

Generalization Theory for Deep Learning



Basic version: finite hypothesis class

Finite hypothesis class: with probability $1 - \delta$ over the choice of a training set of size n , for a bounded loss ℓ , we have

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) - \mathbb{E}_{(x,y) \sim D} [\ell(f(x), y)] \right| = O \left(\sqrt{\frac{\log |\mathcal{F}| + \log 1/\delta}{n}} \right)$$

Pr: for a fixed f , by Hoeffding's inequality
w.p. $1 - \frac{\delta}{|\mathcal{F}|}$, $\text{gen_error}(f) = O \left(\sqrt{\frac{\log \left(\frac{|\mathcal{F}|}{\delta} \right)}{n}} \right)$

Union bound: event₁, ..., event_m

$$P \left(\bigcup_{i=1}^m \text{event}_i \right) \leq \sum_{i=1}^m P(\text{event}_i)$$

choose event_f: $\text{gen_error}(f) \geq \sqrt{\frac{\log \left(\frac{|\mathcal{F}|}{\delta} \right)}{n}}$

$$P \left(\bigcup_{f \in \mathcal{F}} \text{event}_f^+ \right) \leq \sum_{f \in \mathcal{F}} P(\text{event}_f^+) = |\mathcal{F}| \cdot \frac{\delta}{|\mathcal{F}|} = \delta$$

$$\Rightarrow \forall f, P \left(\text{gen_error}(f) < \sqrt{\frac{\log \left(\frac{|\mathcal{F}|}{\delta} \right)}{n}} \right) \geq 1 - \delta \quad \square$$

VC-Dimension

Motivation: Do we need to consider **every** classifier in $\mathcal{F}$?

Intuitively, **pattern of classifications** on the training set should suffice. (Two predictors that predict identically on the training set should generalize similarly).

Let $\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \{+1, -1\}\}$ be a class of binary classifiers.

The **growth function** $\Pi_{\mathcal{F}} : \mathbb{N} \rightarrow \mathbb{F}$ is defined as:

$$\Pi_{\mathcal{F}}(m) = \max_{(x_1, x_2, \dots, x_m)} \left| \left\{ (f(x_1), f(x_2), \dots, f(x_m)) \mid f \in \mathcal{F} \right\} \right|.$$

max : 2^m

The **VC dimension** of $\mathcal{F}$ is defined as:

$$\text{VCdim}(\mathcal{F}) = \max\{m : \Pi_{\mathcal{F}}(m) = 2^m\}.$$

VC-dimension Generalization bound

Theorem (Vapnik-Chervonenkis): with probability $1 - \delta$ over the choice of a training set, for a bounded loss ℓ , we have

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) - \mathbb{E}_{(x,y) \sim D} [\ell(f(x), y)] \right| = O \left(\sqrt{\frac{\text{VCdim}(\mathcal{F}) \log n + \log 1/\delta}{n}} \right)$$

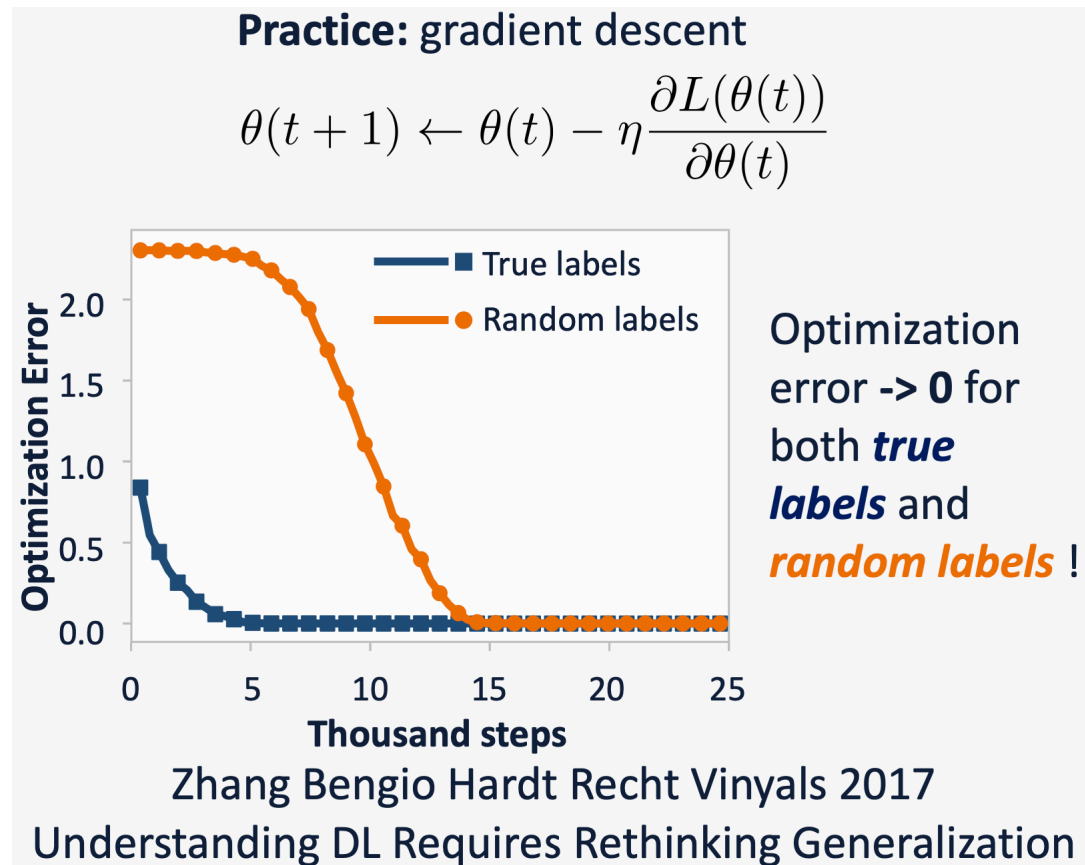
Examples:

