



Weierstraß-Institut für
Angewandte Analysis und Stochastik



Estimation and Inference for smooth DNNs.

Blessing of dimension

Vladimir Spokoiny ,
HSE University and IITP RAS Moscow,
WIAS Berlin

4. Juni 2025

1 Introduction

2 Regression using DNN

- Training procedure
- DNN regression. Theoretical study
- A multilayer network

Given data $(Y_i, \mathbf{X}_i)$ with response Y_i and regressors $\mathbf{X}_i \in \mathbb{R}^q$.

Aim: to build a network with a small prediction error $Y_i - m(\mathbf{X}_i)$.

Shallow network: $\mathbb{X} = \sigma(W\mathbf{X})$,

DNN: $\mathbb{X} = \mathbb{X}^{(K)} = \text{DNN}(\mathbf{X}) \in \mathbb{R}^M$,

$$\text{DNN}(\mathbf{X}) = \sigma(W^{(K)})\sigma(W^{(K-1)}) \dots \sigma(W^{(1)}\mathbf{X})),$$

where M is the number of nodes in the last hidden layer,

$W^{(k)} = (w_{mj})$ is a weighting $M_k \times M_{k-1}$ -matrix with rows $W_m^{(k)}$,
and σ is an entry-wise activating function.

Prediction by **linear regression** on $\mathbb{X}$:

$$m(\mathbf{X}_i) = \mathbb{X}_i^\top \boldsymbol{\beta} = \sigma(W\mathbf{X}_i)^\top \boldsymbol{\beta}.$$

Model:

$$L(\mathbf{Y}, W, \boldsymbol{\beta}) = \frac{1}{2} \sum_{i=1}^n |Y_i - \sigma(W \mathbf{X}_i)^\top \boldsymbol{\beta}|^2$$

MLE:

$$(\widehat{W}, \widehat{\boldsymbol{\beta}}) = \operatorname{argmin}_{W, \boldsymbol{\beta}} L(W, \boldsymbol{\beta}).$$

Issues: Non-convexity, lack of identifiability, overparametrization.

- [Schmidt-Hieber, 2020]. Nonparametric regression using deep neural networks with ReLU activation function, Kolmogorov-Arnold representation,
- [Zuowei Shen et al., 2020], Deep Network Approximation Characterized by Number of Neurons,
- [Fan et al., 2024], How do noise tails impact on deep ReLU networks? polynomial noise, tails, Huber, DNN approximation,
- [Kohler and Krzyzak, 2022], Analysis of the rate of convergence of an over-parametrized deep neural network estimate learned by gradient descent,
- [Shen et al., 2022], Approximation with CNNs in Sobolev Space: with Applications to Classification,
- [Liu et al., 2022], Benefits of Overparameterized Convolutional Residual Networks: Function Approximation under Smoothness Constraint,
- [Simionescu-Badea, 2022], Analysis of the rate of convergence of fully connected deep neural network regression estimates with smooth activation,

- [Chen et al., 2019], Efficient Approximation of Deep ReLU Networks for Functions on Low Dimensional Manifolds,
- [Kohler et al., 2019], Estimation of a Function of Low Local Dimensionality by DNN,
- [Jiao et al., 2021], Deep Nonparametric Regression on Approximately Low-dim Manifolds,
- [Liu et al., 2021], Besov Function Approximation ... on Low-Dim Manifolds ... ,
- [Cloninger and Klock, 2021], A deep network construction that adapts to intrinsic dimensionality beyond the domain,
- [Zhang et al., 2023], Effective Minkowski Dimension of Deep Nonparametric Regression: Function Approximation and Statistical Theories,
- [Shen et al., 2023], RePU DNN, manifolds, Differentiable Neural Networks with RePU Activation: with Applications ...
- [Jiao et al., 2023], Deep nonparametric regression on approximate manifolds

Asymptotic minimax rate results

- are **very poor** (small power of n , hidden constants, ...)
- are **useless** (do not help in practical applications because involve unknown constants)
- are sometimes even **misleading** (rate optimality does not lead to a practically efficient method)
- are **too concervative** (functional classes are too big)
- depend on unknown parameters of a functional class
- do not properly explain the important notions like **effective and critical dimension** or **curse of dimension** (rather **curse of minimax rates**)

- Parameter dimension p large or infinite;
- Moderate sample size n ;
- Finite sample “oracle type” theoretical guaranties;
- Explicit dependence on the “effective dimension”;
- mild and realistic model assumptions
- Efficient numerical procedures

Blessing of dimension:

When and why **overparametrization** can be useful?

Extending the parameter space can improve **geometric properties** of the loss function

without increasing the **effective dimension**.

Let $L(\boldsymbol{v})$ be a random function (loss, empirical risk, negative log-likelihood). Consider

$$\tilde{\boldsymbol{v}} = \underset{\boldsymbol{v}}{\operatorname{argmin}} L(\boldsymbol{v}); \quad \boldsymbol{v}^* = \underset{\boldsymbol{v}}{\operatorname{argmin}} \mathbb{E} L(\boldsymbol{v});$$

Includes MLE, LSE, LAD, MC, and many others procedures.

