

The Mathematics of Large Machine Learning Models

Andrea Montanari

Stanford University

August 11, 2025

Prelude: The role of theory

Current state of Machine Learning

I started working in high-dim statistics and ML and around 2008

- ▶ **2008:** $\approx 1.5 \cdot 10^3$ papers submitted each year in top 3 venues
- ▶ **2024:** $\approx 3 \cdot 10^4$ papers submitted each year in top 3 venues
- ▶ Major discoveries
 - ▶ Machines can learn language
 - ▶ Learn from examples during conversations
 - ▶ Exhibit step-by-step problem-solving

Current state of Machine Learning

I started working in high-dim statistics and ML and around 2008

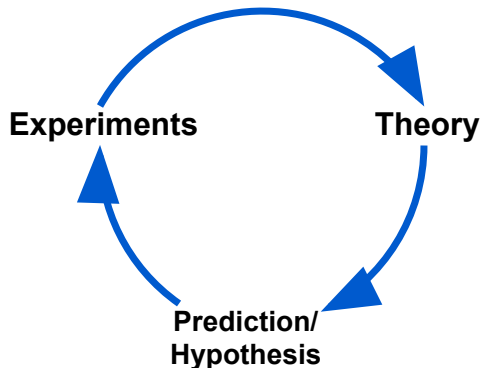
- ▶ **2008:** $\approx 1.5 \cdot 10^3$ papers submitted each year in top 3 venues
- ▶ **2024:** $\approx 3 \cdot 10^4$ papers submitted each year in top 3 venues
- ▶ Major discoveries
 - ▶ Machines can learn language
 - ▶ Learn from examples during conversations
 - ▶ Exhibit step-by-step problem-solving

Current state of Machine Learning

I started working in high-dim statistics and ML and around 2008

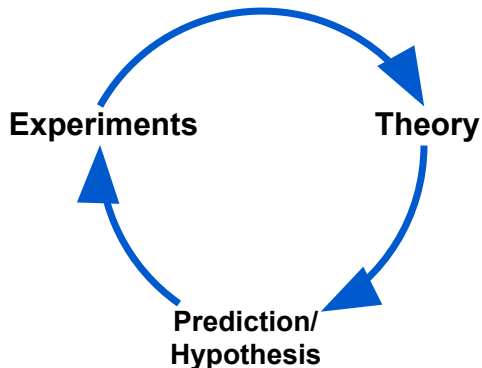
- ▶ **2008:** $\approx 1.5 \cdot 10^3$ papers submitted each year in top 3 venues
- ▶ **2024:** $\approx 3 \cdot 10^4$ papers submitted each year in top 3 venues
- ▶ Major discoveries
 - ▶ Machines can learn language
 - ▶ Learn from examples during conversations
 - ▶ Exhibit step-by-step problem-solving

A success of the scientific method?



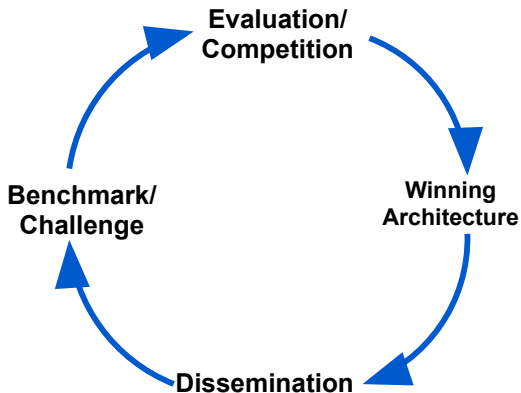
Not really!

A success of the scientific method?

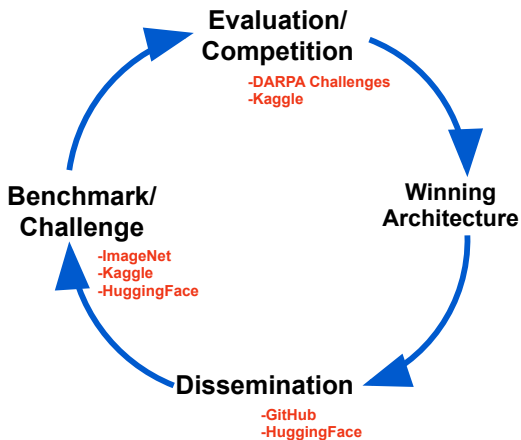


Not really!

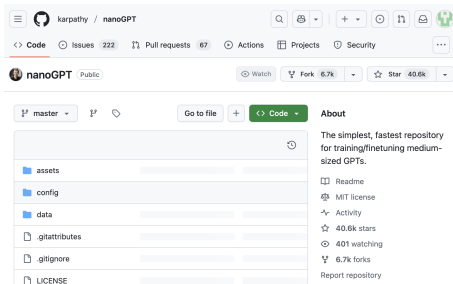
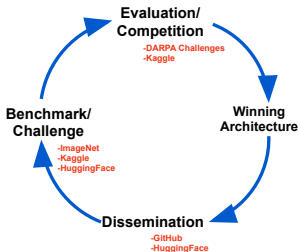
The 'AI research method'



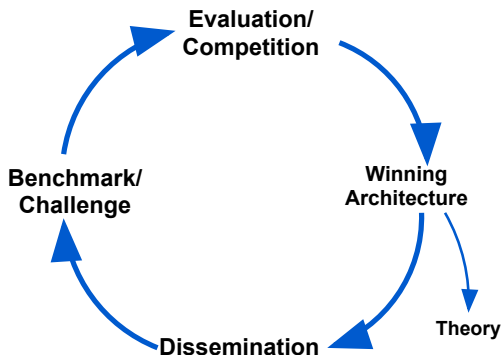
The AI research method and its institutions



The AI research method and its institutions

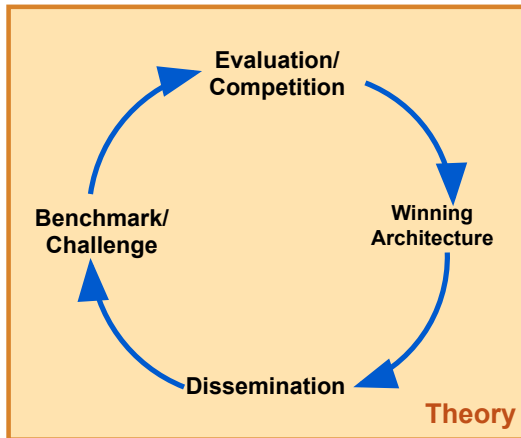


And theory?



- Explain why architecture X is currently winning the race

And theory?



