

Deep learning theory - Lecture 10.

Recap: Single layer NNs

$$f(\vartheta) = \frac{1}{\sqrt{m}} \sum_{j=1}^m a_j \sigma(w_j x + b_j) \quad \vartheta = \{(a_j, w_j, b_j)\}$$

behave like linear models: tangent to $\bar{F}_\vartheta$ at ϑ_0 :

$$\bar{f}(\bar{\vartheta}) = f(\vartheta_0) + D_{\vartheta_0} f(\vartheta_0)(\bar{\vartheta} - \vartheta_0)$$

Thm: As $m \rightarrow \infty$, $\text{Lip } D_{\vartheta} f \sim O(\frac{1}{\sqrt{m}})$ so that

$$- \hat{\mathcal{R}}(\vartheta_t) = e^{-\lambda_* t} \hat{\mathcal{R}}(\vartheta_0)$$

$$- \|f(\vartheta_t) - \bar{f}(\bar{\vartheta}_t)\|_{L^2(\mathcal{P})}^2 = O(\frac{1}{\sqrt{m}}) \text{ unif. in } t.$$

Program: - study ridge^(linear) regression models in overpar. regime
- study kernel methods ($p \rightarrow \infty$) in overpar. regime.

Reminder: Ridge regression: $\mathcal{F}_0 = \{z \mapsto z \cdot \beta, \beta \in \mathbb{R}^d\}$

$$\hat{\mathcal{R}}_n(\beta) = \frac{1}{n} \sum_{j=1}^n (z_j \cdot \beta - y_j)^2 + \lambda \|\beta\|_2^2 = \frac{1}{n} \|z \cdot \beta - y\|_2^2 + \lambda \|\beta\|_2^2$$

$$Z = \begin{pmatrix} \text{---} z_1 \text{---} \\ \vdots \\ \text{---} z_n \text{---} \end{pmatrix} \in \mathbb{R}^{n \times d}$$

λ -convex $\Rightarrow$ unique minimizer $\hat{\beta}_\lambda = \frac{1}{n} \left(\frac{1}{n} Z^T Z + \lambda I \right)^{-1} Z^T y$

Depending on the rank of $Z^T Z$ we have $\left\{ \begin{aligned} &= \frac{1}{n} Z^T \left(\frac{1}{n} Z Z^T + \lambda I \right)^{-1} y \\ &\dots \end{aligned} \right.$

a) $Z^T Z$ full rank $\Rightarrow \hat{\beta}_{0+} = (Z^T Z)^{-1} Z^T y$

b) rank $Z^T Z \neq d \Rightarrow \hat{\beta}_{0+} = \arg \min \{ \|\beta\|_2^2 : Z \cdot \beta = y \}$.

Consider a simple model for distribution of z_d, y_d :

$$y_d = \beta^* \cdot z_d + \varepsilon_d$$

$$\mathbb{E}(\varepsilon_d | w_d) = 0$$

$$\mathbb{E}(z_d) = 0$$

$$\mathbb{E}(\varepsilon_d^2 | w_d) = \tau^2$$

$$\mathbb{E}(z_d z_d^T) = \Sigma$$

Assume
to simplify
 $z_d \stackrel{iid}{\sim} N(0, \Sigma)$

In general (bias-variance decomposition)

$$\mathcal{R}(z, \lambda) = \mathcal{B}(z, \lambda) + \mathcal{V}(z, \lambda)$$

$$\mathbb{E}_\varepsilon \mathcal{R}_{\text{exc}}(z, \lambda) = \mathbb{E}_\varepsilon \mathbb{E}_{z_{\text{new}}} \left(\left(\hat{\beta}(z, \varepsilon) \cdot z_{\text{new}} - \beta^* z_{\text{new}} \right)^2 \right)$$

$$= \mathbb{E}_\varepsilon \mathbb{E}_{z_{\text{new}}} \left(\left(\hat{\beta}(z, \varepsilon) - \beta^* \right)^T z_{\text{new}} z_{\text{new}}^T \left(\hat{\beta}(z, \varepsilon) - \beta^* \right) \right)$$

$$= \mathbb{E}_\varepsilon \left\| \hat{\beta}(z, \varepsilon) - \beta^* \right\|_{\Sigma}^2 = \underbrace{\left\| \hat{\beta}(z, \varepsilon) - \mathbb{E}_\varepsilon(\hat{\beta}(z, \varepsilon)) \right\|_{\Sigma}^2}_{\mathcal{B}(z, \lambda)}$$

$$+ \underbrace{\mathbb{E}_\varepsilon \|\mathbb{E}_\varepsilon(\hat{\beta}(z, \varepsilon)) - \beta^*\|_\Sigma^2}_{V(z, \lambda)}$$