Aim: describe $\tilde{\boldsymbol{v}} - \boldsymbol{v}^*$ (estimation loss) and $L(\tilde{\boldsymbol{v}}) - L(\boldsymbol{v}^*)$ (excess).

Perturbed optimization: $L(\boldsymbol{v})$ is a perturbation of $f(\boldsymbol{v}) = \mathbb{E} L(\boldsymbol{v})$ by a stochastic component

$$\zeta(\boldsymbol{v}) = L(\boldsymbol{v}) - \mathbb{E} L(\boldsymbol{v}).$$

- **Truth** $\mathbf{v}^* \stackrel{\text{def}}{=} \operatorname{argmin}_{\mathbf{v}} \mathbb{E} L(\mathbf{v})$;
- **MLE** $\tilde{\mathbf{v}} \stackrel{\text{def}}{=} \operatorname{argmin}_{\mathbf{v}} L(\mathbf{v})$;
- **Fisher information** $\mathbb{F}(\mathbf{v}) = \nabla^2 \mathbb{E} L(\mathbf{v})$, $\mathbb{F} = \mathbb{F}(\mathbf{v}^*)$;
- **Effective sample size** $\mathbf{N} \stackrel{\text{def}}{=} \lambda_{\min}(\mathbb{F})$;
- **Score** $\zeta(\mathbf{v}) \stackrel{\text{def}}{=} L(\mathbf{v}) - \mathbb{E} L(\mathbf{v})$, $\nabla \zeta = \nabla \zeta(\mathbf{v}^*)$;
- **Effective dimension** $\mathbf{p} \stackrel{\text{def}}{=} \operatorname{tr}(\mathbb{F}^{-1} \operatorname{Var}(\nabla \zeta))$;

■ **Concentration** (requires $p \ll N$)

$$\|\mathbb{F}^{1/2}(\tilde{\mathbf{v}} - \mathbf{v}^*)\| \lesssim \sqrt{p}.$$

■ **Fisher expansion** (requires $p \ll N$)

$$\|\mathbb{F}^{1/2}(\tilde{\mathbf{v}} - \mathbf{v}^* + \mathbb{F}^{-1} \nabla \zeta)\| \lesssim \frac{\|\mathbb{F}^{-1} \nabla \zeta\|^2}{\sqrt{N}}.$$

■ **Wilks expansion** (requires $p \ll N$)

$$|L(\tilde{\mathbf{v}}) - L(\mathbf{v}^*) + \frac{1}{2} \|\mathbb{F}^{-1/2} \nabla \zeta\|^2| \lesssim \frac{\|\mathbb{F}^{-1} \nabla \zeta\|^{3/2}}{\sqrt{N}}.$$

Consider a penalized MLE for a penalty function $\text{pen}_G(\mathbf{v})$

$$\tilde{\mathbf{v}}_G = \underset{\mathbf{v}}{\operatorname{argmin}} L_G(\mathbf{v}) = \underset{\mathbf{v}}{\operatorname{argmin}} \{L(\mathbf{v}) + \text{pen}_G(\mathbf{v})\}$$

Typical examples: $\text{pen}_G(\mathbf{v}) = \|G\mathbf{v}\|^2/2$ and $\text{pen}_\lambda(\mathbf{v}) = \lambda\|\mathbf{v}\|_1$.

Consider three optimization problems

$$\tilde{\mathbf{v}}_G = \underset{\mathbf{v}}{\operatorname{argmin}} L_G(\mathbf{v}),$$

$$\mathbf{v}_G^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} L_G(\mathbf{v}),$$

$$\mathbf{v}^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} L(\mathbf{v}).$$

The penalized MLE $\tilde{\mathbf{v}}_G$ estimates rather $\mathbf{v}_G^*$ than $\mathbf{v}^*$.

$$\mathbf{v}_G^* = \operatorname{argmin}_{\mathbf{v}} \mathbb{E} L_G(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v}} \{ \mathbb{E} L(\mathbf{v}) + \text{pen}_G(\mathbf{v}) \},$$

$$\mathbf{v}^* = \operatorname{argmin}_{\mathbf{v}} \mathbb{E} L(\mathbf{v}).$$

Proposition

Define

$$\mathbb{F}_G \stackrel{\text{def}}{=} \nabla^2 \mathbb{E} L_G(\mathbf{v}_G^*) = \mathbb{F} + \nabla^2 \text{pen}_G(\mathbf{v}_G^*),$$

$$\mathbf{M}_G \stackrel{\text{def}}{=} \nabla \text{pen}_G(\mathbf{v}^*).$$

Then

$$\| \mathbb{F}_G^{1/2} (\mathbf{v}_G^* - \mathbf{v}^* + \mathbb{F}_G^{-1} \mathbf{M}_G) \| \lesssim \frac{\| \mathbb{F}_G^{-1/2} \mathbf{M}_G \|^2}{\sqrt{N}}.$$