- ▶ Ask very fundamental questions about very general phenomena
 - ▶ Does more data always help? How much?
 - ▶ Do more complex models always help? How much?

A few theoretical successes

- ▶ Overfitting and generalization
- ▶ Scaling
- ▶ Generic phenomena in (stochastic) gradient descent
- ▶ Benchmarking/uncertainty quantification
- ▶ ...

A few theoretical successes

- ▶ **Overfitting and generalization**
- ▶ Scaling
- ▶ Generic phenomena in (stochastic) gradient descent
- ▶ Benchmarking/uncertainty quantification
- ▶ ...

Outline

- 1 The generalization problem
- 2 Enters modern machine learning
- 3 Two previews
- 4 Conclusion

The generalization problem

History

Before 1960:

- ▶ A program defines the input-output relation

Rosenblatt, 1960:

- ▶ Input-output pairs $(\mathbf{x}_i, y_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: ‘Perceptron’

History

Before 1960:

- ▶ A program defines the input-output relation

Rosenblatt, 1960:

- ▶ Input-output pairs $(\mathbf{x}_i, y_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: ‘Perceptron’

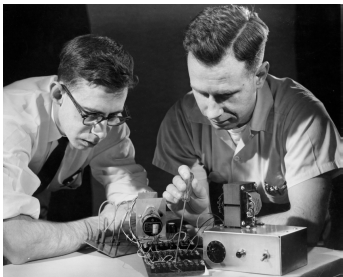
History

Before 1960:

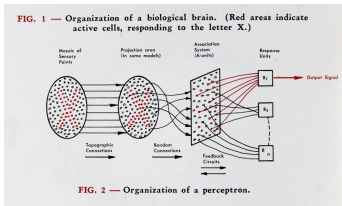
- ▶ A program defines the input-output relation

Rosenblatt, 1960:

- ▶ Input-output pairs $(\mathbf{x}_i, \mathbf{y}_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: ‘Perceptron’



Frank Rosenblatt



Perceptron

History

Before 1960:

- ▶ A program defines the input-output relation

Rosenblatt, 1960:

- ▶ Input-output pairs $(\mathbf{x}_i, y_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: 'Perceptron'

Why does it work? What does it mean to learn?

History

Before 1960:

- ▶ A computer is programmed
- ▶ A program needs to determine the input-output relation

Rosenblatt 1960:

- ▶ Show the machine input-output pairs $(\mathbf{x}_i, y_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: Perceptron

Vapnik–Chervonenkis, 1969-1974:

- ▶ Learning is statistical learning

History

Before 1960:

- ▶ A computer is programmed
- ▶ A program needs to determine the input-output relation

Rosenblatt 1960:

- ▶ Show the machine input-output pairs $(\mathbf{x}_i, y_i)$, $i \leq n$
- ▶ Let the machine **learn** the relation
- ▶ Proof of concept: Perceptron

Vapnik–Chervonenkis, 1969-1974:

- ▶ Learning is statistical learning

Statistical learning

- ▶ **Data:** $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \sim_{\text{iid}} \mathbb{P}$
- ▶ **Input-output relations:** $f(\cdot; \boldsymbol{\theta}_1), f(\cdot; \boldsymbol{\theta}_2), \dots, f(\cdot; \boldsymbol{\theta}_M).$

Statistical learning

- ▶ **Data:** $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \sim_{\text{iid}} \mathbb{P}$
- ▶ **Models:** $f(\cdot; \boldsymbol{\theta}_1), f(\cdot; \boldsymbol{\theta}_2), \dots, f(\cdot; \boldsymbol{\theta}_M).$
- ▶ **Learning:** $\mathcal{R}(\boldsymbol{\theta}) := \mathbb{E}\{\text{dist}(y_{n+1}, f(\mathbf{x}_{n+1}; \boldsymbol{\theta}))\}$:
 - ▶ $\hat{\boldsymbol{\theta}}$ selected from data $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$
 - ▶ Minimize test error $\mathcal{R}(\boldsymbol{\theta})!$

Statistical learning

- ▶ **Data:** $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \sim_{\text{iid}} \mathbb{P}$
- ▶ **Models:** $f(\cdot; \boldsymbol{\theta}_1), f(\cdot; \boldsymbol{\theta}_2), \dots, f(\cdot; \boldsymbol{\theta}_M).$
- ▶ **Learning:** $\mathcal{R}(\boldsymbol{\theta}) := \mathbb{E}\{\text{dist}(y_{n+1}, f(\mathbf{x}_{n+1}; \boldsymbol{\theta}))\}$:
 - ▶ $\hat{\boldsymbol{\theta}}$ selected from data $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$
 - ▶ Minimize test error $\mathcal{R}(\boldsymbol{\theta})!$

Statistical learning

- ▶ **Data:** $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n) \sim_{\text{iid}} \mathbb{P}$
- ▶ **Models:** $f(\cdot; \boldsymbol{\theta}_1), f(\cdot; \boldsymbol{\theta}_2), \dots, f(\cdot; \boldsymbol{\theta}_M).$
- ▶ **Learning:** $\mathcal{R}(\boldsymbol{\theta}) := \mathbb{E}\{\text{dist}(y_{n+1}, f(\mathbf{x}_{n+1}; \boldsymbol{\theta}))\}$:
 - ▶ $\hat{\boldsymbol{\theta}}$ selected from data $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$
 - ▶ Minimize test error $\mathcal{R}(\boldsymbol{\theta})!$

Statistical learning

$\text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta}))$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
$(\mathbf{y}_1, \mathbf{x}_1)$	0	1	$\cdots$	0
$(\mathbf{y}_2, \mathbf{x}_2)$	1	0	$\cdots$	1
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
$(\mathbf{y}_n, \mathbf{x}_n)$	0	1	$\cdots$	1
$(\mathbf{y}_{n+1}, \mathbf{x}_{n+1})$	1	0	$\cdots$	1

	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
empirical risk	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_1)$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_2)$	$\cdots$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_M)$
population risk	$\mathcal{R}(\boldsymbol{\theta}_1)$	$\mathcal{R}(\boldsymbol{\theta}_2)$	$\cdots$	$\mathcal{R}(\boldsymbol{\theta}_M)$

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta})), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_{n+1}, \mathbf{f}(\mathbf{x}_{n+1}; \boldsymbol{\theta})).$$

Statistical learning

$\text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta}))$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
$(\mathbf{y}_1, \mathbf{x}_1)$	0	1	$\cdots$	0
$(\mathbf{y}_2, \mathbf{x}_2)$	1	0	$\cdots$	1
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
$(\mathbf{y}_n, \mathbf{x}_n)$	0	1	$\cdots$	1
$(\mathbf{y}_{n+1}, \mathbf{x}_{n+1})$	1	0	$\cdots$	1

