

9.520/6.7910: Statistical Learning Theory and Applications

- Class: **Tue, Thu 11:00 am - 12:30 pm**, 46-3002
- Office Hours: **Tue, Thu 1:30pm - 3:00pm, 46-5155D**
- Web: <https://poggio-lab.mit.edu/9-520/>
- Contact: 9.520@mit.edu, pierb@mit.edu
- 9.520/6.7910 will use Canvas: <https://canvas.mit.edu/courses/34219>
- Also check Canvas announcements for updates

9.520: the parts of the course

The first part is about:

- classical regularization and regularized least squares + kernel machines, SVM
- logistic regression, square and exponential loss
- large margin (minimum norm)
- stochastic gradient methods,
- overparametrization, implicit regularization
- approximation/estimation errors.

The second part is about deep networks and, particular, about:

- approximation theory -- why deep networks avoid the curse of dimensionality and are a universal parametric approximator s
- optimization theory -- how weights and activation evolve in time and across layers during training with SGD
- learning theory -- how generalization in deep networks can be explained in terms of the complexity control implicit in regularized (or unregularized) SGD and in terms of the details of the sparse compositional architecture itself.

In these modern times...
the theory we need should be concerned with
fundamental principles of intelligence,
not just theorems for the sake of math

The stunning success of deep learning (DL) over the past decade is a puzzling natural phenomenon of great scientific interest. This success has been amply demonstrated through a plethora of convincing experiments, and yet it is crying out for a thorough scientific understanding through foundational research (C.P.)

Conceptual Framework of Machine Learning

1. choose $\mathbb{H}_\lambda$, a space $\mathbb{H}_\lambda$ of functions with parameters λ - in which to search for f .
The functions in $\mathbb{H}_\lambda$ must have a *non-exponential* number of parameters (in the effective dimensionality) AND be powerful enough to approximate a broad range of target functions (realm of approximation theory);
2. optimize the values of the parametric approximator by minimizing a loss function on a training set (realm of optimization theory);
3. study generalization properties of the trained representation (domain of learning theory).

Conceptual Framework of Machine Learning

Note that the functions in $\mathbb{H}_\lambda$ should have a *non-exponential* number of parameters (in the effective dimensionality). This means that if I can find such a space AND if $\mathbb{H}_\lambda$ can approximate a broad range of target functions then I have made not only *approximation possible* but I have made also *optimization easier* and *generalization more likely*. The reason is # parameters.

In the modern era.. “new” puzzles about approximation, generalization, optimization

- No curse of dimensionality
- VC bounds overparametrization
- Optimization unreasonably easy, non-convex non-PL but optimization works

In reality...

The puzzles about *overparametrization* and *generalization* are not new (remember lorenzo with norms not # parameters)

Multiple layers ==> *compositional sparsity* is new

In the modern era.. “new” puzzles about approximation, generalization, optimization

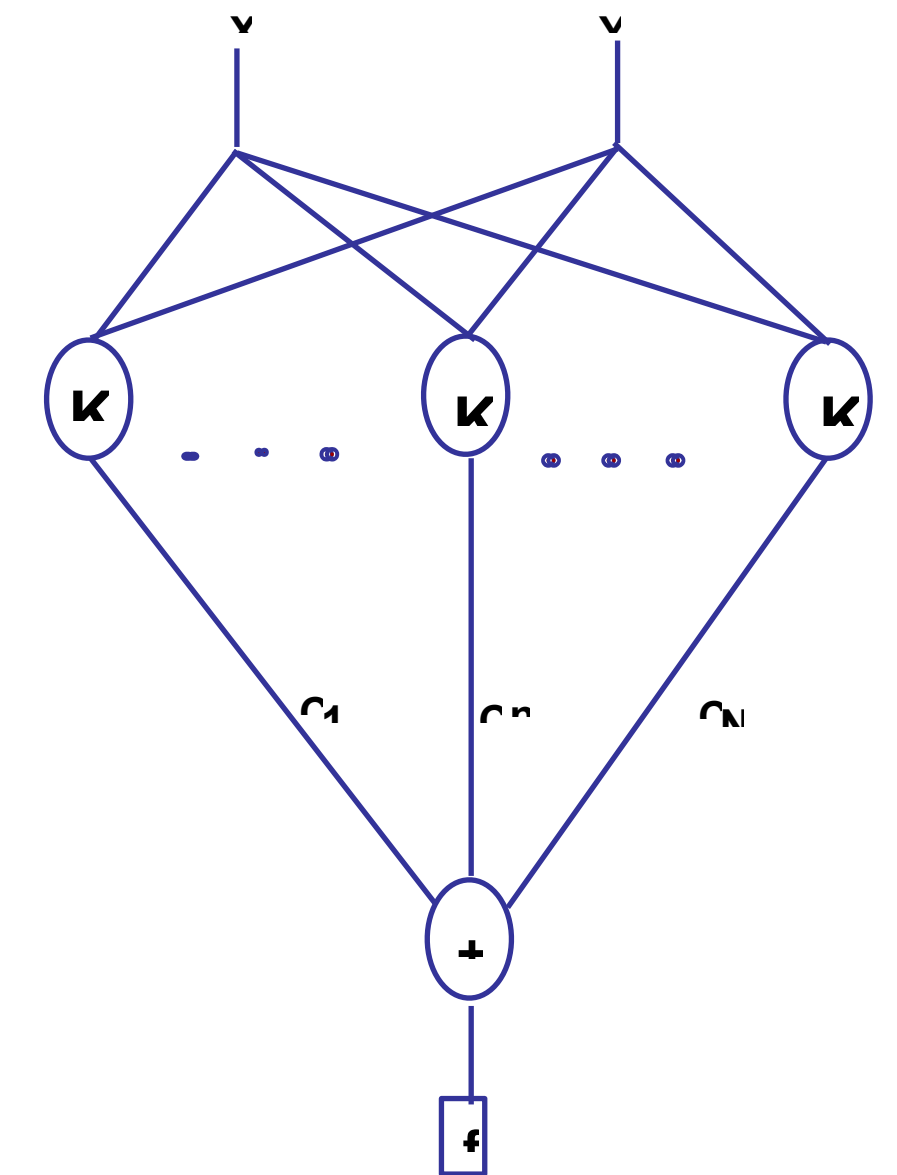
- No curse of dimensionality
- VC bounds overparametrization
- Optimization unreasonably easy, non-convex non-PL but optimization works relaxing adding neurons