Fix $Q: \mathbb{R}^p \rightarrow \mathbb{R}^q$

$$\mathbb{p}_Q \stackrel{\text{def}}{=} \text{tr} \text{Var}(Q \mathbf{F}_G^{-1} \nabla \zeta), \quad \text{variance}$$

$$\mathbf{b}_Q = \|Q \mathbf{F}_G^{-1} \mathbf{M}_G\|, \quad \text{bias}$$

$$\begin{aligned} \mathcal{R}_Q &\stackrel{\text{def}}{=} \mathbb{E} \{ \|Q \mathbf{F}_G^{-1} (\nabla \zeta + \mathbf{M}_G)\|^2 \mathbb{I}_{\Omega(\mathbf{x})} \} \quad \text{squared risk} \\ &\leq \mathbb{p}_Q + \mathbf{b}_Q^2. \end{aligned}$$

Then $\tilde{\mathbf{v}}_G - \mathbf{v}^* \approx \mathbf{F}_G^{-1} (\nabla \zeta + \mathbf{M}_G)$ and

$$(1 - \alpha_Q)^2 \mathcal{R}_Q \leq \mathbb{E} \{ \|Q (\tilde{\mathbf{v}}_G - \mathbf{v}^*)\|^2 \mathbb{I}_{\Omega(\mathbf{x})} \} \leq (1 + \alpha_Q)^2 \mathcal{R}_Q.$$

- **Stochastic linearity**: the stochastic component $\zeta(\mathbf{v}) = L(\mathbf{v}) - \mathbb{E}L(\mathbf{v})$ is linear in $\mathbf{v}$ or $\nabla\zeta(\mathbf{v}) \equiv \nabla\zeta$;
- **Sub-gaussian** tails for $\nabla\zeta$;
- **Convexity**: $\mathbb{E}L_G(\mathbf{v})$ is strongly convex in $\mathbf{v}$;
- **Smoothness**: $\mathbb{E}L_G(\mathbf{v})$ is third order smooth in uniform gateaux sense (**self-concordance**);

- The results are **dimension- and coordinate-free**. The **ambient dimension** $p \leq \infty$ is unimportant; only the **effective dimension** p matters. Can be controlled by a choice of penalization, replacing model reduction. The results require $p \ll N$.
- All the remainder are **explicit** and can be handled. Helpful for further analysis of iterative procedures like stochastic optimization, EM-procedures, etc.
- Conditions are restrictive. **Convexity** is the most critical, but can be ensured by extending the parameter set (**blessing of dimension**) and **localization**.
- The main tool is **perturbed optimization**, [Spokoiny, 2025c, Spokoiny, 2025b].

1 Introduction

2 Regression using DNN

- Training procedure
- DNN regression. Theoretical study
- A multilayer network

Model:

$$L(\mathbf{Y}, W, \boldsymbol{\beta}) = -\frac{1}{2} \sum_{i=1}^n |Y_i - \sigma(W \mathbf{X}_i)^\top \boldsymbol{\beta}|^2$$

MLE:

$$(\widehat{W}, \widehat{\boldsymbol{\beta}}) = \underset{W, \boldsymbol{\beta}}{\operatorname{argmin}} L(W, \boldsymbol{\beta}).$$

Issue: The major assumptions of stochastic linearity and convexity are not fulfilled.

Decouple the structural relations $\mathbf{X} = \sigma(W\mathbf{X})$ and $m = \mathbf{X}^\top \beta$ and replace them by **structural penalty** $\lambda^* \|\mathbf{X} - \sigma(W\mathbf{X})\|^2/2$ and $\lambda^* \|m - \mathbf{X}^\top \beta\|^2/2$ (later $\lambda^* = 1$): for $\mathbf{v} = (W, \mathbf{X}, \beta, m)$

$$\mathcal{L}(\mathbf{v}) = \frac{1}{2} \|\mathbf{Y} - m\|^2 + \frac{1}{2} \|\mathbf{X} - \sigma(W\mathbf{X})\|^2 + \frac{1}{2} \|m - \mathbf{X}^\top \beta\|^2.$$

Noise reduction by presmoothing by $\mathcal{S} : \mathbf{Y} \rightarrow \mathbf{Z} \stackrel{\text{def}}{=} \mathcal{S}\mathbf{Y}$:

$$\mathcal{L}(\mathbf{v}) = \frac{\mu^2}{2} \|\mathbf{Z} - z\|^2 + \frac{1}{2} \|\mathbf{X} - \sigma(W\mathbf{X})\|^2 + \frac{1}{2} \|z - \mathcal{S}\mathbf{X}^\top \beta\|^2.$$

Identifiability issue: partially resolved by ridge penalties:

$$\begin{aligned}\mathcal{L}(\mathbf{v}) = & \frac{\mu^2}{2} \|\mathbf{Z} - \mathbf{z}\|^2 + \frac{1}{2} \|\mathbf{z} - \mathcal{S} \mathbf{X}^\top \boldsymbol{\beta}\|^2 + \frac{1}{2} \|\mathbf{X} - \sigma(W \mathbf{X})\|^2 \\ & + \frac{\lambda_{\mathbf{W}}}{2} \|W\|^2 + \frac{\lambda_{\mathbf{X}}}{2} \|\mathbf{X}\|^2 + \frac{\lambda_{\boldsymbol{\beta}}}{2} \|\boldsymbol{\beta}\|^2 + \frac{\lambda_{\mathbf{z}}}{2} \|\mathbf{z}\|^2.\end{aligned}$$

The estimator $\tilde{\mathbf{v}} = (\widetilde{W}, \widetilde{\mathbf{X}}, \widetilde{\boldsymbol{\beta}}, \widetilde{\mathbf{z}})$ and the background truth $\mathbf{v}^* = (W^*, \mathbf{X}^*, \boldsymbol{\beta}^*, \mathbf{z}^*)$:

$$\tilde{\mathbf{v}} = \underset{\mathbf{v}}{\operatorname{argmin}} \mathcal{L}(\mathbf{v}), \quad \mathbf{v}^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} \mathcal{L}(\mathbf{v}).$$

The main results describe the error of estimation $\tilde{\mathbf{v}} - \mathbf{v}^*$.