- Linear functions: $\text{VC-dim} = O(\text{dimension})$
- Neural network: VC-dimension of fully-connected net with width W and H layers is $\widetilde{\Theta}(WH)$ (Bartlett et al., '17).

Problems with VC-dimension bound

if $n < \# \text{params}$

1. In over-parameterized regime, bound $\gg 1$.
2. Cannot explain the random noise phenomenon:
 - Neural networks that fit random labels and that fit true labels have the same VC-dimension.



PAC Bayesian Generalization Bounds

Setup: Let P be a prior over function in class $\mathcal{F}$, let Q be the posterior (after algorithm's training).

Theorem: with probability $1 - \delta$ over the choice of a training set, for a bounded loss ℓ , we have

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) - \mathbb{E}_{(x,y) \sim D} [\ell(f(x), y)] \right| = O \left(\sqrt{\frac{KL(Q || P) + \log 1/\delta}{n}} \right)$$

prior: NN's (universal)
gaussian

Q : gaussian

Rademacher Complexity

Intuition: how well can a classifier class **fit random noise**?

(Empirical) **Rademacher complexity:** For a training set $S = \{x_1, x_2, \dots, x_n\}$, and a class $\mathcal{F}$, denote:

$$\hat{R}_n(S) = \mathbb{E}_{\sigma} \sup_{f \in \mathcal{F}} \sum_{i=1}^n \sigma_i f(x_i) .$$

where $\sigma_i \sim \text{Unif}\{+1, -1\}$ (Rademacher R.V.).

(Population) **Rademacher complexity:**

$$R_n = \mathbb{E}_S \left[\hat{R}_n(s) \right] .$$

Rademacher Complexity Generalization Bound

Theorem: with probability $1 - \delta$ over the choice of a training set, for a bounded loss ℓ , we have

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) - \mathbb{E}_{(x,y) \sim D} [\ell(f(x), y)] \right| = O \left(\frac{\hat{R}_n}{n} + \sqrt{\frac{\log 1/\delta}{n}} \right)$$