For our model we can compute

$$\hat{\beta}(z, \lambda) = \frac{1}{n} S_\lambda z^\top y \quad \text{for } S_\lambda = \left(\frac{1}{n} z^\top z + \lambda \mathbb{I}_d \right)^{-1}$$

so that

$$B(\lambda, z) = \lambda^2 \langle \beta_*, S_\lambda \Sigma S_\lambda \beta_* \rangle$$

$$V(\lambda, z) = \frac{1}{n} \text{Tr} \left(\frac{1}{n} S_\lambda^2 \hat{\Sigma} \Sigma \right). \quad \hat{\Sigma} = \frac{1}{n} z^\top z$$

These quantities are Random in z , but they concentrate around deterministic values

Write formulas for $\lambda = 0_+$ for simplicity:

Thm: Assume that $C^{-1} \leq \frac{n}{p} \leq C$, $\lambda_{\max}(\Sigma) \leq C$

whp.

$$|B_n(z, \lambda) - \bar{B}_n(\lambda_*)| \leq \bar{C} n^{-\frac{1}{6}}$$

$$|V_n(z, \lambda) - \bar{V}_n(\lambda_*)| \leq \bar{C} n^{-\frac{1}{6}}$$

$$\frac{1}{p} \sum \lambda_j^{-1}(\Sigma) \leq C$$

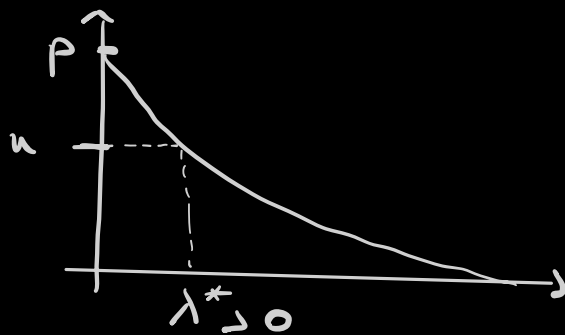
$$\|\Sigma\| = 1$$

where

$$\mathcal{B}_n(\lambda_*) = \frac{\lambda_*^2 \beta_*^T (\Sigma + \lambda_* \mathbb{1})^{-2} \Sigma \beta_*}{1 - \frac{1}{n} \text{Tr}(\Sigma^2 (\Sigma + \lambda_* \mathbb{1})^{-2})}$$

$$\mathcal{V}_n(\lambda_*) = \frac{\frac{1}{n} \text{Tr}(\Sigma^2 (\Sigma + \lambda_* \mathbb{1})^{-1})}{1 - \frac{1}{n} \text{Tr}(\Sigma^2 (\Sigma + \lambda_* \mathbb{1})^{-2})}$$

and λ^* solves $n = \text{Tr}(\Sigma (\Sigma + \lambda_* \mathbb{1})^{-1})$. (*)



Note (implicit regularization) The model

deterministic $y_s = \Sigma^{1/2} \beta^* + \frac{\omega}{\sqrt{n}} \varepsilon'_s$ $\varepsilon'_s \stackrel{iid}{\sim} \mathcal{N}(0, \mathbb{1})$
 fitted through

$$\hat{\beta}(\lambda_*) = \arg \min_{\beta} \underbrace{\frac{1}{n} \|\Sigma^{1/2} \beta - y\|^2 + \lambda_* \|\beta\|_2^2}_{\hat{\mathcal{R}}(\lambda_*, \varepsilon')}$$

has the risk decomposition $\hat{\mathcal{R}}(\lambda_*, \varepsilon')$

$$\mathcal{R}_n(\lambda_*) = \mathbb{E}_{\varepsilon'} \hat{\mathcal{R}}(\lambda_*, \varepsilon') = \mathcal{B}_n(\lambda_*) + \mathcal{V}_n(\lambda_*)$$

$\Rightarrow$ asymptotically, it is like fitting a simple model with positive regularization! deterministic in $\mathcal{Z}$.

Bounding terms B_n, V_n .

