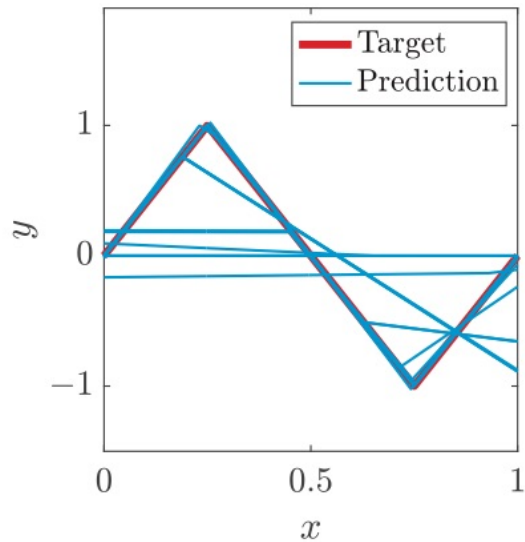
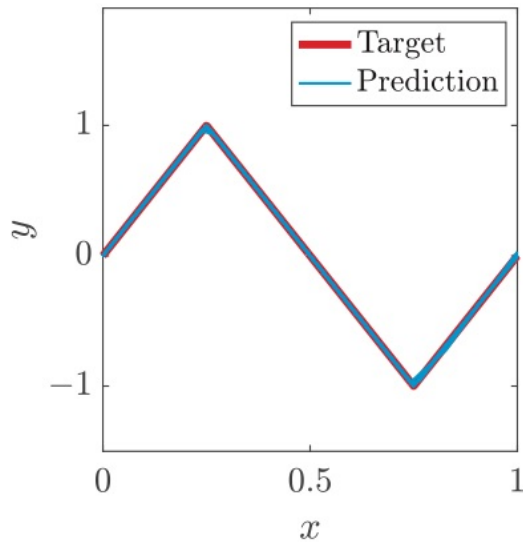