$$\hat{R}_n \asymp \sqrt{H}$$

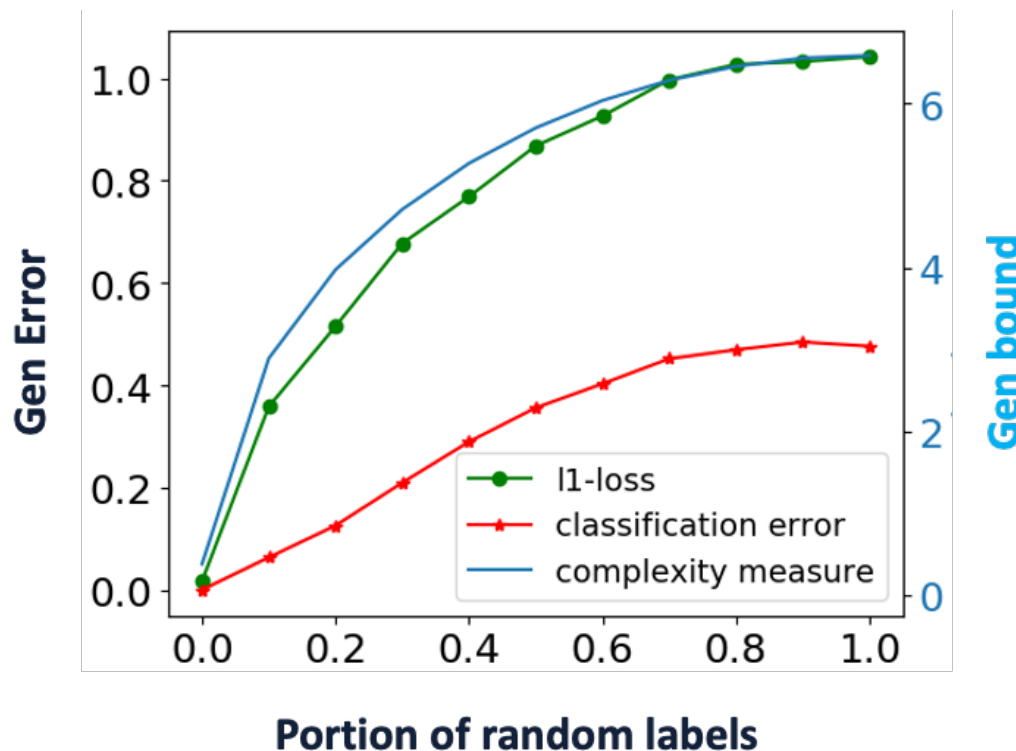
and

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) - \mathbb{E}_{(x,y) \sim D} [\ell(f(x), y)] \right| = O \left(\frac{R_n}{n} + \sqrt{\frac{\log 1/\delta}{n}} \right)$$

Kernel generalization bound

$$y: \mathbb{Q}^h$$

Use Rademacher complexity theory, we can obtain a generalization bound $O(\sqrt{y^\top (H^*)^{-1} y / n})$ where $y \in \mathbb{R}^n$ are n labels, and $H^* \in \mathbb{R}^{n \times n}$ is the kernel (e.g., NTK) matrix.



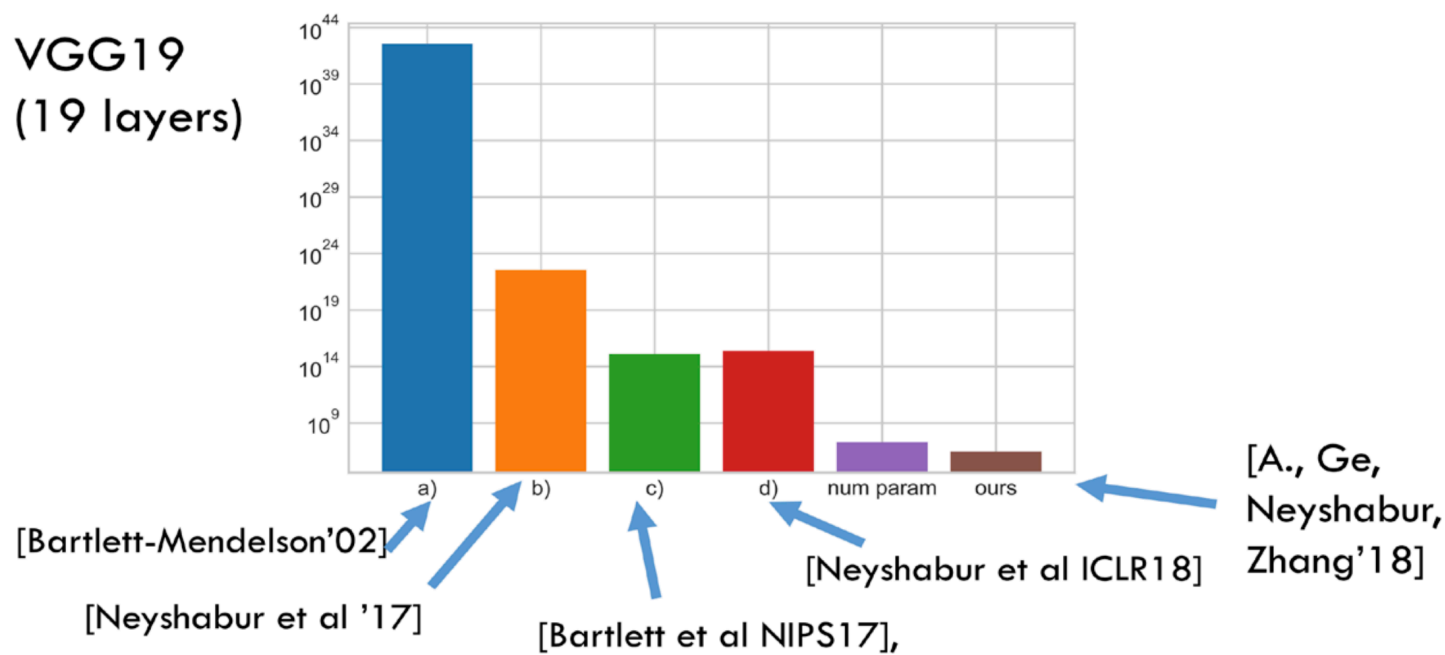
If $y \propto v_i$
 v_i : top eigenvectors
bound: $\sqrt{\frac{1}{6n}}$

