

What makes NN good learners?

• 2L NN

RQ: Why are NNs so good at learning from data?

→ NN adapt to **STRUCTURE** in the DATA

Spoiler Feature Learning $\leftrightarrow$ Spectral Statistics of the weights

1) Motivation and Settings

Supervised Learning

$$\begin{aligned} D &= \{(x_i, y_i) \in \mathbb{R}^d \times \mathbb{R} : i \in [n]\} \\ &= \{1, \dots, n\} \end{aligned}$$

Define some data distribution

$$(x, y) \sim P_{x, y}$$

Goal. Learn $F : \mathbb{R}^d \rightarrow \mathbb{R}$ "generalizes" well
 $x \mapsto f(x) = \hat{y}$

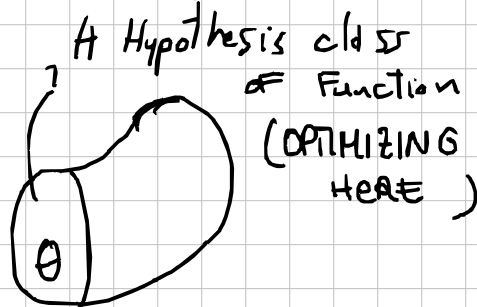
$$\min_F E[l(y, f(x))]$$

- 2) Problems:
- I) Don't know the data distribution $p_{x,y}$
 - II) How to optimise over $\{f: \mathbb{R}^d \rightarrow \mathbb{R}\}$?

ERM Empirical Risk Minimisation

$$\min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n \ell(y_i, f(x_i))$$

TREATS Generalization



QUESTION: How big n do i need?

Thm (Tsybakov 2008)

Assume $y_i := f^*(x_i) + \varepsilon_i$ $E[\varepsilon_i] = 0$

→ f^* Lipschitziana

$E[\varepsilon_i^2] = \sigma^2$

→ $x_i \sim \text{Unif}([0,1]^d)$

soft assumptions

Then: $\inf_{\hat{f}} \sup_{f^*} E[(\hat{f}(x_i) - f^*(x_i))^2] \gtrsim \frac{1}{n^{2/2+d}}$

$R(\hat{f}) < \sigma \gtrsim \frac{1}{n^{2/2+d}}$

CURSE OF DIMENSIONALITY

Model: We need assumptions

→ Data $p_{x,y}$
→ Hypothesis (Architecture)
→ Algorithms

} TYPICAL CASE ANALYSIS

Real question: How much data to learn a NN with SGD/GD

~ 2018 NTK Theory / Lazy Training (Jacobs et al. 2018)

ERM for NNs $\Leftrightarrow$ Kernel Methods

~ 2018-2022 Limitation of Kernels in HD
Thm (Mei, Montanari 2022)

Isotropic data $p_y \sim N(0, 1, S^{d-1}, [0, 1]^d)$

you have $n = \Theta(d^K)$, $K \gg 0$. Then a kernel can only learn a polynomial of degree K .

Beyond Kernels

Can we learn better than a kernel?

Focus on 2L NN

$$f_\theta(x) = \sum_{j=1}^{p_1} a_j \sigma(\langle \underset{w_j^T x}{w_j}, x \rangle)$$

Data: [Smooth index model]

Smooth Function

$$y_i = f_{\star}(x_i) + \epsilon$$

$$f_{\star}(x) = g(\langle w_1^{\star}, x \rangle, \dots, \langle w_r^{\star}, x \rangle)$$

ORTHOGONAL

$$= g(W^{\star} x)$$

$$\text{where } \langle w_k^{\star}, w_k^{\star} \rangle = \delta_{kk}$$

$$g: \mathbb{R}^r \rightarrow \mathbb{R}$$

MANIFOLD HYPOTHESIS $\rightarrow$ we hope

$$r \ll d$$

How much data do i need??

Algorithm: 2 step GD Training

1) Init (a^0, w^0)

2) Train W For T steps with fix a^0 .

3) Optimize over a for fixed W^T

~ 2023 (Ba et al., Dardi et. al.)