For $\mathbf{v} = (W, \mathbf{X}, \boldsymbol{\beta}, \mathbf{z})$, consider

$$\begin{aligned} \mathcal{L}(\mathbf{v}) = & \frac{\mu^2}{2} \|\mathbf{Z} - \mathbf{z}\|^2 + \frac{1}{2} \|\mathbf{z} - \mathcal{S} \mathbf{X}^\top \boldsymbol{\beta}\|^2 + \frac{1}{2} \|\mathbf{X} - \sigma(W \mathbf{X})\|^2 \\ & + \frac{\lambda_W}{2} \|W\|^2 + \frac{\lambda_X}{2} \|\mathbf{X}\|^2 + \frac{\lambda_\beta}{2} \|\boldsymbol{\beta}\|^2 + \frac{\lambda_z}{2} \|\mathbf{z}\|^2. \end{aligned}$$

Sufficient condition for local convexity: for $\mathbf{v} \in U(\mathbf{v}^*)$:

$$\mathcal{F}(\mathbf{v}) = \nabla^2 \mathcal{L}(\mathbf{v}) > 0.$$

$$\begin{aligned}\mathcal{L}(\mathbf{v}) = & \frac{\mu^2}{2} \|\mathbf{Z} - \mathbf{z}\|^2 + \frac{1}{2} \|\mathbf{z} - \mathcal{S} \mathbf{X}^\top \boldsymbol{\beta}\|^2 + \frac{1}{2} \|\mathbf{X} - \sigma(W \mathbf{X})\|^2 \\ & + \frac{\lambda_{\mathbf{w}}}{2} \|\mathbf{W}\|^2 + \frac{\lambda_{\mathbf{x}}}{2} \|\mathbf{X}\|^2 + \frac{\lambda_{\boldsymbol{\beta}}}{2} \|\boldsymbol{\beta}\|^2 + \frac{\lambda_{\mathbf{z}}}{2} \|\mathbf{z}\|^2.\end{aligned}$$

Lemma

Fix $\mathbf{v} = (W, \mathbf{X}, \boldsymbol{\beta}, \mathbf{z})$ with $\mathbf{X} = \sigma(W \mathbf{X})$ and $\mathbf{z} = \mathcal{S} \mathbf{X}^\top \boldsymbol{\beta}$. Then for any $\mathbf{u} = (\mathbf{w}, \mathbf{x}, \mathbf{b}, \mathbf{h})$

$$\begin{aligned}\left. \frac{d^2}{dt^2} \mathcal{L}(\mathbf{v} + t\mathbf{u}) \right|_{t=0} &= (\mu^2 + \lambda_{\mathbf{z}}) \|\mathbf{h}\|^2 + \lambda_{\mathbf{x}} \|\mathbf{x}\|^2 + \lambda_{\boldsymbol{\beta}} \|\mathbf{b}\|^2 + \lambda_{\mathbf{w}} \|\mathbf{w}\|^2 \\ &+ \sum_{i=1}^n \sum_{m=1}^M \left\{ \sigma'(W_m \mathbf{X}_i) \mathbf{w}_m \mathbf{X}_i - \mathbf{x}_{mi} \right\}^2 + \|\mathbf{h} - \mathcal{S}(\mathbf{X}^\top \mathbf{b} + \mathbf{x}^\top \boldsymbol{\beta})\|^2.\end{aligned}$$

Define

$$\mathbb{F}_m(\mathbf{v}) = \mathbb{F}_m(W_m) \stackrel{\text{def}}{=} \sum_{i=1}^n \sigma'(W_m \mathbf{X}_i)^2 \mathbf{X}_i \mathbf{X}_i^\top + \frac{4}{3} \lambda_{\mathbf{w}} \mathbb{I}_d,$$

$$\mathcal{F}_{\mathbf{w}\mathbf{w}}(\mathbf{v}) = \mathcal{F}_{\mathbf{w}\mathbf{w}}(W) = \text{block}\{\mathbb{F}_1(W_1), \dots, \mathbb{F}_M(W_M)\}.$$