Norm-based Rademacher complexity bound

Theorem: If the activation function σ is ρ -Lipschitz. Let $\mathcal{F} = \{x \mapsto W_{H+1}\sigma(W_h\sigma(\cdots\sigma(W_1x)\cdots), \|W_h^T\|_{1,\infty} \leq B \forall h \in [H]\}$ then $R_n(\mathcal{S}) \leq \|X^\top\|_{2,\infty} (2\rho B)^{H+1} \sqrt{2 \ln d}$ where $X = [x_1, \dots, x_n] \in \mathbb{R}^{d \times n}$ is the input data matrix.

Comments on generalization bounds

- When plugged in real values, the bounds are rarely non-trivial (i.e., smaller than 1)
- “*Fantastic Generalization Measures and Where to Find them*” by Jiang et al. '19 : large-scale investigation of the correlation of extant generalization measures with true generalization.



Comments on generalization bounds

- Uniform convergence may be unable to explain generalization of deep learning [Nagarajan and Kolter, '19]
 - Uniform convergence: a bound for all $f \in \mathcal{F}$
 - Exists example that 1) can generalize, 2) uniform convergence fails.

$$\text{bias}^2 + \text{var}$$

- Rates:
 - Most bounds: $1/\sqrt{n}$.
 - Local Rademacher complexity: $1/n$.

Separation between NN and kernel

- For approximation and optimization, neural network has no advantage over kernel. Why NN gives better performance: **generalization.**
- [Allen-Zhu and Li '20] Construct a class of functions $\mathcal{F}$ such that $y = f(x)$ for some $f \in \mathcal{F}$:
 - no kernel is sample-efficient; *need $\exp(d)$ sample*
 - Exists a neural network that is sample-efficient. *poly(d) sample*

Separation between NN and kernel

Defn Kernel method is linear method with an embedding
 $\phi: \mathbb{R}^d \rightarrow \mathcal{H}$, Hilbert space
 $\Rightarrow$ it turns an element $f \in \mathcal{H}$ into a prediction function
$$g = \langle f, \phi(x) \rangle$$
$$\stackrel{\circ}{=} f(x)$$

The method uses samples, n $\{x_i\}_{i=1}^n$, $x_i \in \mathbb{R}^d$
observed $\{y_i\}_{i=1}^n$

$$f \in \text{span}(\{\phi(x_i)\}_{i=1}^n)$$

Example: $\underset{f}{\text{argmin}} \frac{1}{n} \sum_{i=1}^n (y_i - \langle f, \phi(x_i) \rangle)^2 + \lambda \|f\|_{\mathcal{H}}^2$
 $\Rightarrow f \in \text{span}(\{\phi(x_i)\})$

Separation between NN and kernel

Thm $\exists$ a class of functions $\mathcal{C} \subseteq \{\mathbb{R}^d \rightarrow \mathbb{R}\}$
and a distribution over $\mathbb{R}^d$ s.t.

i) $\forall$ Kernel method, if it satisfies that

$\forall c \in \mathcal{C}$, given $y_i = c(x_i)$
if $\mathbb{E}_{\mathcal{U}} [(c(x) - \langle f, \phi(x) \rangle)^2] \leq \frac{1}{d-1}$
then you need $n \geq 2$

ii) $\exists$ simple procedure (not Kernel) that
you can output correct c , using $n = d$ samples

★ NN can simulate this procedure

Separation between NN and kernel

Pf: μ : uniform distribution over $\{-1, 1\}^d$

$$\mathcal{C} = \{C_S = \prod_{i \in S} x_i, S \subset \{1, \dots, d\}\}$$

Pf of part ii), choose $\begin{pmatrix} -1 \\ \vdots \\ 1 \end{pmatrix}, \begin{pmatrix} -1 \\ \vdots \\ 1 \end{pmatrix} \dots \begin{pmatrix} -1 \\ \vdots \\ 1 \end{pmatrix}$
 $e_1 \quad \quad \quad e_d$

$$\Rightarrow y_i = C_S(e_i), \text{ if } i \in S, C_S(e_i) = -1$$
$$\text{if } i \notin S, C_S(e_i) = 1$$