	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
empirical risk	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_1)$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_2)$	$\cdots$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_M)$
population risk	$\mathcal{R}(\boldsymbol{\theta}_1)$	$\mathcal{R}(\boldsymbol{\theta}_2)$	$\cdots$	$\mathcal{R}(\boldsymbol{\theta}_M)$

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta})), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_{n+1}, \mathbf{f}(\mathbf{x}_{n+1}; \boldsymbol{\theta})).$$

Statistical learning

$\text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta}))$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
$(\mathbf{y}_1, \mathbf{x}_1)$	0	1	$\cdots$	0
$(\mathbf{y}_2, \mathbf{x}_2)$	1	0	$\cdots$	1
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
$(\mathbf{y}_n, \mathbf{x}_n)$	0	1	$\cdots$	1
$(\mathbf{y}_{n+1}, \mathbf{x}_{n+1})$	1	0	$\cdots$	1

	$\mathbf{f}(\cdot; \boldsymbol{\theta}_1)$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_2)$	$\cdots$	$\mathbf{f}(\cdot; \boldsymbol{\theta}_M)$
empirical risk	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_1)$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_2)$	$\cdots$	$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}_M)$
population risk	$\mathcal{R}(\boldsymbol{\theta}_1)$	$\mathcal{R}(\boldsymbol{\theta}_2)$	$\cdots$	$\mathcal{R}(\boldsymbol{\theta}_M)$

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, \mathbf{f}(\mathbf{x}_i; \boldsymbol{\theta})), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_{n+1}, \mathbf{f}(\mathbf{x}_{n+1}; \boldsymbol{\theta})).$$

Statistical learning

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i)), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i; \boldsymbol{\theta})).$$

- ▶ Central limit theorem: $|\mathcal{R}(\boldsymbol{\theta}_i) - \hat{\mathcal{R}}_n(\boldsymbol{\theta}_i)| = O(1/\sqrt{n})$
- ▶ Vapnik–Chervonenkis, 1969-1974: $\mathcal{F} := \{f(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \mathcal{S}\}$

$$\max_{\boldsymbol{\theta} \in \mathcal{S}} |\hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta})| \leq \text{const.} \sqrt{\frac{\text{dvc}(\mathcal{F})}{n}}$$

Statistical learning

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i)), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i; \boldsymbol{\theta})).$$

- ▶ Central limit theorem: $|\mathcal{R}(\boldsymbol{\theta}_i) - \hat{\mathcal{R}}_n(\boldsymbol{\theta}_i)| = O(1/\sqrt{n})$
- ▶ Vapnik–Chervonenkis, 1969-1974: $\mathcal{F} := \{f(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \mathcal{S}\}$

$$\max_{\boldsymbol{\theta} \in \mathcal{S}} |\hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta})| \leq \text{const.} \sqrt{\frac{\text{dvc}(\mathcal{F})}{n}}$$

Statistical learning

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i)), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i; \boldsymbol{\theta})).$$

- ▶ Central limit theorem: $|\mathcal{R}(\boldsymbol{\theta}_i) - \hat{\mathcal{R}}_n(\boldsymbol{\theta}_i)| = O(1/\sqrt{n})$
- ▶ Vapnik–Chervonenkis, 1969-1974: $\mathcal{F} := \{f(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \mathcal{S}\}$

$$\max_{\boldsymbol{\theta} \in \mathcal{S}} |\hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta})| \leq \text{const.} \sqrt{\frac{\text{d}_{\text{VC}}(\mathcal{F})}{n}}$$

Statistical learning

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{n} \sum_{i=1}^n \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i)), \quad \mathcal{R}(\boldsymbol{\theta}) = \mathbb{E} \text{dist}(\mathbf{y}_i, f(\mathbf{x}_i; \boldsymbol{\theta})).$$

- ▶ Central limit theorem: $|\mathcal{R}(\boldsymbol{\theta}_i) - \hat{\mathcal{R}}_n(\boldsymbol{\theta}_i)| = O(1/\sqrt{n})$
- ▶ Vapnik–Chervonenkis, 1969-1974: $\mathcal{F} := \{f(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \mathcal{S}\}$

$$\max_{\boldsymbol{\theta} \in \mathcal{S}} |\hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta})| \lesssim \sqrt{\frac{d_{\text{VC}}(\mathcal{F})}{n}}$$

Key insight: Uniform CLT-style bounds

$$\max_{\boldsymbol{\theta} \in \mathcal{S}} |\hat{\mathcal{R}}_n(\boldsymbol{\theta}) - \mathcal{R}(\boldsymbol{\theta})| \lesssim \sqrt{\frac{d_{\text{VC}}(\mathcal{F})}{n}}$$

- ▶ Simultaneously for infinitely many $\boldsymbol{\theta}$'s!
- ▶ Bound does not depend on the number of parameters p ($\boldsymbol{\theta} \in \mathbb{R}^p$)
- ▶ Only depends on the function class $\mathcal{F}$ via $d_{\text{VC}}(\mathcal{F})$.

Key insight: Uniform CLT-style bounds

$$\max_{f \in \mathcal{F}} |\hat{\mathcal{R}}_n(f) - \mathcal{R}(f)| \lesssim \sqrt{\frac{d_{\text{VC}}(\mathcal{F})}{n}}$$

- ▶ Simultaneously for infinitely many θ 's!
- ▶ Bound does not depend on the number of parameters p ($\theta \in \mathbb{R}^p$)
- ▶ Only depends on the function class $\mathcal{F}$ via $d_{\text{VC}}(\mathcal{F})$.

A couple of examples

Example 1: Linear model

$$\mathcal{F}_{\text{lin}}(\rho) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \langle \boldsymbol{\theta}, \mathbf{x} \rangle : \|\boldsymbol{\theta}\| \leq \rho \right\},$$
$$\max_{f \in \mathcal{F}_{\text{lin}}(\rho)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \rho \sqrt{\frac{d}{n}}.$$

Example 2: 2-layer networks

[Bartlett 1996]

$$\mathcal{F}_{2\text{lin}}(\gamma) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \sum_{i=1}^N \mathbf{a}_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle) : \|\mathbf{w}_i\| = 1 \forall i \in [N], \|\mathbf{a}\|_1 \leq \gamma \right\},$$
$$\max_{f \in \mathcal{F}_{2\text{lin}}(\gamma)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \gamma \sqrt{\frac{d}{n}}.$$

Independent of N!

