B-spline with indices  $k, m \in \mathbb{N}^+$ ,  $\mathbf{j} \in \mathbb{R}^d$ . Here  $k$  is a scaling factor,  $\mathbf{j}$  is a shifting vector,  $m$  is the degree of the B-spline and  $d$  is the dimension (see a formal definition in Appendix C.2).

Since  $s \geq d/p + 1$  (by Assumption 3), setting  $r = +\infty, m = \lceil s \rceil + 1$  in Lemma 5 (see Appendix C.3) and applying Lemma 4 gives

$$\|\tilde{f}_i - f_i \circ \phi_i^{-1}\|_{L^\infty} \leq C c_0 N^{-s/d} \quad (20)$$

for some constant  $C$  depending on  $s, p, q$  and  $d$ .

Combining (18) and (19), we approximate  $f^*$  by

$$\tilde{f}^* \equiv \sum_{i=1}^{C_{\mathcal{M}}} \tilde{f}_i \circ \phi_i \times \mathbb{1}_{U_i} = \sum_{i=1}^{C_{\mathcal{M}}} \sum_{j=1}^N \tilde{f}_{i,j} \circ \phi_i \times \mathbb{1}_{U_i}. \quad (21)$$

Such an approximation has error

$$\|\tilde{f}^* - f^*\|_{L^\infty} \leq C C_{\mathcal{M}} c_0 N^{-s/d}.$$

**Step 3: Implement local approximations in Step 2 by CNNs.** In Step 2, (21) gives a natural approximation of  $f^*$ . In the sequel, we aim to implement all ingredients of  $\tilde{f}_{i,j} \circ \phi_i \times \mathbb{1}_{U_i}$  using CNNs. In particular, we show that CNNs can implement the cardinal B-spline  $\tilde{f}_{i,j}$ , the linear projection  $\phi_i$ , the indicator function  $\mathbb{1}_{U_i}$ , and the multiplication operation.

• **Implement  $\mathbb{1}_{U_i}$  by CNNs.** Recall our construction of  $U_i$  in Step 1. For any  $\mathbf{x} \in \mathcal{M}$ , we have  $\mathbb{1}_{U_i}(\mathbf{x}) = 1$  if  $d_i^2(\mathbf{x}) = \|\mathbf{x} - \mathbf{c}_i\|_2^2 \leq \omega^2$ ; otherwise  $\mathbb{1}_{U_i}(\mathbf{x}) = 0$ .

To implement  $\mathbb{1}_{U_i}$ , we rewrite it as the composition of a univariate indicator function  $\mathbb{1}_{[0, \omega^2]}$  and the distance function  $d_i^2$ :

$$\mathbb{1}_{U_i}(\mathbf{x}) = \mathbb{1}_{[0, \omega^2]} \circ d_i^2(\mathbf{x}) \quad \text{for} \quad \mathbf{x} \in \mathcal{M}. \quad (22)$$

We show that CNNs can efficiently implement both  $\mathbb{1}_{[0, \omega^2]}$  and  $d_i^2$ . Specifically, given  $\theta \in (0, 1)$  and  $\Delta \geq 8DB^2\theta$ , there exist CNNs that yield functions  $\tilde{\mathbb{1}}_\Delta$  and  $\tilde{d}_i^2$  satisfying

$$\|\tilde{d}_i^2 - d_i^2\|_{L^\infty} \leq 4B^2D\theta \quad (23)$$

and

$$\tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2(\mathbf{x}) = \begin{cases} 1, & \text{if } \mathbf{x} \in U_i, d_i^2(\mathbf{x}) \leq \omega^2 - \Delta, \\ 0, & \text{if } \mathbf{x} \notin U_i, \\ \text{between 0 and 1,} & \text{otherwise.} \end{cases} \quad (24)$$

We also characterize the network sizes for realizing  $\tilde{\mathbb{I}}_\Delta$  and  $\tilde{d}_i^2$ : The network for  $\tilde{\mathbb{I}}_\Delta$  has  $O(\log(\omega^2/\Delta))$  layers, 2 channels and all weight parameters bounded by  $\max(2, |\omega^2 - 4B^2D\theta|)$ ; the network for  $\tilde{d}_i^2$  has  $O(\log(1/\theta) + D)$  layers,  $6D$  channels and all weight parameters bounded by  $4B^2$ . More technical details are provided in Lemma 9 in Appendix C.6.