Typical shallow networks: Kernel Machines

$$\min_{f \in H} \left[\frac{1}{\ell} \sum_{i=1}^{\ell} V(f(x_i) - y_i) + \lambda \|f\|_K^2 \right]$$

implies

$$f(\mathbf{x}) = \sum_i^l \alpha_i K(\mathbf{x}, \mathbf{x}_i)$$

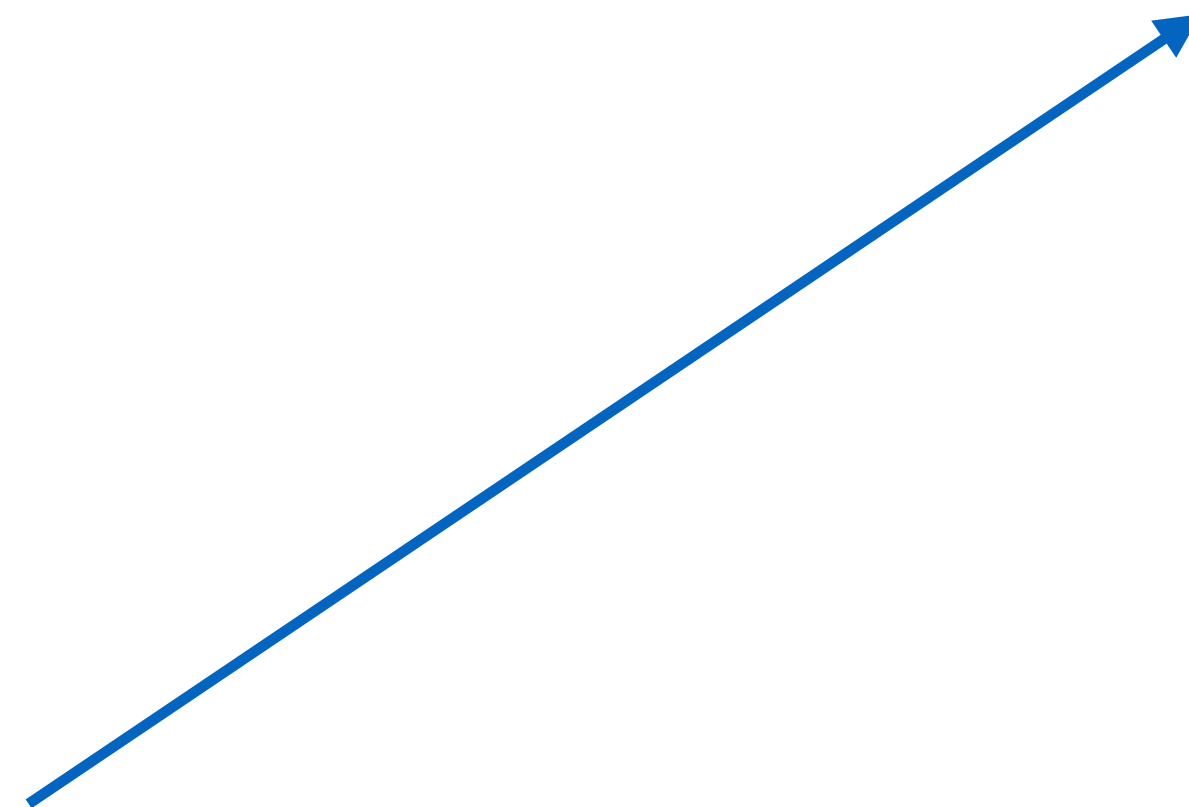


Equation includes splines, Radial Basis Functions and Support Vector Machines (depending on choice of V).