A couple of examples

Example 1: Linear model

$$\mathcal{F}_{\text{lin}}(\rho) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \langle \boldsymbol{\theta}, \mathbf{x} \rangle : \|\boldsymbol{\theta}\| \leq \rho \right\},$$
$$\max_{f \in \mathcal{F}_{\text{lin}}(\rho)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \rho \sqrt{\frac{d}{n}}.$$

Example 2: 2-layer networks

[Bartlett 1996]

$$\mathcal{F}_{2\text{lin}}(\gamma) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \sum_{i=1}^N \mathbf{a}_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle) : \|\mathbf{w}_i\| = 1 \forall i \in [N], \|\mathbf{a}\|_1 \leq \gamma \right\},$$
$$\max_{f \in \mathcal{F}_{2\text{lin}}(\gamma)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \gamma \sqrt{\frac{d}{n}}.$$

Independent of N!

A couple of examples

Example 1: Linear model

$$\mathcal{F}_{\text{lin}}(\rho) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \langle \boldsymbol{\theta}, \mathbf{x} \rangle : \|\boldsymbol{\theta}\| \leq \rho \right\},$$
$$\max_{f \in \mathcal{F}_{\text{lin}}(\rho)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \rho \sqrt{\frac{d}{n}}.$$

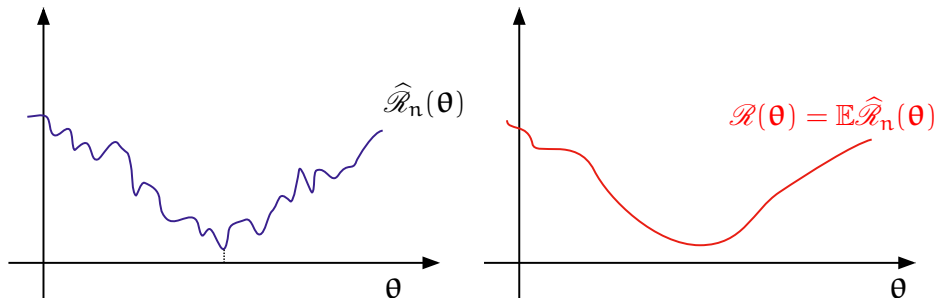
Example 2: 2-layer networks

[Bartlett 1996]

$$\mathcal{F}_{2\text{lin}}(\gamma) := \left\{ f(\mathbf{x}; \boldsymbol{\theta}) = \sum_{i=1}^N a_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle) : \|\mathbf{w}_i\| = 1 \forall i \in [N], \|\mathbf{a}\|_1 \leq \gamma \right\},$$
$$\max_{f \in \mathcal{F}_{2\text{lin}}(\gamma)} \left| \widehat{\mathcal{R}}_n(f) - \mathcal{R}(f) \right| \lesssim \gamma \sqrt{\frac{d}{n}}.$$

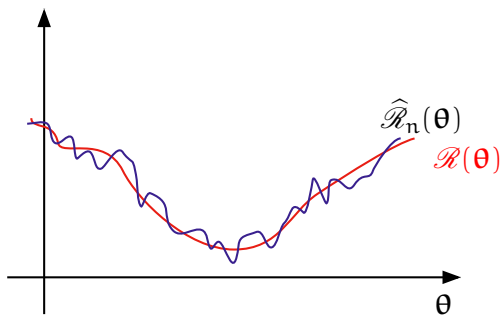
Independent of N!

Statistical learning: Implications



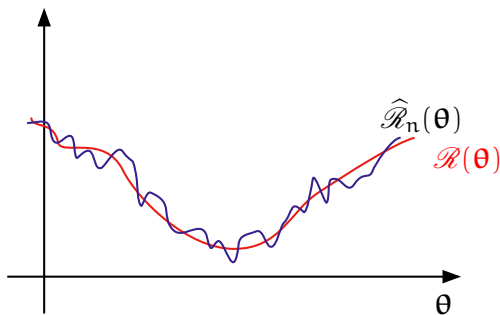
- ▶ Want to minimize $\mathcal{R}(\theta) := \mathbb{E}\{\ell(y_{n+1}, f(x_{n+1}; \theta))\}$ $\ell = \text{dist}$
- ▶ We have access to $\hat{\mathcal{R}}_n(\theta) := n^{-1} \sum_{i \leq n} \ell(y_i, f(x_i; \theta))$

Statistical learning: Implications



► Vapnik–Chervonenkis

Statistical learning: Implications

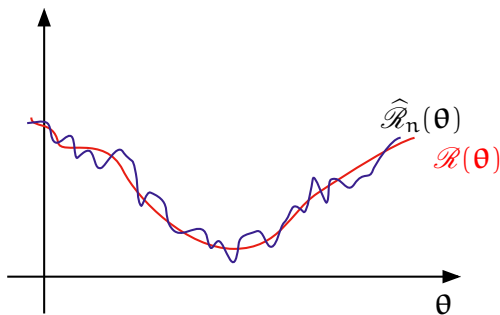


- Vapnik–Chervonenkis:

$$\max_{\theta \in \mathcal{S}} |\mathcal{R}(\theta) - \hat{\mathcal{R}}_n(\theta)| = O\left(\sqrt{\frac{d_{VC}(\mathcal{F})}{n}}\right)$$

- $\Rightarrow$ Can minimize $\hat{\mathcal{R}}_n(\theta)$ instead of $\mathcal{R}(\theta)$
- Choose $\mathcal{S}$ so that $d_{VC}(\mathcal{S}) \ll n$

Statistical learning: Implications

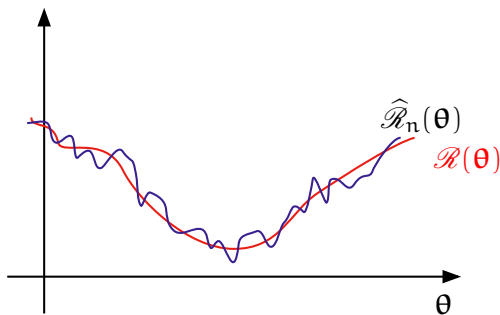


- Vapnik–Chervonenkis:

$$\max_{\theta \in \mathcal{S}} |\mathcal{R}(\theta) - \hat{\mathcal{R}}_n(\theta)| = O\left(\sqrt{\frac{d_{VC}(\mathcal{F})}{n}}\right)$$

- $\Rightarrow$ Can minimize $\hat{\mathcal{R}}_n(\theta)$ instead of $\mathcal{R}(\theta)$
- Choose $\mathcal{S}$ so that $d_{VC}(\mathcal{S}) \ll n$