• **Implement  $\tilde{f}_{i,j} \circ \phi_i$  by CNNs.** Since  $\phi_i$  is a linear projection, it can be realized by a single-layer perceptron. By Lemma 8 (see Appendix C.5), this single-layer perceptron can be realized by a CNN, denoted by  $\phi_i^{\text{CNN}}$ .

For  $\tilde{f}_{i,j}$ , Proposition 3 (see Appendix C.8) shows that for any  $\delta \in (0, 1)$  and  $2 \leq K \leq d$ , there exists a CNN  $\tilde{f}_{i,j}^{\text{CNN}} \in \mathcal{F}^{\text{CNN}}(L, J, K, \kappa, \kappa)$  with

$$L = O\left(\log \frac{1}{\delta}\right), J = O(1), \kappa = O\left(\delta^{-(\log 2)(\frac{2d}{sp-d} + \frac{c_1}{d})}\right)$$

such that when setting  $N = C_1 \delta^{-d/s}$ , we have

$$\left\| \sum_{j=1}^N \tilde{f}_{i,j}^{\text{CNN}} - f_i \circ \phi_i^{-1} \right\|_{L^\infty(\phi_i(U_i))} \leq \delta, \quad (25)$$

where  $C_1$  is a constant depending on  $s, p, q$  and  $d$ . The constant hidden in  $O(\cdot)$  depends on  $d, s, \frac{2d}{sp-d}, p, q, c_0$ . The CNN class  $\mathcal{F}^{\text{CNN}}$  is defined in Appendix B.

• **Implement the multiplication  $\times$  by a CNN.** According to Lemma 7 (see Appendix C.4) and Lemma 8, for any  $\eta \in (0, 1)$ , the multiplication operation  $\times$  can be approximated by a CNN  $\tilde{\times}$  with  $L^\infty$  error  $\eta$ :

$$\|a \times b - \tilde{\times}(a, b)\|_{L^\infty} \leq \eta. \quad (26)$$

Such a CNN has  $O(\log 1/\eta)$  layers, 6 channels. All parameters are bounded by  $\max(2c_0^2, 1)$ .

**Step 4: Implement  $\tilde{f}^*$  by a ConvResNet.** We assemble all CNN approximations in Step 3 together and show that the whole approximation can be realized by a ConvResNet.

• **Assemble all ingredients together.** Assembling all CNN approximations together gives an approximation of  $\tilde{f}_{i,j} \circ \phi_i \times \mathbb{1}_{U_i}$  as

$$\tilde{f}_{i,j}^\circ \equiv \tilde{\times}(\tilde{f}_{i,j}^{\text{CNN}} \circ \phi_i^{\text{CNN}}, \tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2). \quad (27)$$

After substituting (27) into (21), we approximate the target function  $f^*$  by

$$\tilde{f}^\circ \equiv \sum_{i=1}^{C_M} \sum_{j=1}^N \tilde{f}_{i,j}^\circ. \quad (28)$$

The approximation error of  $\tilde{f}^\circ$  is analyzed in Lemma 12 (see Appendix C.9). According to Lemma 12, the approximation error can be bounded as follows:

$$\begin{aligned} \|\tilde{f}^\circ - f^*\|_{L^\infty} &\leq \sum_{i=1}^{C_M} (A_{i,1} + A_{i,2} + A_{i,3}) \quad \text{with} \\ A_{i,1} &= \sum_{j=1}^N \left\| \tilde{\times}(\tilde{f}_{i,j}^{\text{CNN}} \circ \phi_i^{\text{CNN}}, \tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2) - (\tilde{f}_{i,j}^{\text{CNN}} \circ \phi_i^{\text{CNN}}) \times (\tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2) \right\|_{L^\infty} \leq N\eta, \\ A_{i,2} &= \left\| \left( \sum_{j=1}^N (\tilde{f}_{i,j}^{\text{CNN}} \circ \phi_i^{\text{CNN}}) \right) \times (\tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2) - f_i \times (\tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2) \right\|_{L^\infty} \leq \delta, \\ A_{i,3} &= \|f_i \times (\tilde{\mathbb{I}}_\Delta \circ \tilde{d}_i^2) - f_i \times \mathbb{1}_{U_i}\|_{L^\infty} \leq \frac{c(\pi + 1)}{\omega(1 - \omega/\tau)} \Delta, \end{aligned}$$