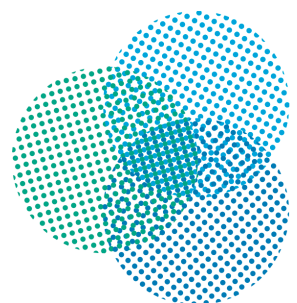
Degree (or rate) of approximation (theorem)

$$f(x_1, x_2, \dots, x_d)$$

Both shallow and deep network can approximate a function of d variables equally well (density). The number of parameters in both cases depends exponentially on d as $N = O(\epsilon^{-\frac{d}{s}})$ or $\epsilon = O(N^{-\frac{s}{d}})$



Curse of dimensionality



Key *sparse target function* property of the world

A generic function associated to the CIFAr problem $f : \mathcal{R}^{1000} \rightarrow \mathcal{R}$ requires a very large memory (10 bins per coordinate corresponds to 10^{1000} , numbers which is $\gg N_{Edd} \approx 10^{80}$ the number of all the protons in the universe...there are $\sim 10^{15}$ synapses in the human brain).

Outline: puzzles about approximation, optimization, generalization

- No Curse dimensionality (puzzle for approx/optimization)
- VC bounds are “very” vacuous for overparametrization
- Optimization unreasonably easy, non-convex non-PL but optimization works

Generalization bounds

In general generalization bounds are *quantitatively* meaningless but *qualitatively* useful.

- Quantitatively meaningless: because it is easier and better to do empirical estimates of test error on data not used in training
- Qualitatively useful because they often provide an insight into mechanisms of generalization

Rademacher

Rademacher Complexity Let $\mathbb{H}$ be a set of real-valued functions $h : \mathcal{X} \rightarrow \mathbb{R}$ defined over a set $\mathcal{X}$. Given a fixed sample $S \in \mathcal{X}^m$ the empirical Rademacher complexity of $\mathbb{H}$ is defined as follows:

$$\mathcal{R}_S(\mathbb{H}) := \frac{2}{m} \mathbb{E}_\sigma \left[\sup_{h \in \mathbb{H}} \left| \sum_{i=1}^m \sigma_i h(x_i) \right| \right].$$

The expectation is taken over $\sigma = (\sigma_1, \dots, \sigma_m)$, where, $\sigma_i \in \{\pm 1\}$ are i.i.d. and uniformly distributed samples.

From Class 6

$$\mathbb{E}_S \left[\sup_{f \in \sqrt{\frac{C_0}{\lambda}}} \left| \frac{1}{n} \sum_i l(y_i; f(x_i)) - \mathbb{E}_{x,y} [l(y, f(x))] \right| \right]$$

$$\leq 2 \mathbb{E}_{S, \sigma} \left[\sup_{f \in \sqrt{\frac{C_0}{\lambda}}} \left| \frac{1}{n} \sum_i \sigma_i l(y_i, f(x_i)) \right| \right]$$

Rademacher Complexity
of $\mathcal{H}$

Expected error

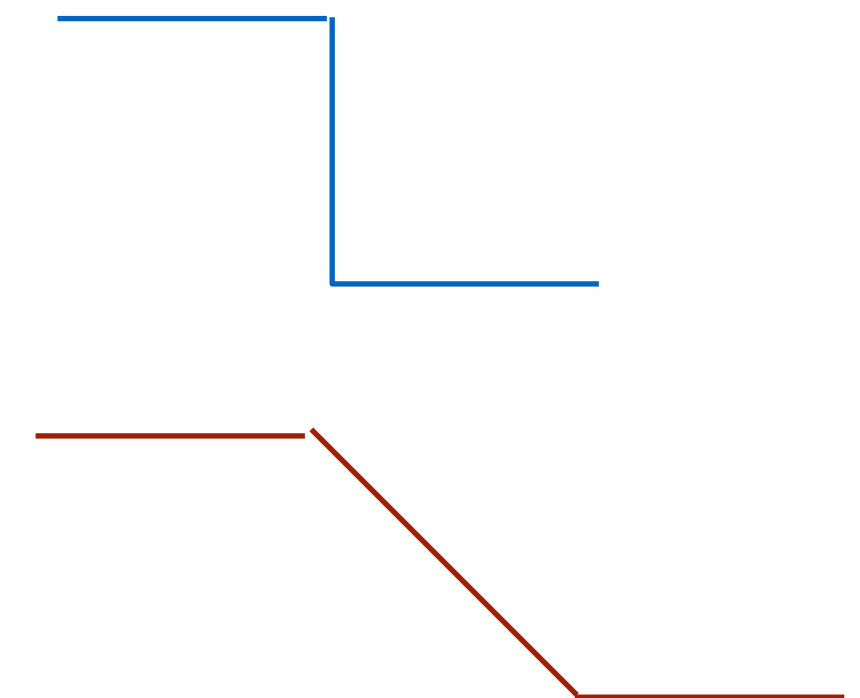
A typical generalization bound that holds with probability at least $(1 - \delta)$,
 $\forall f \in \mathbb{F}$ has the form:

$$|L(g) - \hat{L}(g)| \leq \frac{1}{\gamma} \mathbb{R}_N(\mathbb{G}) + c_2 \sqrt{\frac{\ln(\frac{1}{\delta})}{2N}}$$

where $L(g) = \mathbf{E}[\ell_{\text{gamma}}(g(x), y)]$ is the expected loss and
 $\hat{L}(g) = \frac{1}{N} \sum_n \ell_{\text{gamma}}(g(x_n), y_n)$ is the empirical loss

with

$$\ell_{\text{gamma}}(y, y') = \begin{cases} 1, & \text{if } yy' \leq 0, \\ 1 - \frac{yy'}{\gamma}, & \text{if } 0 \leq yy' \leq \gamma, \\ 0, & \text{if } yy' \geq \gamma. \end{cases}$$



Rademacher Bounds

Theorem Let P be a distribution over $\mathbb{R}^d \times \{\pm 1\}$. Let $\mathbb{F} = \{f_W \mid \prod_{i=1}^L \|W_i\| \leq 1\}$. Let $g(x) = \rho f(x)$ and $\mathcal{S} = \{(x_i, y_i)\}_{i=1}^N$ be a dataset of i.i.d. samples selected from P . Then, with probability at least $1 - \delta$ over the selection of $\mathcal{S}$, for any g_W , we have

$$|L(g) - \hat{L}(g)| \leq 2\mathcal{R}_{\mathcal{S}}(\mathbb{G}) + O\sqrt{\frac{1}{\delta N}}$$

Since $\mathcal{R}(\mathbb{G}) = \rho\mathcal{R}(\mathbb{F}) = \rho\frac{\mathcal{R}(\mathbb{F})}{N}$ because $g = \rho f$ we have

$$|L(g) - \hat{L}(g)| \leq 2\rho\mathcal{R}_{\mathcal{S}}(\mathbb{F}) + O\sqrt{\frac{1}{\delta N}}$$

Parenthesis

The key observation is that the Rademacher complexity satisfies the property:

$\mathbb{R}_N(\mathbb{G}) = \rho_1 \cdots \rho_K \mathbb{R}_N(\mathbb{F})$ where $g(x)$ is the function represented by the network and $f(x)$ is the function with normalized weight matrices.

Then, for the unnormalized cross-entropy loss the bound gives

$$L(g) \leq \prod_{k=1}^K \rho_k \mathbb{R}_N(\mathbb{F}) + c_2 \sqrt{\frac{\ln\left(\frac{1}{\delta}\right)}{2N}},$$

since $\hat{L}(g) \approx 0$, whereas for the normalized cross-entropy loss it yields

$$L(f) \leq \hat{L}(f) + \mathbb{R}_N(\mathbb{F}) + c_2 \sqrt{\frac{\ln\left(\frac{1}{\delta}\right)}{2N}}$$

Parenthesis

Comparison

- VC bound: combinatorial, worst-case;
- Rademacher bound: data-dependent, usually tighter;
- Relation: $VC(\mathcal{H}) = d$ implies $\mathfrak{R}_n(\mathcal{H}) = O(\sqrt{d \log n/n})$.

VC Dimension of Deep Networks

* Neural networks with W parameters and piecewise-linear activations (e.g. ReLU):

$$\text{VC}(\mathcal{H}) = \Theta(W \log W).$$

* Sample complexity:

$$n = O\left(\frac{W \log W + \log(1/\delta)}{\epsilon^2}\right).$$

Upper direction: metric entropy $\Rightarrow$ small chains $\Rightarrow$ small $\mathbb{R}ad$ (Dudley).

Lower direction: large $\mathbb{R}ad$ forces large entropy at some scale (Sudakov).

Hence, **they bound each other** up to constants

Another footnote

Training is a minimization

**In Class 4 we saw when optimization
algorithm are meant to work and why**

Both these two things break
for Neural Networks:

- The convexity.
- The existence of B .

- Let us now see *how*:

Both these two things break
for Neural Networks:

- The convexity.
- The existence of B .