Statistical learning: Implications

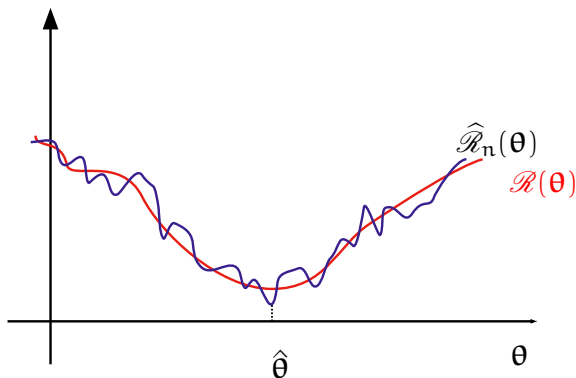


- Vapnik–Chervonenkis:

$$\max_{\theta \in \mathcal{S}} |\mathcal{R}(\theta) - \hat{\mathcal{R}}_n(\theta)| = O\left(\sqrt{\frac{d_{VC}(\mathcal{F})}{n}}\right)$$

- $\Rightarrow$ Can minimize $\hat{\mathcal{R}}_n(\theta)$ instead of $\mathcal{R}(\theta)$
- Choose $\mathcal{S}$ so that $d_{VC}(\mathcal{S}) \ll n$

Statistical learning

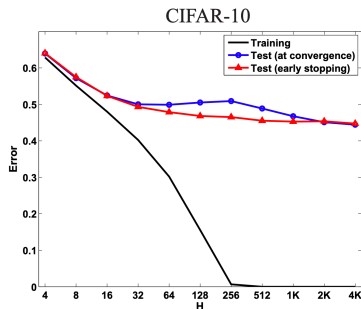
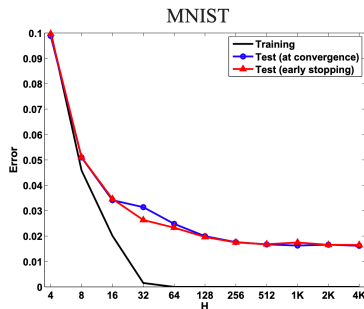


► Learned model $\hat{\theta}$

$$\text{Generalization Error} := \mathcal{R}(\hat{\theta}) - \hat{\mathcal{R}}_n(\hat{\theta})$$

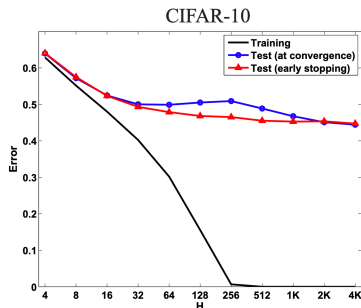
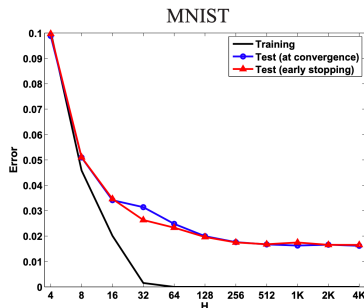
Enters modern machine learning

A small experiment



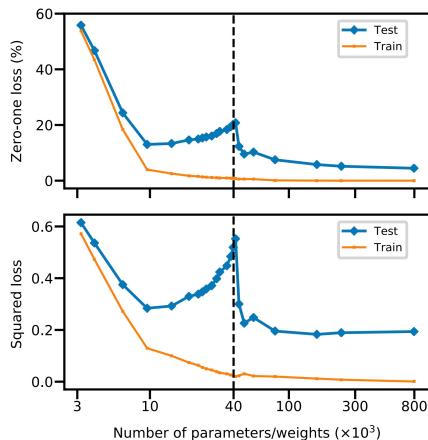
- ▶ x -axis $\propto$ Number of parameters
- ▶ Converges to vanishing training error
- ▶ Resulting weights depend on the initialization

A small experiment

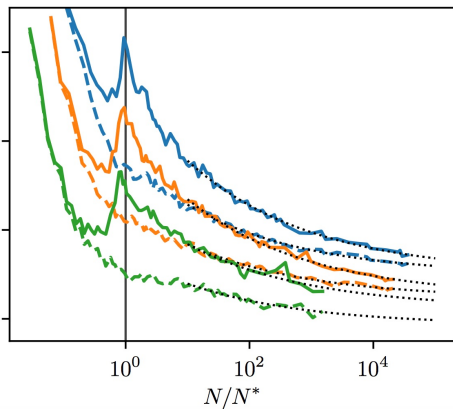


- ▶ x -axis $\propto$ Number of parameters
- ▶ Converges to vanishing training error
- ▶ Resulting weights depend on the initialization

Another version of the same experiment

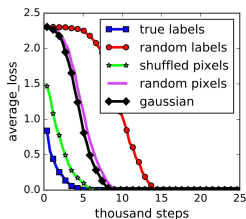


Yet another version

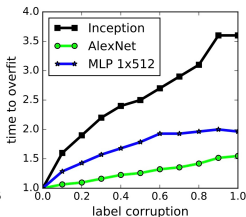


MNIST: 50,000 images in 2 different classes. 5-layers Neural Net. Quadratic hinge loss. Spigler, Geiger, d'Ascoli, Sagun, Biroli, Wyart, 2018

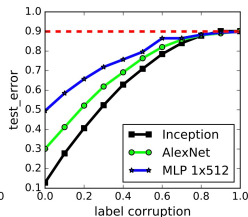
A larger scale experiment



(a) learning curves



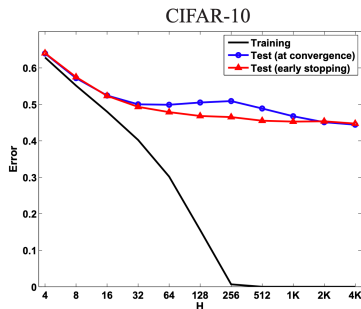
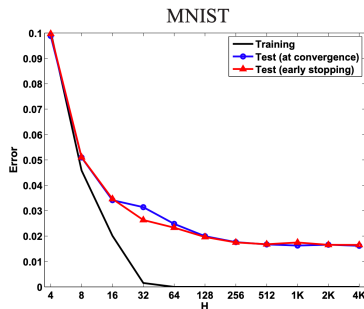
(b) convergence slowdown



(c) generalization error growth

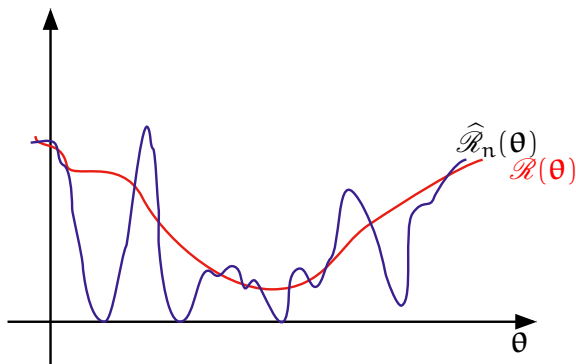
- ▶ Model complex enough to ‘interpolate’ random labels
- ▶ Despite this, does well on uncorrupted test samples
- ▶ Test error $\gg$ Train error ≈ 0

What does this tell us about the landscape?