where  $\delta, \eta, \Delta$  and  $\theta$  are defined in (25), (26), (24) and (23), respectively. For any  $\varepsilon \in (0, 1)$ , with properly chosen  $\delta, \eta, \Delta$  and  $\theta$  as in (53) in Lemma 12, one has

$$\|\tilde{f}^\circ - f^*\|_{L^\infty} \leq \varepsilon. \quad (29)$$

With these choices, the network size of each CNN is quantified in Appendix C.10.

• **Realize  $\tilde{f}^\circ$  by a ConvResNet.** Lemma 17 (see Appendix C.15) shows that for every  $\tilde{f}_{i,j}^\circ$ , there exists  $\tilde{f}_{i,j}^{\text{CNN}} \in \mathcal{F}^{\text{CNN}}(L, J, K, \kappa_1, \kappa_2)$  with  $L = O(\log 1/\varepsilon + D + \log D)$ ,  $J = O(D)$ ,  $\kappa_1 = O(1)$ ,  $\log \kappa_2 = O(\log^2 1/\varepsilon)$  such that  $\tilde{f}_{i,j}^{\text{CNN}}(\mathbf{x}) = \tilde{f}_{i,j}^\circ(\mathbf{x})$  for any  $\mathbf{x} \in \mathcal{M}$ . As a result, the function  $\tilde{f}^\circ$  in (28) can be expressed as a sum of CNNs:

$$\tilde{f}^\circ = \tilde{f}^{\text{CNN}} \equiv \sum_{i=1}^{C_M} \sum_{j=1}^N \tilde{f}_{i,j}^{\text{CNN}}, \quad (30)$$

where  $N$  is chosen of  $O(\varepsilon^{-d/s})$  (see Proposition 3 and Lemma 12). Lemma 18 (see Appendix C.16) shows that  $\tilde{f}^{\text{CNN}}$  can be realized by  $\tilde{f} \in \mathcal{C}(M, L, J, \kappa_1, \kappa_2)$  with

$$\begin{aligned} M &= O(\varepsilon^{-d/s}), L = O(\log(1/\varepsilon) + D + \log D), \\ J &= O(D), \kappa_1 = O(1), \log \kappa_2 = O(\log^2(1/\varepsilon)). \end{aligned}$$

## 5. Conclusion

Our results show that ConvResNets are adaptive to low-dimensional geometric structures of data sets. Specifically, we establish a universal approximation theory of ConvResNets for Besov functions on a  $d$ -dimensional manifold  $\mathcal{M}$ . Our network size depends on the intrinsic dimension  $d$  and only weakly depends on  $D$ . We also establish a statistical theory of ConvResNets for binary classification when the given data are located on  $\mathcal{M}$ . The classifier is learned by minimizing the empirical logistic loss. We prove that if the

ConvResNet architecture is properly chosen, the excess risk of the learned classifier decays at a fast rate depending on the intrinsic dimension of the manifold.

Our ConvResNet has many practical properties: it has a fixed filter size and a fixed number of channels. Moreover, it does not require any cardinality constraint, which is beneficial to training.

Our analysis can be extended to multinomial logistic regression for multi-class classification. In this case, the network will output a vector where each component represents the likelihood of an input belonging to certain class. By assuming that each likelihood function is in the Besov space, we can apply our analysis to approximate each function by a ConvResNet.