Lemma

Let $\mu^2 \geq 3$, $\lambda_{\mathbf{x}} \geq 4 + \frac{3b_0^2}{2}$. For $\mathbf{v} = (W, \mathbb{X}, \boldsymbol{\beta}, \mathbf{z})$ with $\mathbb{X} = \sigma(W \mathbf{X})$, $\mathbf{z} = \mathcal{S} \mathbb{X}^\top \boldsymbol{\beta}$, and $\|\boldsymbol{\beta}\| \leq b_0$

$$\mathcal{F}(\mathbf{v}) \geq \text{block}\left\{ \frac{3}{4} \mathcal{F}_{\mathbf{w}\mathbf{w}}, \left(\lambda_{\mathbf{x}} - 3 - \frac{3b_0^2}{2} \right) \mathbb{I}_{\mathbf{x}}, \right. \\ \left. \frac{1}{2} \mathbb{X} \mathbf{A} \mathbb{X}^\top + \lambda_{\boldsymbol{\beta}} \mathbb{I}_{\boldsymbol{\beta}}, (\mu^2 - 3) \mathbb{I}_{\mathbf{z}} \right\}.$$

Introduce a full dimensional metric tensor $\mathcal{D}$ by

$$\mathcal{D}^2 \stackrel{\text{def}}{=} \text{block}\{D_{\mathbf{w}}^2, \mathbb{I}_{\mathbf{x}}, D_{\boldsymbol{\beta}}^2, \mathbb{I}_{\mathbf{z}}\},$$

$$D_{\mathbf{w}}^2 \stackrel{\text{def}}{=} \frac{3}{4} \text{block}\{\mathbb{F}_1, \dots, \mathbb{F}_M\}, \quad D_{\boldsymbol{\beta}}^2 \stackrel{\text{def}}{=} \frac{1}{2} \mathbf{X}^* \mathbf{A} \mathbf{X}^{*\top} + \lambda_{\boldsymbol{\beta}} \mathbb{I}_{\boldsymbol{\beta}},$$

$$\mathbb{F}_m(W_m) = \sum_{i=1}^n \sigma'(W_m \mathbf{X}_i)^2 \mathbf{X}_i \mathbf{X}_i^{\top} + \frac{4}{3} \lambda_{\mathbf{w}} \mathbb{I}_d.$$

With $\sigma(t) = \alpha^{-1} \log(1 + e^{\alpha t})$, define

$$N^{-1/2} = \max_{m \leq M, i \leq n} \left\{ \frac{\sigma''(W_m^* \mathbf{X}_i)}{\sigma'(W_m^* \mathbf{X}_i)} \|\mathbb{F}_m^{-1/2} \mathbf{X}_i\| \right\}, \quad N_1^{-1/2} = \|D_{\boldsymbol{\beta}}^{-1}\|.$$

Under $\|\boldsymbol{\beta}\| \leq b_0$, $\mathbf{X} \mathbf{A} \mathbf{X}^{\top} \leq 2 \mathbf{X}^* \mathbf{A} \mathbf{X}^{*\top}$, smoothness cond. holds with

$$\tau_3 \stackrel{\text{def}}{=} 6\sqrt{e} N^{-1/2} + 6(b_0 + 1) N_1^{-1/2}.$$

For the score vector $\nabla\zeta$, it holds $\nabla\zeta = (0, 0, 0, \nabla_z\zeta)$ with

$$\nabla_z\zeta = \mathbf{Z} - \mathbb{E}\mathbf{Z}, \quad \text{Var}(\nabla_z\zeta) = \text{Var}(\mathcal{S}\mathbf{Y}) = \mathcal{S} \text{Var}(\mathbf{Y}) \mathcal{S}^\top.$$

With $\mathcal{F} = \mathcal{F}(\mathbf{v}^*) = \nabla^2 \mathcal{L}(\mathbf{v}^*)$, the **effective dimension**

$$\mathfrak{p} \stackrel{\text{def}}{=} \text{tr}\{\mathcal{F}^{-1} \text{Var}(\nabla\zeta)\}.$$

Then $\mathcal{F}^{-1} \leq \mathcal{D}^{-2}$ and

$$\|\mathcal{D}\mathcal{F}^{-1}\nabla\zeta\| \leq \|\mathcal{D}^{-1}\nabla\zeta\| = \|\mathcal{S}\epsilon\|,$$

$$\mathfrak{p} \leq \text{tr}\{\mathcal{D}^{-2} \text{Var}(\nabla\zeta)\} \leq \mathbb{E}\|\mathcal{S}\epsilon\|^2.$$

The **effective radius** $r_{\mathcal{D}}$ fulfills with $\mathbb{W}^2 \stackrel{\text{def}}{=} \text{Var}(\mathcal{S}\epsilon)$

$$r_{\mathcal{D}} \leq z(\mathbb{W}^2, \mathbf{x}) \leq \sqrt{\text{tr } \mathbb{W}^2} + \sqrt{2\mathbf{x} \|\mathbb{W}^2\|} \approx \sqrt{\mathfrak{p}}.$$

Identifiability of the product $\mathbb{X}^\top \beta$ after restricting to a bounded set

$$\mathcal{R}^\circ \stackrel{\text{def}}{=} \{v = (W, \mathbb{X}, \beta, z) : \|\beta\| \leq b_0, \mathbb{X} \mathbb{A} \mathbb{X}^\top \leq 2 \mathbb{X}^* \mathbb{A} \mathbb{X}^{*\top}\}.$$

Here $\mathbb{A} = \mathcal{S}^\top \mathcal{S}$ and b_0 is a fixed constant, e.g. $b_0 = 2$.