- ▶ Many (near-)global minima with $\hat{\mathcal{R}}_n(\theta) \approx 0$.
- ▶ Close to a random initialization
- ▶ At global minima $0 \approx \hat{\mathcal{R}}_n(\theta) \ll \mathcal{R}(\theta)$.

The actual landscape: A cartoon

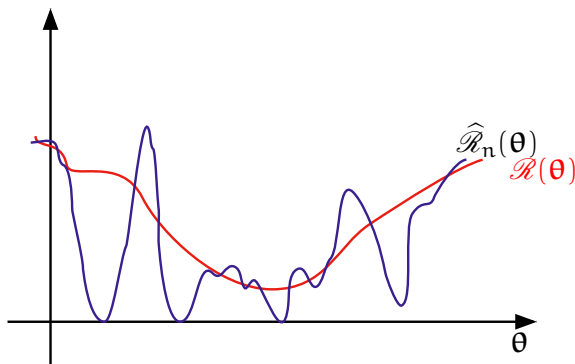


- ▶ How does generalization work?
- ▶ Why does minimizing $\hat{\mathcal{R}}_n(\theta)$ yields models with small $\mathcal{R}(\theta)$?

Many ‘bad’ global optima

Explanation must involve the optimization dynamics

The actual landscape: A cartoon

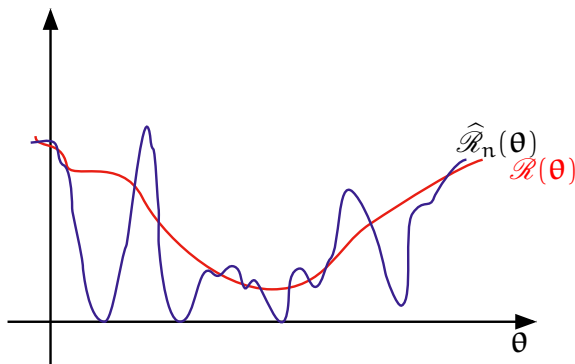


- ▶ How does generalization work?
- ▶ Why does minimizing $\hat{\mathcal{R}}_n(\theta)$ yields models with small $\mathcal{R}(\theta)$?

Many ‘bad’ global optima

Explanation must involve the optimization dynamics

The actual landscape: A cartoon

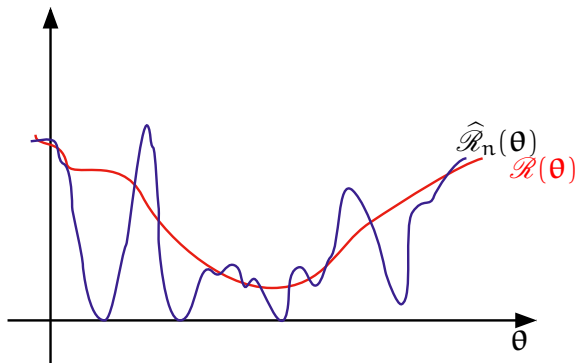


- ▶ How does generalization work?
- ▶ Why does minimizing $\hat{\mathcal{R}}_n(\theta)$ yields models with small $\mathcal{R}(\theta)$?

Many ‘bad’ global optima

Explanation must involve the optimization dynamics

The actual landscape



Two theories

Neural tangent theory:

Global minima near a random initialization behave well

Feature learning:

Few steps of gradient descent sufficient to converge close to $\operatorname{argmin}_{\theta} \mathcal{R}(\theta)$. (Overfitting cannot take place so quickly.)

Two-layer neural networks

$$f(\mathbf{x}; \boldsymbol{\theta}) = \sum_{j=1}^N a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle), \quad \boldsymbol{\theta} = ((a_i)_{i \leq N}, (\mathbf{w}_i)_{i \leq N}) \in \mathbb{R}^N \times (\mathbb{S}^{d-1})^N.$$

- ▶ $p = Nd$
- ▶ Square loss train error:

$$\hat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2.$$

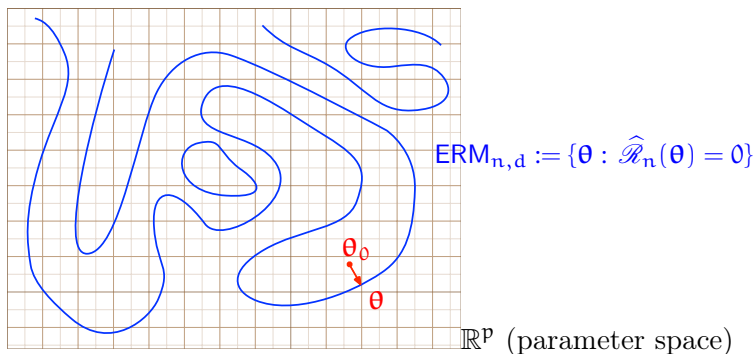
- ▶ Square loss test error:

$$\mathcal{R}(\boldsymbol{\theta}) := \frac{1}{2n} \mathbb{E}\{(y - f(\mathbf{x}; \boldsymbol{\theta}))^2\}.$$

Data $\{(\mathbf{x}_i, y_i) : i \leq n\}$ iid, $\varepsilon_i \sim N(0, \tau^2)$

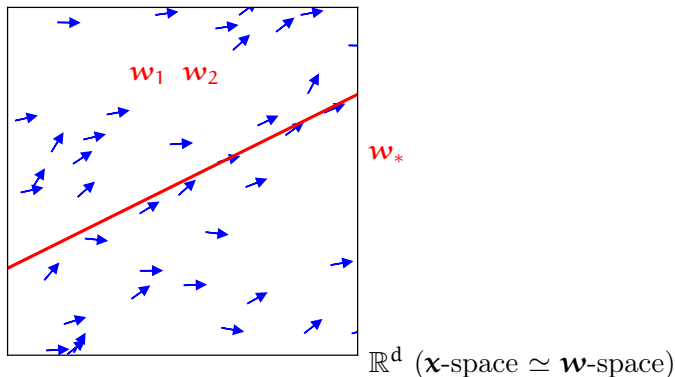
$$\mathbf{x}_i \sim N(0, \mathbf{I}_d), \quad y_i = \varphi(\langle \mathbf{w}_*, \mathbf{x}_i \rangle) + \varepsilon_i$$