## References

- Alipanahi, B., Delong, A., Weirauch, M. T., and Frey, B. J. Predicting the sequence specificities of DNA- and RNA-binding proteins by deep learning. *Nature Biotechnology*, 33:831–838, 2015.
- Barron, A. R. Universal approximation bounds for superpositions of a sigmoidal function. *IEEE Transactions on Information theory*, 39(3):930–945, 1993.
- Blalock, D., Ortiz, J. J. G., Frankle, J., and Gutttag, J. What is the state of neural network pruning? *arXiv preprint arXiv:2003.03033*, 2020.
- Chen, L.-C., Papandreou, G., Kokkinos, I., Murphy, K., and Yuille, A. L. DeepLAB: Semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected CRFs. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 40(4):834–848, 2017.
- Chen, M., Jiang, H., Liao, W., and Zhao, T. Nonparametric regression on low-dimensional manifolds using deep ReLU networks. *arXiv preprint arXiv:1908.01842*, 2019a.
- Chen, M., Jiang, H., Liao, W., and Zhao, T. Efficient approximation of deep ReLU networks for functions on low dimensional manifolds. In *Advances in Neural Information Processing Systems*, pp. 8172–8182, 2019b.
- Chen, M., Liu, H., Liao, W., and Zhao, T. Doubly robust off-policy learning on low-dimensional manifolds by deep neural networks. *arXiv preprint arXiv:2011.01797*, 2020.
- Chui, C. K. and Mhaskar, H. N. Deep nets for local manifold learning. *Frontiers in Applied Mathematics and Statistics*, 4:12, 2018.
- Cloninger, A. and Klock, T. ReLU nets adapt to intrinsic dimensionality beyond the target domain. *arXiv preprint arXiv:2008.02545*, 2020.
- Conway, J. H., Sloane, N. J. A., and Bannai, E. *Sphere-packings, Lattices, and Groups*. Springer-Verlag, Berlin, Heidelberg, 1987. ISBN 0-387-96617-X.
- Cybenko, G. Approximation by superpositions of a sigmoidal function. *Mathematics of control, signals and systems*, 2(4):303–314, 1989.
- DeVore, R. A. and Lorentz, G. G. *Constructive Approximation*, volume 303. Springer Science & Business Media, 1993.
- DeVore, R. A. and Popov, V. A. Interpolation of Besov spaces. *Transactions of the American Mathematical Society*, 305(1):397–414, 1988.
- Dispa, S. Intrinsic characterizations of Besov spaces on lipschitz domains. *Mathematische Nachrichten*, 260(1): 21–33, 2003.
- Dũng, D. Optimal adaptive sampling recovery. *Advances in Computational Mathematics*, 34(1):1–41, 2011.
- Fang, Z., Feng, H., Huang, S., and Zhou, D.-X. Theory of deep convolutional neural networks II: Spherical analysis. *Neural Networks*, 131:154–162, 2020.
- Federer, H. Curvature measures. *Transactions of the American Mathematical Society*, 93(3):418–491, 1959.
- Geer, S. A. and van de Geer, S. *Empirical Processes in M-estimation*, volume 6. Cambridge University press, 2000.
- Geller, D. and Pesenson, I. Z. Band-limited localized parseval frames and Besov spaces on compact homogeneous manifolds. *Journal of Geometric Analysis*, 21(2):334–371, 2011.
- Girshick, R. Fast R-CNN. In *Proceedings of the IEEE International Conference on Computer Vision*, pp. 1440–1448, 2015.
- Gong, S., Boddeti, V. N., and Jain, A. K. On the intrinsic dimensionality of image representations. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 3987–3996, 2019.
- Graves, A., Mohamed, A.-r., and Hinton, G. Speech recognition with deep recurrent neural networks. In *2013 IEEE International Conference on Acoustics, Speech and Signal Processing*, pp. 6645–6649. IEEE, 2013.
- Hamers, M. and Kohler, M. Nonasymptotic bounds on the  $L_2$  error of neural network regression estimates. *Annals of the Institute of Statistical Mathematics*, 58(1):131–151, 2006.