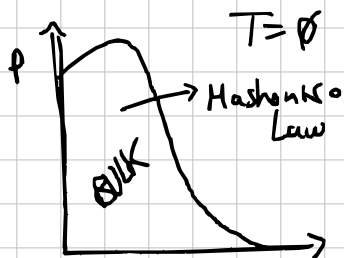
$W \in \mathbb{R}^{P \times d}$ Singular values of $\{ \lambda \}$

$$P(\lambda) = \frac{1}{P} \sum_{i=1}^P \delta(\lambda - \delta_i) \quad \delta_i \ll 1$$

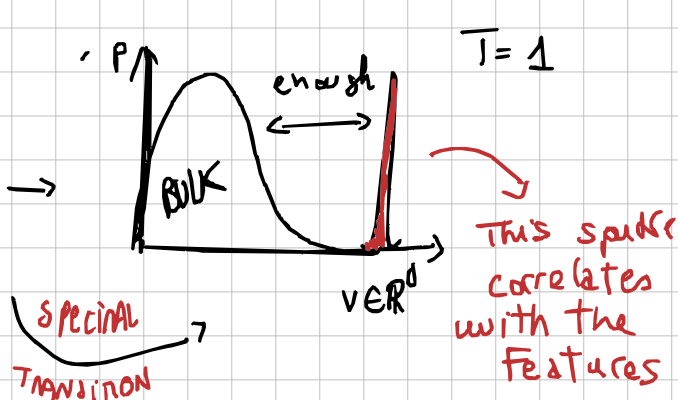
if i only
learn $W^{\star} x$,
a kernel
would be
enough

$\Rightarrow$ You get

Mashenko Pattern Law (RANDOM MATRIX Follow This)



RANDOM MATRICES
in HD



SPECIAL
TRANSITION

$$V \approx \sum_{k=1} \mu_k W_k^* \text{ Hermite Polynomial}$$

Thm. (Dandi et al. 2023; Cai et al. 2024)

After 1 step $n = O(d)$, f_θ can learn any non-linear component of f^*

$$f^*(x) \approx \varphi(\langle v, x \rangle) + \tilde{E}$$

CF Kernel $f^*(x) \approx \langle v, x \rangle + \tilde{E}$

[There is also an asymptote in that regime]

3) Quadratic NN

Question: Can we solve ERM problem for a 2L NN? I need to simplify more!!

Assume: $f_\theta(x)$ is quadratic (1L NN)

$$f_\theta(x) = \sum_{j=1}^p (\langle w_j, x \rangle^2 - \|w_j\|_2^2)$$

Data:

$$f_*(x) = \sum_{k=1}^r a_k^* \left(\langle w_k^*, x \rangle - \|w_k^*\|^2 \right)$$

$$x \sim N(0, \mathbb{I})$$

Note: $f_0(x) = \sum_{j=1}^p w_j^T (xx^T - \mathbb{I}) w_j$ Convex in S

$$= T_c(SG) \quad \text{where } S = W^T W \in \mathbb{R}^{d \times d}$$

$$= \langle S, G \rangle_F \quad G = xx^T - \mathbb{I}$$

→ Frobenius Norm

Similarly $f_*(x) = T_c(S_* G)$ where $S_* = W_*^T W_* \in \mathbb{R}^{d \times d}$

$$\therefore \text{ERM} \approx \min_{S \succeq 0} \frac{1}{n} \sum (\text{Tr}(S - S_*) G_i^T \epsilon)^2 + \lambda \text{Tr}(S)$$

$$\text{Tr}(S) = \sum_j \lambda_j$$

REGULARIZATION
[NUCLEAR NORM]

$\frac{2}{\sqrt{p}}$
L1
spectrum

differs differently induces SPARSITY than from original space

~ Results 2025 (Eriba et al. 2025)

Minimized when $d \rightarrow \infty$

$$\hat{S} = \text{soft+threshold} (S_* + dG) \xrightarrow{\text{Gaussian}}$$