Theorem

Let $\|\mathcal{S}^\top \mathcal{S}\|_{\text{op}} \leq 1$. If $\lambda_z \geq 3$ and $\lambda_{\mathbb{X}} \geq 4 + 3b_0^2/2$, let $r_{\mathcal{D}} \tau_3 \leq 4/9$. For any $\mathcal{Q}: \mathbb{R}^{\bar{p}} \rightarrow \mathbb{R}^q$, on a random set $\Omega(\mathbf{x})$ with $\mathbb{P}(\Omega(\mathbf{x})) \geq 1 - e^{-x}$

$$\|\mathcal{Q}(\tilde{v} - v^* + \mathcal{F}^{-1} \nabla \zeta)\| \leq \|\mathcal{Q} \mathcal{D}^{-1}\| \frac{3\tau_3}{4} \|\mathcal{S} \varepsilon\|^2.$$

where $\mathcal{D}^2 = \text{block}\{D_{\mathbb{W}}^2, \mathbb{I}_{\mathbb{X}}, D_{\beta}^2, \mathbb{I}_z\}$,

$$D_{\mathbb{W}}^2 \stackrel{\text{def}}{=} \frac{3}{4} \text{block}\{\mathbb{F}_1, \dots, \mathbb{F}_M\}, \quad D_{\beta}^2 \stackrel{\text{def}}{=} \frac{1}{2} \mathbb{X}^* \mathbb{A} \mathbb{X}^{*\top} + \lambda_{\beta} \mathbb{I}_{\beta}.$$

Theorem

Introduce

$$\mathbb{p}_Q \stackrel{\text{def}}{=} \mathbb{E} \{ \|Q \mathcal{F}^{-1} \nabla \zeta\|^2 \mathbb{1}_{\Omega(\mathbf{x})} \} ,$$

assume $\mathbb{E} \{ \|S\varepsilon\|^4 \mathbb{1}_{\Omega(\mathbf{x})} \} \leq C_4^2 \mathbb{p}^2$, and define

$$\alpha_Q \stackrel{\text{def}}{=} \frac{\|Q \mathcal{D}^{-1}\| (3/4) \tau_3 C_4 \mathbb{p}}{\sqrt{\mathbb{p}_Q}} .$$

If $\alpha_Q < 1$ then

$$(1 - \alpha_Q)^2 \mathbb{p}_Q \leq \mathbb{E} \{ \|Q (\tilde{\mathbf{v}} - \mathbf{v}^*)\|^2 \mathbb{1}_{\Omega(\mathbf{x})} \} \leq (1 + \alpha_Q)^2 \mathbb{p}_Q .$$

Denote $\mathcal{I} = \mathcal{F}^{-1}$. Then

$$\mathcal{F}^{-1} \nabla \zeta = (\mathcal{I}_{\mathbf{w}z} \mathcal{S}\epsilon, \mathcal{I}_{\mathbf{x}z} \mathcal{S}\epsilon, \mathcal{I}_{\beta z} \mathcal{S}\epsilon, \mathcal{I}_{zz} \mathcal{S}\epsilon),$$

$$D_{\mathbf{w}}^2 \stackrel{\text{def}}{=} \frac{3}{4} \text{block}\{\mathbf{F}_1, \dots, \mathbf{F}_M\}, \quad D_{\beta}^2 \stackrel{\text{def}}{=} \frac{1}{2} \mathbf{X}^* \mathbf{A} \mathbf{X}^{*\top} + \lambda_{\beta} \mathbb{I}_{\beta}.$$

Corollary

On $\Omega(\mathbf{x})$, with $\mathcal{I} = \mathcal{F}^{-1}$,

$$\|D_{\mathbf{w}}(\widetilde{W} - W^* + \mathcal{I}_{\mathbf{w}z} \mathcal{S}\epsilon)\| \leq (3/4) \tau_3 \|\mathcal{S}\epsilon\|^2,$$

$$\|\widetilde{\mathbf{X}} - \mathbf{X}^* + \mathcal{I}_{\mathbf{x}z} \mathcal{S}\epsilon\| \leq (3/4) \tau_3 \|\mathcal{S}\epsilon\|^2,$$

$$\|D_{\beta}(\widetilde{\beta} - \beta^* + \mathcal{I}_{\beta z} \mathcal{S}\epsilon)\| \leq (3/4) \tau_3 \|\mathcal{S}\epsilon\|^2,$$

$$\|\widetilde{z} - z^* + \mathcal{I}_{zz} \mathcal{S}\epsilon\| \leq (3/4) \tau_3 \|\mathcal{S}\epsilon\|^2$$