$$f(x) = \sum_{j=1}^m \eta_j (w_j x + b_j) +$$

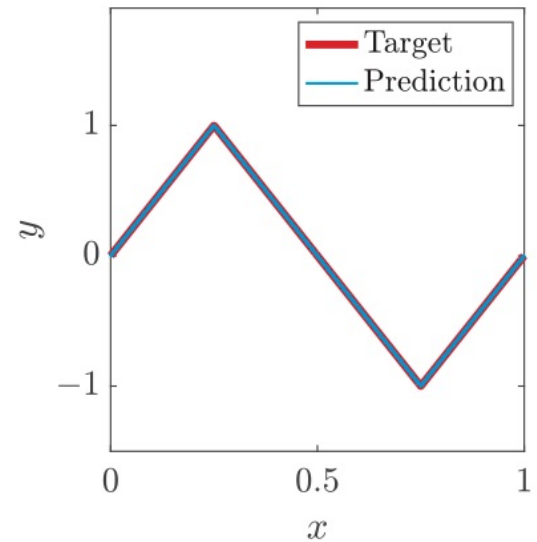
Prediction functions – $m = 5$



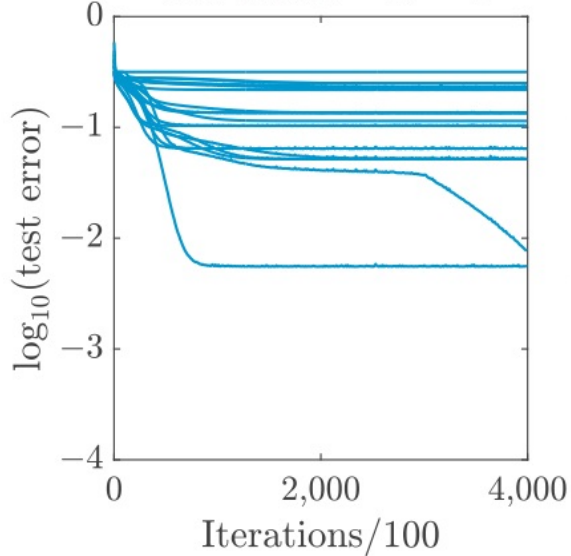
Prediction functions – $m = 20$



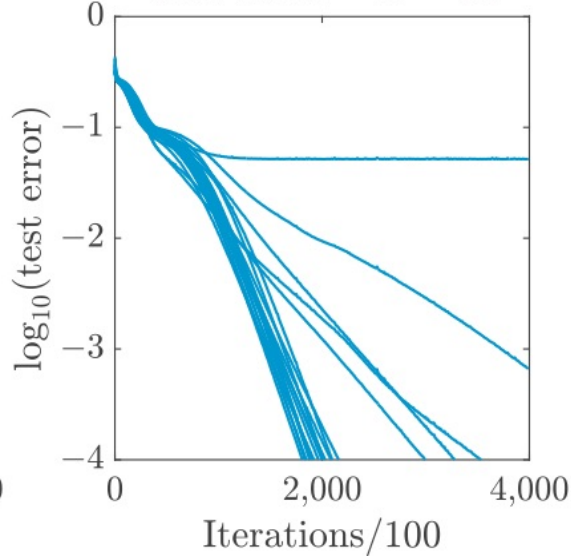
Prediction functions – $m = 100$



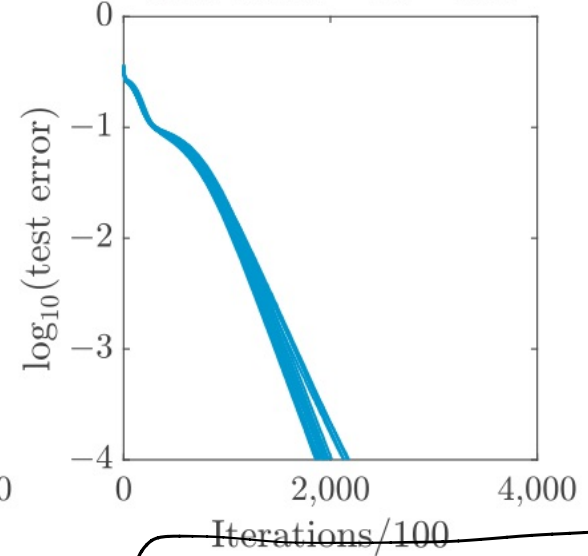
Test errors – $m = 5$



Test errors – $m = 20$



Test errors – $m = 100$



(overparameterized)

→ Chapter 12

Exam: December 11 9^h $\rightarrow$ 11^h30

2 sides of A4 page, of notes
short answers only!

HW solutions online by Dec.

Data (x_i, y_i) $i \in 1, \dots, n$

$$\hat{R}(f) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, f(x_i)) + \Omega(f)$$

log LR square loss $(y_i - f(x_i))^2$

$y \in \{-1, 1\}$ 0-1 loss $\mathbb{1}_{y_i \neq f(x_i)} \Rightarrow$ logistic loss $\log(1 + e^{-y_i f(x_i)})$

Empirical minimization:

Model: $f_\theta: \mathcal{X} \rightarrow \mathbb{R}$

- linear models: $f_\theta(x) = \theta^\top \phi(x)$ $\phi(x) \in \mathbb{R}^d$
Chapter 3, 4, 5, 8
- kernel methods (positive definite)
Chapter 7 $f_\theta(x) = \langle \theta, \phi(x) \rangle$ $\phi(x) \in \mathcal{H}$

kernel $k(x, x') = \langle \phi(x), \phi(x') \rangle$.

↳ linear
polynomial
gaussian

human
recognition

- Today: neural networks

"feature" learning

Neural networks: single hidden layer fully connected $n \in \mathbb{R}^d$

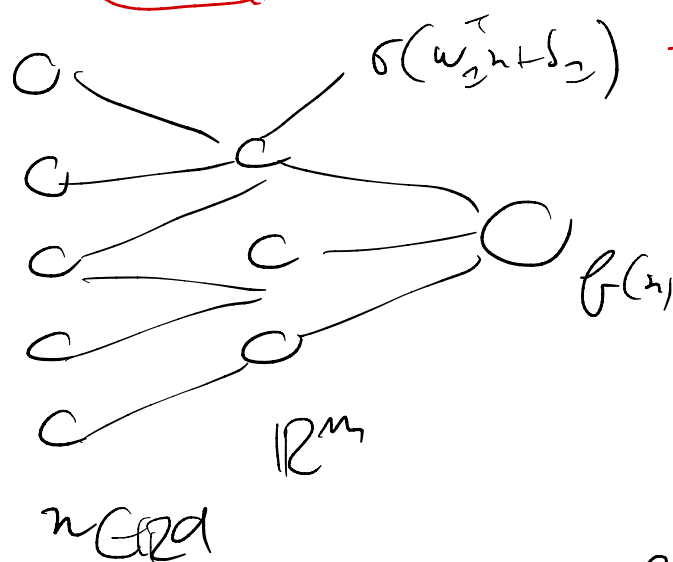
$$f(x) = \sum_{j=1}^m \eta_j \cdot \sigma(w_j^T x + b_j)$$