‘Neural tangent’ theory



- ▶ Fast convergence to $\hat{\mathcal{R}}_n(\hat{\boldsymbol{\theta}}) \approx 0$
- ▶ Overfitting $\hat{\mathcal{R}}_n(\hat{\boldsymbol{\theta}}) \ll \mathcal{R}(\hat{\boldsymbol{\theta}})$
- ▶ Fits min-norm kernel regression (see next)
- ▶ Statistically suboptimal.

$$f(\mathbf{x}; \theta) = \sum_{j=1}^N a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle)$$



- Converges (ideally) to Bayes error $\hat{\mathcal{R}}_n(\hat{\theta}) \approx \tau^2/2$.
- No overfitting $\hat{\mathcal{R}}_n(\hat{\theta}) \approx \mathcal{R}(\hat{\theta})$
- Learns latent direction $\mathbf{w}_*$: nonlinear.
- Expected to be statistically optimal.

Two theories

Neural tangent theory:

Global minima near a random initialization behave well

Feature learning:

Few steps of GD sufficient to converge to $\operatorname{argmin}_{\boldsymbol{\theta}} \mathcal{R}(\boldsymbol{\theta})$.

Punchline: Each is correct in a different dynamical regime

Two theories

Neural tangent theory:

Global minima near a random initialization behave well

Feature learning:

Few steps of GD sufficient to converge to $\operatorname{argmin}_{\boldsymbol{\theta}} \mathcal{R}(\boldsymbol{\theta})$.

Punchline: Each is correct in a different dynamical regime

Two previews

Preview #1: Generalization in the neural tangent regime

Neural tangent model for 2-layer networks

Linearize around initialization $\mathbf{W}^0$

$$f(\mathbf{x}; \mathbf{a}, \mathbf{W}) = \sum_{j=1}^N a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle)$$

$$f_{\text{NT}}(\mathbf{x}; \bar{\mathbf{b}}, \mathbf{b}) = \sum_{j=1}^N \langle \mathbf{b}_j, \mathbf{x} \rangle \sigma'(\langle \mathbf{w}_j^0, \mathbf{x} \rangle) + \sum_{j=1}^N \bar{b}_j \sigma(\langle \mathbf{w}_j^0, \mathbf{x} \rangle).$$

Gradient descent $\longrightarrow$ closest interpolant

$$\hat{\mathbf{b}} = \operatorname{argmin} \left\{ \|\mathbf{b}\| : y_i = f_{\text{NT}}(\mathbf{x}_i; \mathbf{b}) \quad \forall i \leq n \right\}$$

Neural tangent model for 2-layer networks

Linearize around initialization $\mathbf{W}^0$

$$f(\mathbf{x}; \mathbf{a}, \mathbf{W}) = \sum_{j=1}^N a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle)$$

$$f_{\text{NT}}(\mathbf{x}; \bar{\mathbf{b}}, \mathbf{b}) = \sum_{j=1}^N \langle \mathbf{b}_j, \mathbf{x} \rangle \sigma'(\langle \mathbf{w}_j^0, \mathbf{x} \rangle) + \sum_{j=1}^N \bar{b}_j \sigma(\langle \mathbf{w}_j^0, \mathbf{x} \rangle).$$

Gradient descent $\longrightarrow$ closest interpolant

$$\hat{\mathbf{b}} = \operatorname{argmin} \left\{ \|\mathbf{b}\| : y_i = f_{\text{NT}}(\mathbf{x}_i; \mathbf{b}) \quad \forall i \leq n \right\}$$

Neural tangent model for 2-layer networks

Gradient descent $\longrightarrow$ closest interpolant

$$\hat{\mathbf{b}} = \operatorname{argmin} \left\{ \|\mathbf{b}\| : y_i = f_{\text{NT}}(\mathbf{x}_i; \mathbf{b}) \quad \forall i \leq n \right\}$$

Equivalently

$$\mathbf{z}_i = \boldsymbol{\Phi}(\mathbf{x}_i) := (\mathbf{x}_i^\top \sigma'(\langle \mathbf{w}_1^0, \mathbf{x}_i \rangle), \mathbf{x}_i^\top \sigma'(\langle \mathbf{w}_2^0, \mathbf{x}_i \rangle), \dots, \mathbf{x}_i^\top \sigma'(\langle \mathbf{w}_N^0, \mathbf{x}_i \rangle))^\top,$$

$$\hat{\mathbf{b}}_\lambda = \operatorname{argmin}_{\mathbf{b}} \left\{ \|\mathbf{y} - \mathbf{Z}\mathbf{b}\|_2^2 + \lambda \|\mathbf{b}\|_2^2 \right\}, \quad \lambda \downarrow 0.$$

Neural tangent model for 2-layer networks

Gradient descent $\longrightarrow$ closest interpolant

$$\hat{\mathbf{b}} = \operatorname{argmin} \left\{ \|\mathbf{b}\| : y_i = f_{\text{NT}}(\mathbf{x}_i; \mathbf{b}) \quad \forall i \leq n \right\}$$

Equivalently

$$\mathbf{z}_i = \boldsymbol{\Phi}(\mathbf{x}_i) := (\mathbf{x}_i^T \sigma'(\langle \mathbf{w}_1^0, \mathbf{x}_i \rangle), \mathbf{x}_i^T \sigma'(\langle \mathbf{w}_2^0, \mathbf{x}_i \rangle), \dots, \mathbf{x}_i^T \sigma'(\langle \mathbf{w}_N^0, \mathbf{x}_i \rangle))^T,$$

$$\hat{\mathbf{b}}_\lambda = \operatorname{argmin}_{\mathbf{b}} \left\{ \|\mathbf{y} - \mathbf{Z}\mathbf{b}\|_2^2 + \lambda \|\mathbf{b}\|_2^2 \right\}, \quad \lambda \downarrow 0.$$

Simpler setting: Random features model

$$\min_{\mathbf{b}} \left\{ \sum_{i=1}^n (y_i - f_{\text{RF}}(\mathbf{x}_i; \mathbf{b}))^2 + \lambda \|\mathbf{b}\|_2^2 \right\}, \quad f_{\text{RF}}(\mathbf{x}; \mathbf{b}) = \sum_{i=1}^N b_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle).$$

Theorem (Mei, M. 2019)