Construction can be extended to a K -layer network using recurrence

$$\mathbb{X}^{(k)} = \sigma(W^{(k)} \mathbb{X}^{(k-1)})$$

for $k = 1, \dots, K$ and $\mathbb{X}^{(0)} = \mathbf{X}$. With $\mathbb{X} = (\mathbb{X}^{(1)}, \dots, \mathbb{X}^{(K)})$ and $W = (W^{(1)}, \dots, W^{(K)})$, this leads to the log-likelihood

$$\begin{aligned} \mathcal{L}(\mathbf{v}) = \mathcal{L}(W, \mathbb{X}, \beta, \mathbf{z}) = & \frac{\mu^2}{2} \|\mathbf{Z} - \mathbf{z}\|^2 \\ & + \frac{1}{2} \|\mathbf{z} - \mathcal{S} \mathbb{X}^{(K)\top} \beta\|^2 + \frac{1}{2} \sum_{k=1}^K \|\mathbb{X}^{(k)} - \sigma(W^{(k)} \mathbb{X}^{(k-1)})\|^2 \\ & + \frac{\lambda_{\mathbb{X}, K}}{2} \|\mathbb{X}^{(K)}\|^2 + \frac{\lambda_{\beta}}{2} \|\beta\|^2 + \frac{\lambda_{\mathbf{z}}}{2} \|\mathbf{z}\|^2 + \frac{\lambda_{\mathbf{W}}}{2} \|\mathbf{W}\|^2 + \frac{\lambda_{\mathbf{X}}}{2} \|\mathbf{X}\|^2. \end{aligned}$$

For any $\mathbf{u} = (\mathbf{w}, \mathbf{x}, \mathbf{b}, \mathbf{h})$ with $\mathbf{w} = (\mathbf{w}^{(k)}) \in \mathbb{R}^{n_{\mathbf{w}}}$, $\mathbf{x} = (\mathbf{x}^{(k)}) \in \mathbb{R}^{n \times d}$, $\mathbf{b} \in \mathbb{R}^{M_K}$, and $\mathbf{h} \in \mathbb{R}^q$, it holds

$$\begin{aligned} \left. \frac{d^2}{dt^2} \mathcal{L}(\mathbf{v} + t\mathbf{u}) \right|_{t=0} &= \lambda_{\mathbf{x}} \|\mathbf{x}\|^2 + \lambda_{\mathbf{w}} \|\mathbf{w}\|^2 + (1 + \lambda_{\mathbf{h}}) \|\mathbf{h}\|^2 + \lambda_{\beta} \|\mathbf{b}\|^2 \\ &+ \sum_{k=1}^K \sum_{m=1}^{M_k} \sum_{i=1}^n \left\{ \sigma'(W_m^{(k)} \mathbf{x}_i^{(k-1)}) \left(\mathbf{w}_m^{(k)} \mathbf{x}_i^{(k-1)} + W_m^{(k)} \mathbf{x}_i^{(k-1)} \right) - \mathbf{x}_{mi}^{(k)} \right\}^2 \\ &+ \left\| \mathbf{h} - \mathcal{S}(\mathbf{x}^{(K)})^\top \mathbf{b} + \mathbf{x}^{(K)}^\top \beta \right\|^2 + \lambda_{\mathbf{x}, K} \|\mathbf{x}^{(K)}\|^2. \end{aligned}$$

With $\mathbf{v} = (W^{(1)}, \mathbb{X}^{(1)}, \dots, W^{(K)}, \mathbb{X}^{(K)})$ and $\mathbf{v}^{(k)} = (W^{(k)}, \mathbb{X}^{(k-1)})$, define for $k \leq K$ and $m \leq M_k$

$$\mathbb{F}_m^{(k)}(\mathbf{v}^{(k)}) \stackrel{\text{def}}{=} \sum_{i=1}^n \sigma'(W_m^{(k)} \mathbb{X}_i^{(k-1)})^2 \mathbb{X}_i^{(k-1)} \mathbb{X}_i^{(k-1)\top} + \lambda_k \mathbb{I}_{M_k},$$

$$\mathbb{F}^{(k)}(\mathbf{v}^{(k)}) \stackrel{\text{def}}{=} \text{block}\{\mathbb{F}_1^{(k)}(\mathbf{v}^{(k)}), \dots, \mathbb{F}_{M_k}^{(k)}(\mathbf{v}^{(k)})\}.$$

Similarly define for $k = 2, \dots, K$ and $i = 1, \dots, n$

$$\mathbb{G}_i^{(k)}(\mathbf{v}^{(k)}) \stackrel{\text{def}}{=} \sum_{m=1}^{M_k} \sigma'(W_m^{(k)} \mathbb{X}_i^{(k-1)})^2 W_m^{(k)\top} W_m^{(k)},$$

$$\mathbb{G}^{(k)}(\mathbf{v}^{(k)}) \stackrel{\text{def}}{=} \text{block}\{\mathbb{G}_1^{(k)}(\mathbf{v}^{(k)}), \dots, \mathbb{G}_n^{(k)}(\mathbf{v}^{(k)})\}.$$

Also set $\mathbb{G}^{(1)}(\mathbf{v}^{(1)}) \equiv 0$.