$\Rightarrow$ identify which i is in S

Separation between NN and kernel

part i) $\mathcal{C}$ is a basis for $\{f: \{-1, 1\}^d \rightarrow \mathbb{R}\}$
w.r.t. μ (by symmetry)

$$C_S, C_{S'} \quad \mathbb{E}_{X \sim \mu} [C_S(X) \cdot C_{S'}(X)] = \begin{cases} 0 & \text{if } S \neq S' \\ 1 & \text{if } S = S' \end{cases}$$

$$\Rightarrow \forall f, \exists \{a_i\}, f = \sum a_i \cdot C_{S_i}$$

$$\text{Goal: } \mathbb{E}_{X \sim \mu} \left[\left(C_{S^*}(X) - \langle f, \Phi(X) \rangle \right)^2 \right]$$

Since f is in span $\{ \Phi(x_i) \}$

$$f = \sum_{i=1}^n a_i \Phi(x_i), \quad f(x) = \sum_{i=1}^n a_i \langle \Phi(x_i), \Phi(x) \rangle$$

$$x \mapsto \langle \Phi(x_i), \Phi(x) \rangle$$

$$= \sum_{S \in [d]} \lambda_{i,S} C_S(X)$$

Separation between NN and kernel

$$\begin{aligned} & \mathbb{E}_{x \sim \mathcal{M}} \left[\left(f^*(x) - \langle f, \phi(x) \rangle \right)^2 \right] \quad (*) \\ &= \mathbb{E}_{x \sim \mathcal{M}} \left[\left(f^*(x) - \sum_{f \in \mathcal{H}} \sum_{i=1}^n a_i \lambda_{i,f} f(x) \right)^2 \right] \\ &= \left(1 - \sum_{i=1}^n a_i \lambda_{i,f^*} \right)^2 + \sum_{f \neq f^*} \left(\sum_{i=1}^n a_i \lambda_{i,f} \right)^2 \end{aligned}$$

by assumption $(*) \leq \frac{1}{9}$

$$\begin{aligned} \Rightarrow \left(1 - \sum_{i=1}^n a_i \lambda_{i,f^*} \right)^2 &\leq \frac{1}{9} \\ \sum_{f \neq f^*} \left(\sum_{i=1}^n a_i \lambda_{i,f} \right)^2 &\leq \frac{1}{9} \end{aligned}$$

Separation between NN and kernel

Notations:

$$\Lambda : 2^d \times 4$$

$$\Lambda_{S,i} = \lambda_{i,S}$$

$$A = n \times 2^d$$

$$A_{i,S^*} = \alpha_{i,S^*}$$

$$\Omega = \Lambda A : 2^d \times 2^d \text{ of rank at most } n$$

$$\Omega_{S^*,S} = \sum_{i=1}^n \alpha_{i,S^*} \lambda_{i,S}$$

$$\Rightarrow \Omega_{S^*,S} \leq \frac{2}{3}$$

$$\sum_{S \neq S^*} \Omega_{S,S^*}^2 \leq \frac{1}{9}$$



Separation between NN and kernel

$$\Omega = \text{diag}(\Omega) + \Omega' \quad , \quad \Omega' : \text{off-diagonal}$$

$$\|\Omega'\|_F^2 = \sum \text{eigen}^2(\Omega') \leq \frac{2^d}{9}$$

$\Rightarrow \Omega'$ has at most $\frac{2^d}{9}$ eigenvalues $\geq \frac{2}{3}$
 consider subspace with eigenvalue $< \frac{2}{3}$
 $\Rightarrow \dim \geq \frac{3}{4} \cdot 2^d$

$\forall x \in \text{this subspace}$

$$\|\Omega x\|_2 = \|\text{diag}(\Omega)x + \Omega'x\|_2$$

$$> \|\text{diag}(\Omega)x\|_2 - \|\Omega'x\|_2$$

$$\frac{2}{3} \|x\|_2 - \frac{2}{9} \|x\|_2 > 0$$

$$\text{rank}(\Omega) \geq \frac{3}{4} \cdot 2^d, \quad n \geq \frac{3}{4} \cdot 2^d \geq 2^{d-1} \quad \square$$