Start with denominator:

$$\frac{1}{n} \text{Tr}(\Sigma^2 (\Sigma + \lambda_* \mathbb{I})^{-2}) \leq \frac{1}{n} \text{Tr}(\Sigma (\Sigma + \lambda_* \mathbb{I})^{-1}) = 1 \quad \swarrow \text{cond}^*$$

assume that $\frac{1}{n} T_n(\Sigma^2 (\Sigma + \lambda_* \mathbb{1})^2) < 1 - \frac{1}{C_*} (**)$

$$\Rightarrow \begin{cases} \mathcal{B}_n(\lambda_*) \leq \lambda_*^2 \beta_*' (\Sigma + \lambda_* \mathbb{1})^2 \Sigma \beta_* \\ \mathcal{V}_n(\lambda_*) \leq \frac{T^2}{n} T_n(\Sigma + \lambda_* \mathbb{1})^{-1} \Sigma^2 \end{cases}$$

Consider λ_* small, $\{\sigma_k\} = \sigma(\Sigma)$ $\Sigma = \sum \sigma_d \sigma_d \sigma_d^T$
 $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_k \geq \dots$

Define k_* such that $\sigma_k \geq \lambda_* \geq \sigma_{k+1}$

Prop: If $(**)$ holds

$$\mathcal{V}_n \lesssim C_* T^2 \frac{k_*}{n}$$

$$\mathcal{B}_n(\lambda_*) \leq C_* \left(\sigma_{k_*}^2 \|\beta_{*, \leq k_*}\|_{\Sigma^{-1}}^2 + \|\beta_{*, \geq k_*}\|_{\Sigma}^2 \right)$$

where $\beta_{*, \leq k_*} = \sum_{d=1}^{k_*} (\beta_* \cdot \sigma_d) \cdot \sigma_d$

$\Rightarrow$ even if $p \gg n$ we can have.

$$\mathcal{V}_n \xrightarrow{n \rightarrow \infty} 0 \quad \left(\frac{k_*}{n} \rightarrow 0 \right)$$

$$\mathcal{B}_n \xrightarrow{n \rightarrow \infty} 0 \quad \left(\beta_* \text{ concentrated on top eigenvalues of } \Sigma \right)$$

"Example": $\beta_* = \begin{pmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ 0 \end{pmatrix}$ $\Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$

Benign overfitting

$$\begin{aligned}\Rightarrow n &= \sum \frac{\sigma_e}{\sigma_e + \lambda_*} \geq \sum_{e=1}^{h_*} \frac{\sigma_e}{2\sigma_e} + \sum_{e=h_*+1}^p \frac{\sigma_e}{2\lambda_*} \\ &= h_* + \frac{1}{\sigma_e} \sum_{e=h_*+1}^p \sigma_e \asymp h_* + r_h^{(1)}\end{aligned}$$

Example: $\sigma_e \asymp h^{-\alpha} \xrightarrow{\alpha > 1} r_h^{(1)} = \frac{1}{h^{-\alpha}} \int_h^\infty x^{-\alpha} dx \asymp h$

$$\Rightarrow n \asymp h$$

$$\begin{aligned}\Rightarrow V_n &\asymp c_* \frac{T^2}{n} \text{Tr}(\Sigma^2 (\Sigma + \lambda_* \mathbb{1})^{-2}) \\ &= c_* \frac{T^2}{n} \left(h_* + \frac{1}{\sigma_e} \sum_{e=h_*+1}^p \sigma_e^2 \right) \asymp c_* \frac{T^2}{n} h_*\end{aligned}$$

$$\begin{aligned}\Rightarrow \mathcal{R}_n(\beta) &\leq c_* \sum \frac{\lambda_*^2 \sigma_e}{(\sigma_e + \lambda_*)^2} (\sigma_e^T \beta_*)^2 \quad \xrightarrow{\beta_e} \overline{\beta}_e \\ &= c_* \frac{2}{\sigma_e} \left(\sum_{e=1}^{h_*} \overline{\beta}_e^2 \frac{1}{\sigma_e} + \sum_{e=h_*+1}^p \overline{\beta}_e^2 \frac{\sigma_e^2}{\sigma_e^2} \right)\end{aligned}$$

Example (Latent space model)

$$\begin{cases} y_i = \theta_*^T x_i + \xi_i & \xi_i \sim \mathcal{N}(0, \tau^2) \\ z_i = W x_i + u_i & u_i \sim \mathcal{N}(0, \mathbb{1}). \end{cases}$$

$$x_i \in \mathcal{N}(0, \mathbb{1}_d) \quad W^T W = \frac{p}{d}$$

assume $\frac{p}{n} \rightarrow \gamma \in (0, \infty)$, $\frac{d}{p} \rightarrow \psi \in (0, 1)$