- Han, S., Pool, J., Tran, J., and Dally, W. J. Learning both weights and connections for efficient neural networks. *arXiv preprint arXiv:1506.02626*, 2015.
- Han, S., Pool, J., Narang, S., Mao, H., Gong, E., Tang, S., Elsen, E., Vajda, P., Paluri, M., Tran, J., et al. Dsd: Dense-sparse-dense training for deep neural networks. *arXiv preprint arXiv:1607.04381*, 2016.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 770–778, 2016.
- Hinton, G. E. and Salakhutdinov, R. R. Reducing the dimensionality of data with neural networks. *Science*, 313 (5786):504–507, 2006.
- Hu, T., Shang, Z., and Cheng, G. Sharp rate of convergence for deep neural network classifiers under the teacher-student setting. *arXiv preprint arXiv:2001.06892*, 2020.
- Huang, F., Ash, J., Langford, J., and Schapire, R. Learning deep resnet blocks sequentially using boosting theory. In *International Conference on Machine Learning*, pp. 2058–2067, 2018.
- Jaffard, S., Meyer, Y., and Ryan, R. D. *Wavelets: tools for science and technology*. SIAM, 2001.
- Jiang, F., Jiang, Y., Zhi, H., Dong, Y., Li, H., Ma, S., Wang, Y., Dong, Q., Shen, H., and Wang, Y. Artificial intelligence in healthcare: past, present and future. *Stroke and vascular neurology*, 2(4):230–243, 2017.
- Kim, Y., Ohn, I., and Kim, D. Fast convergence rates of deep neural networks for classification. *arXiv preprint arXiv:1812.03599*, 2018.
- Kohler, M. and Krzyżak, A. Adaptive regression estimation with multilayer feedforward neural networks. *Nonparametric Statistics*, 17(8):891–913, 2005.
- Kohler, M. and Langer, S. Statistical theory for image classification using deep convolutional neural networks with cross-entropy loss. *arXiv preprint arXiv:2011.13602*, 2020.
- Kohler, M. and Mehnert, J. Analysis of the rate of convergence of least squares neural network regression estimates in case of measurement errors. *Neural Networks*, 24(3): 273–279, 2011.
- Kohler, M., Krzyżak, A., and Walter, B. On the rate of convergence of image classifiers based on convolutional neural networks. *arXiv preprint arXiv:2003.01526*, 2020.
- Krizhevsky, A., Sutskever, I., and Hinton, G. E. ImageNet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems*, pp. 1097–1105, 2012.
- LeCun, Y., Boser, B., Denker, J. S., Henderson, D., Howard, R. E., Hubbard, W., and Jackel, L. D. Backpropagation applied to handwritten zip code recognition. *Neural Computation*, 1(4):541–551, 1989.
- Lee, H., Ge, R., Ma, T., Risteski, A., and Arora, S. On the ability of neural nets to express distributions. In *Conference on Learning Theory*, pp. 1271–1296, 2017.
- Lee, J. M. *Riemannian manifolds: an introduction to curvature*, volume 176. Springer Science & Business Media, 2006.
- Long, J., Shelhamer, E., and Darrell, T. Fully convolutional networks for semantic segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 3431–3440, 2015.
- Lu, Z., Pu, H., Wang, F., Hu, Z., and Wang, L. The expressive power of neural networks: A view from the width. In *Advances in Neural Information Processing Systems*, pp. 6231–6239, 2017.
- McCaffrey, D. F. and Gallant, A. R. Convergence rates for single hidden layer feedforward networks. *Neural Networks*, 7(1):147–158, 1994.
- Mhaskar, H. N. and Micchelli, C. A. Approximation by superposition of sigmoidal and radial basis functions. *Advances in Applied mathematics*, 13(3):350–373, 1992.
- Miotto, R., Wang, F., Wang, S., Jiang, X., and Dudley, J. T. Deep learning for healthcare: review, opportunities and challenges. *Briefings in Bioinformatics*, 19(6):1236–1246, 2018.
- Montanelli, H. and Yang, H. Error bounds for deep ReLU networks using the Kolmogorov–Arnold superposition theorem. *Neural Networks*, 129:1–6, 2020.
- Nakada, R. and Imaizumi, M. Adaptive approximation and estimation of deep neural network with intrinsic dimensionality. *arXiv preprint arXiv:1907.02177*, 2019.
- Nitanda, A. and Suzuki, T. Functional gradient boosting based on residual network perception. In *International Conference on Machine Learning*, pp. 3819–3828, 2018.
- Niyogi, P., Smale, S., and Weinberger, S. Finding the homology of submanifolds with high confidence from random samples. *Discrete & Computational Geometry*, 39(1-3): 419–441, 2008.