Lemma

Let $\mathbf{X}^{(k)} \equiv \sigma(W^{(k)} \mathbf{X}^{(k-1)})$, $k = 1, \dots, K$. Also, assume $\|\beta\| \leq b_0$ and

$$\max_{i=1,\dots,n} \max_{k=1,\dots,K} \frac{3}{2} \|\mathbf{G}_i^{(k)}(\mathbf{v}^{(k)})\|_{\text{op}} \leq G_0 \leq \lambda_{\mathbf{X}} - 4.$$

Then it holds

$$\mathcal{F}(\mathbf{v}) \geq \text{block} \left\{ \frac{1}{2} \mathbf{F}^{(1)}(\mathbf{v}^{(1)}), \mathbb{I}_{\mathbf{X}^{(1)}}, \dots, \frac{1}{2} \mathbf{F}^{(K)}(\mathbf{v}^{(K)}), \mathbb{I}_{\mathbf{X}^{(K)}}, \right. \\ \left. \frac{1}{2} \mathbf{X}^{(K)} \mathbf{A} \mathbf{X}^{(K)\top} + \lambda_{\beta} \mathbb{I}_{\beta}, (\lambda_z - 2) \mathbb{I}_z \right\}.$$

- (Parametric) lower bounds. Please, ask about
- LASSO-type procedures. Issue $\|v\|_1$ is non-smooth.
 Idea: replace by a smooth penalty ensuring $p \asymp \#(\text{active set})$.
- Other models including
 - dimension reduction, manifold learning
 - density estimation
 - error-in-operator models
 - matrix completion
 - inference for stochastic processes
 - statistical nonlinear inverse problems, PDEs
 - ...
- adaptation, parameter/penalty choice
- higher-order expansions, estimation of smooth functionals.



Banach, S. (1938).

Über homogene Polynome in (L^2) .
Studia Mathematica, 7(1):36–44.



Chen, M., Jiang, H., Liao, W., and Zhao, T. (2019).

Efficient approximation of deep relu networks for functions on low dimensional manifolds.
ArXiv, abs/1908.01842:null.



Cloninger, A. and Klock, T. (2021).

A deep network construction that adapts to intrinsic dimensionality beyond the domain.
Neural Networks, 141:404–419.



Fan, J., Gu, Y., and Zhou, W.-X. (2024).

How do noise tails impact on deep ReLU networks?
The Annals of Statistics, 52(4):1845–1871.



Jiao, Y., Shen, G., Lin, Y., and Huang, J. (2021).

Deep nonparametric regression on approximately low-dimensional manifolds.
arXiv: Statistics Theory, page null.



Jiao, Y., Shen, G., Lin, Y., and Huang, J. (2023).

Deep nonparametric regression on approximate manifolds: Nonasymptotic error bounds with polynomial prefactors.
The Annals of Statistics, 51:691–716.



Kohler, M. and Krzyżak, A. (2022).

Analysis of the rate of convergence of an over-parametrized deep neural network estimate learned by gradient descent.



Kohler, M., Krzyżak, A., and Langer, S. (2019).
Estimation of a function of low local dimensionality by deep neural networks.
IEEE Transactions on Information Theory, 68:4032–4042.



Liu, H., Chen, M., Er, S., Liao, W., Zhang, T.-M., and Zhao, T. (2022).
Benefits of overparameterized convolutional residual networks: Function approximation under smoothness constraint.
ArXiv, abs/2206.04569:null.



Liu, H., Chen, M., Zhao, T., and Liao, W. (2021).
Besov function approximation and binary classification on low-dimensional manifolds using convolutional residual networks.



Puchkin, N., Noskov, F., and Spokoiny, V. (2025).
Sharper dimension-free bounds on the Frobenius distance between sample covariance and its expectation.
Bernoulli, 31(2):1664–1691.
<https://arxiv.org/abs/2308.14739>.



Schmidt-Hieber, J. (2020).
Nonparametric regression using deep neural networks with ReLU activation function.
The Annals of Statistics, 48(4):1875–1897.



Shen, G., Jiao, Y., Lin, Y., and Huang, J. (2022).
Approximation with cnns in sobolev space: with applications to classification.



Shen, G., Jiao, Y., Lin, Y., and Huang, J. (2023).
Differentiable neural networks with repu activation: with applications to score estimation and isotonic regression.
ArXiv, abs/2305.00608:null.

- 
- Simionescu-Badea, C. (2022).
Analysis of the rate of convergence of fully connected deep neural network regression estimates with smooth activation function.
- 
- Spokoiny, V. (2024).
Sharp deviation bounds and concentration phenomenon for the squared norm of a sub-gaussian vector.
<https://arxiv.org/abs/2305.07885>.
- 
- Spokoiny, V. (2025a).
Deviation bounds for the norm of a random vector under exponential moment conditions with applications.
Probability, Uncertainty and Quantitative Risk, 10(1):135–158.
<https://arxiv.org/abs/2309.02302v1>.
- 
- Spokoiny, V. (2025b).
Marginal minimization and sup-norm expansions in perturbed optimization.
[arxiv.org/2505.02562](https://arxiv.org/abs/2505.02562).
- 
- Spokoiny, V. (2025c).
Sharp bounds in perturbed smooth optimization.
[arxiv.org/2504.11834](https://arxiv.org/abs/2504.11834).
- 
- Zhang, Z., Chen, M., Wang, M., Liao, W., and Zhao, T. (2023).
Effective minkowski dimension of deep nonparametric regression: Function approximation and statistical theories.
ArXiv, abs/2306.14859:null.
- 
- Zuwei Shen, Z. S., Haizhao Yang, H. Y., and Shijun Zhang, S. Z. (2020).
Deep network approximation characterized by number of neurons.
Communications in Computational Physics, 28(5):1768–1811.