Assume $n, N, d \rightarrow \infty$

$$\frac{N}{d} \rightarrow \psi_1, \quad \frac{n}{d} \rightarrow \psi_2.$$

Decompose $\sigma(x) = \sigma_0 + \sigma_1 x + \sigma^{\text{NL}}(x)$ where (for $G \sim N(0, 1)$)

$\mathbb{E}[G \sigma^{\text{NL}}(G)] = \mathbb{E}[\sigma^{\text{NL}}(G)] = 0$, $\zeta^2 := \frac{\sigma_1^2}{\mathbb{E}[\sigma^{\text{NL}}(G)^2]}$. Then there are explicit functions $\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, $\mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, such that

$$R(\hat{f}_\lambda) = F_1^2 \mathcal{B}(\zeta, \psi_1, \psi_2, \lambda) + (\tau^2 + F_*^2) \mathcal{V}(\zeta, \psi_1, \psi_2, \lambda) + F_*^2 + o_d(1).$$

Simpler setting: Random features model

$$\min_{\mathbf{b}} \left\{ \sum_{i=1}^n (y_i - f_{\text{RF}}(\mathbf{x}_i; \mathbf{b}))^2 + \lambda \|\mathbf{b}\|_2^2 \right\}, \quad f_{\text{RF}}(\mathbf{x}; \mathbf{b}) = \sum_{i=1}^N b_i \sigma(\langle \mathbf{w}_i, \mathbf{x} \rangle).$$

Theorem (Mei, M. 2019)

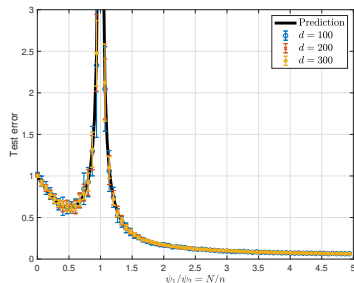
Assume $n, N, d \rightarrow \infty$

$$\frac{N}{d} \rightarrow \psi_1, \quad \frac{n}{d} \rightarrow \psi_2.$$

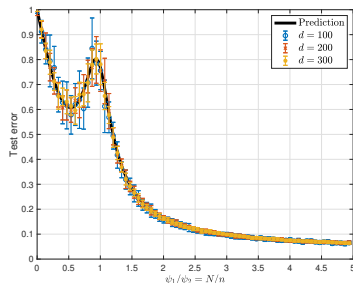
Decompose $\sigma(\mathbf{x}) = \sigma_0 + \sigma_1 \mathbf{x} + \sigma^{\text{NL}}(\mathbf{x})$ where (for $G \sim \mathcal{N}(0, 1)$)
 $\mathbb{E}[G \sigma^{\text{NL}}(G)] = \mathbb{E}[\sigma^{\text{NL}}(G)] = 0$, $\zeta^2 := \frac{\sigma_1^2}{\mathbb{E}[\sigma^{\text{NL}}(G)^2]}$. Then there are explicit
functions $\mathcal{B}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, $\mathcal{V}(\zeta, \psi_1, \psi_2, \bar{\lambda})$, such that

$$R(\hat{f}_\lambda) = F_1^2 \mathcal{B}(\zeta, \psi_1, \psi_2, \lambda) + (\tau^2 + F_*^2) \mathcal{V}(\zeta, \psi_1, \psi_2, \lambda) + F_*^2 + o_d(1).$$

Insight: $N/n \gg 1$ optimal



$\lambda = 0+$



$\lambda > 0$

- ▶ Solid line: Theoretical prediction
- ▶ Mathematics tool: Random matrix theory

Preview #2: Dynamical decoupling of feature learning and overfitting

Gradient flow dynamics

Gradient flow

$$\dot{\boldsymbol{\theta}}(t) = -\frac{n}{d} \nabla \widehat{\mathcal{R}}_n(\boldsymbol{\theta}(t)).$$

Square loss train error

$$\widehat{\mathcal{R}}_n(\boldsymbol{\theta}) := \frac{1}{2n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i; \boldsymbol{\theta}))^2, \quad \boldsymbol{\theta} = (\mathbf{a}, \mathbf{W}).$$

Two-layer network

$$f(\mathbf{x}; \boldsymbol{\theta}) = \frac{1}{N} \sum_{j=1}^N a_j \sigma(\langle \mathbf{w}_j, \mathbf{x} \rangle).$$

Dynamical approach

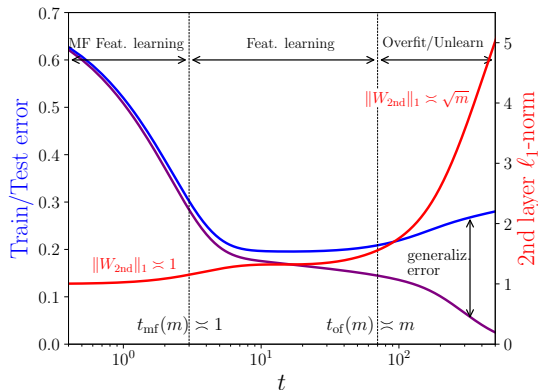
Gradient flow

$$\dot{\boldsymbol{\theta}}(\mathbf{t}) = -\frac{n}{d} \nabla \hat{\mathcal{R}}_n(\boldsymbol{\theta}(\mathbf{t})) .$$

1. Dynamics in a random landscape
2. Exact asymptotics as $n, d \rightarrow \infty$, $n/d \rightarrow \bar{\alpha}$ by rigorizing tools from theoretical physics
[DMFT, dynamical mean-field theory: Celentano, Cheng, M, 2022]
3. Large width limit $N, \bar{\alpha} \rightarrow \infty$, $\bar{\alpha}/N = \alpha$:

[M, Urbani, 2025]

Dynamical decoupling: Learning $\rightarrow$ Overfitting



$t \asymp 1$:

- ▶ Test error $\approx$ Train error
- ▶ Learns latent direction $\mathbf{w}_*$
- ▶ Mean field asymptotically correct (Mei, M, Nguyen, 2018; Chizat, Bach, 2018; Rotskoff, Vanden Eijnden, 2018)

$t \asymp m \equiv N$:

- ▶ Test error $\gg$ Train error; Train error $\rightarrow 0$
- ▶ *Unlearns* latent direction $\mathbf{w}_*$
- ▶ Neural tangent kernel learning

Conclusion

Conclusion

- ▶ Important to ask fundamental questions about AI.
- ▶ What does it mean that ‘AI models learn’?
- ▶ Requires to understand training dynamics

Thank you!

Conclusion

- ▶ Important to ask fundamental questions about AI.
- ▶ What does it mean that ‘AI models learn’?
- ▶ Requires to understand training dynamics

Thank you!