- Ohn, I. and Kim, Y. Smooth function approximation by deep neural networks with general activation functions. *Entropy*, 21(7):627, 2019.
- Oono, K. and Suzuki, T. Approximation and non-parametric estimation of ResNet-type convolutional neural networks. In *International Conference on Machine Learning*, pp. 4922–4931, 2019.
- Osher, S., Shi, Z., and Zhu, W. Low dimensional manifold model for image processing. *SIAM Journal on Imaging Sciences*, 10(4):1669–1690, 2017.
- Park, C. Convergence rates of generalization errors for margin-based classification. *Journal of Statistical Planning and Inference*, 139(8):2543–2551, 2009.
- Petersen, P. and Voigtlaender, F. Equivalence of approximation by convolutional neural networks and fully-connected networks. *Proceedings of the American Mathematical Society*, 148(4):1567–1581, 2020.
- Schmidt-Hieber, J. Nonparametric regression using deep neural networks with ReLU activation function. *arXiv preprint arXiv:1708.06633*, 2017.
- Schmidt-Hieber, J. Deep ReLU network approximation of functions on a manifold. *arXiv preprint arXiv:1908.00695*, 2019.
- Sermanet, P., Eigen, D., Zhang, X., Mathieu, M., Fergus, R., and LeCun, Y. Overfeat: Integrated recognition, localization and detection using convolutional networks. *arXiv preprint arXiv:1312.6229*, 2013.
- Shaham, U., Cloninger, A., and Coifman, R. R. Provable approximation properties for deep neural networks. *Applied and Computational Harmonic Analysis*, 44(3):537–557, 2018.
- Shen, X. and Wong, W. H. Convergence rate of sieve estimates. *The Annals of Statistics*, pp. 580–615, 1994.
- Simonyan, K. and Zisserman, A. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
- Spivak, M. D. *A comprehensive introduction to differential geometry*. Publish or Perish, 1970.
- Suzuki, T. Adaptivity of deep ReLU network for learning in besov and mixed smooth besov spaces: optimal rate and curse of dimensionality. In *International Conference on Learning Representations*, 2019.
- Suzuki, T. and Nitanda, A. Deep learning is adaptive to intrinsic dimensionality of model smoothness in anisotropic besov space. *arXiv preprint arXiv:1910.12799*, 2019.
- Tenenbaum, J. B., De Silva, V., and Langford, J. C. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290(5500):2319–2323, 2000.
- Triebel, H. *Theory of Function Spaces*. Modern Birkhäuser Classics. Birkhäuser Basel, 1983.
- Triebel, H. *Theory of function spaces II*. Monographs in Mathematics. Birkhäuser Basel, 1992.
- Tu, L. *An Introduction to Manifolds*. Universitext. Springer New York, 2010. ISBN 9781441973993.
- Wu, Y., Schuster, M., Chen, Z., Le, Q. V., Norouzi, M., Macherey, W., Krikun, M., Cao, Y., Gao, Q., Macherey, K., et al. Google’s neural machine translation system: Bridging the gap between human and machine translation. *arXiv preprint arXiv:1609.08144*, 2016.
- Yarotsky, D. Error bounds for approximations with deep ReLU networks. *Neural Networks*, 94:103–114, 2017.
- Young, T., Hazarika, D., Poria, S., and Cambria, E. Recent trends in deep learning based natural language processing. *IEEE Computational Intelligence Magazine*, 13(3):55–75, 2018.
- Zagoruyko, S. and Komodakis, N. Wide residual networks. *arXiv preprint arXiv:1605.07146*, 2016.
- Zhang, H., Wu, C., Zhang, Z., Zhu, Y., Zhang, Z., Lin, H., Sun, Y., He, T., Mueller, J., Manmatha, R., et al. ResNeSt: Split-attention networks. *arXiv preprint arXiv:2004.08955*, 2020.
- Zhou, D.-X. Theory of deep convolutional neural networks: Downsampling. *Neural Networks*, 124:319–327, 2020a.
- Zhou, D.-X. Universality of deep convolutional neural networks. *Applied and Computational Harmonic Analysis*, 48(2):787–794, 2020b.
- Zhou, J. and Troyanskaya, O. G. Predicting effects of noncoding variants with deep learning-based sequence model. *Nature Methods*, 12(10):931–934, 2015.