input weights
 $\in \mathbb{R}^{m \times (d+1)}$

$b_j =$ "biases"
= "intercept"
= constant term

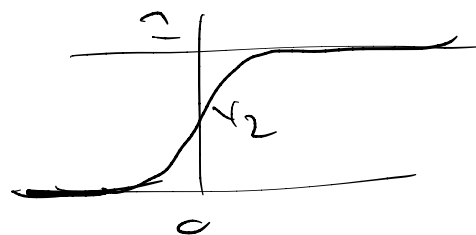
hidden neurons

output weight
 $\in \mathbb{R}^m$



σ : activation function

Sigmoid = $\sigma(x) = \frac{1}{1 + e^{-x}}$



Rectified linear unit = ReLU

$\sigma(x) = \max(0, x)$



! No activation at the last layer
often for classification

logistic loss
on $f(x)$

$f(x) = \sum_j \eta_j \cdot \sigma(w_j^T x + b_j)$ $g(x) = \text{sigmoid}(f(x)) \in \mathbb{R}$

$\Rightarrow$ (cross-entropy loss) $\hat{=}$ maximum likelihood

$p(y=1) = g(x)$

$- \mathbb{1}_{y=1} \log g(x) - \mathbb{1}_{y=-1} \log(1 - g(x))$

Algorithm: minimize $R(\theta)$ with $\ell_0(n) = \sum_{j=1}^n y_j \sigma(w_j^T x + b_j)$

$$\theta = (w, b)$$

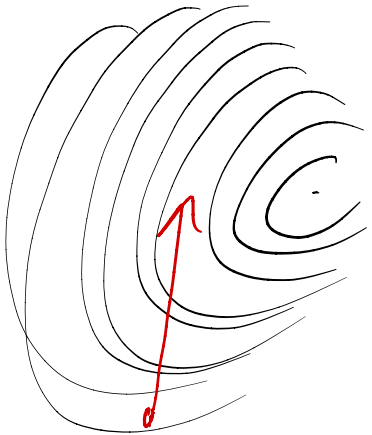
$$\in \mathbb{R}^{M \times (d+2)}$$

using SGD or ADAM or --

two types: single pass $\rightarrow$ directly minimizing the test risk
multiple passes

$\hookrightarrow$ not enough data
minimize training risk
need to control overfitting

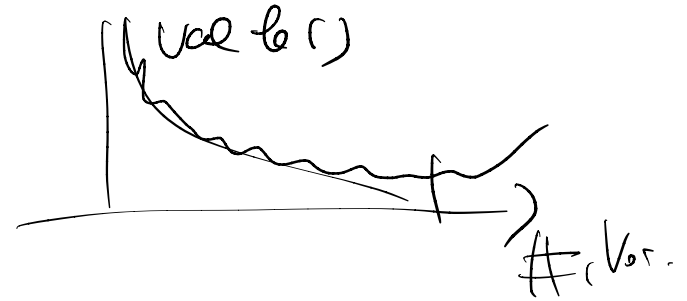
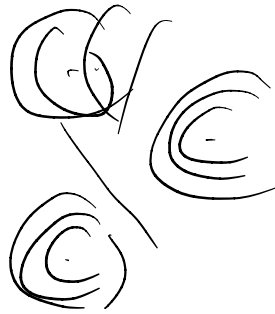
Convex



0 gradient $\Rightarrow$ global optm.

$$R(\bar{\theta}_T) - R(\theta_*) \leq \frac{C}{\sqrt{T}}$$

Non Convex



$$\mathbb{E} \| \nabla R(\bar{\theta}_T) \|^2 \leq \frac{1}{\sqrt{T}}$$

$t \in (c, T)$

Error decomposition: $\underbrace{R(\hat{f}) - R^*}_{\text{excess risk}} = \underbrace{R(\hat{f}) - \inf_{g \in F} R(g)}_{\text{estimation error}} + \underbrace{\inf_{g \in F} R(g) - R^*}_{\text{approximation error}}$

for $\hat{f}$ an approximate ERM

$\leq 2 \cdot \underbrace{\text{uniform deviation}}_{\text{empirical risk}} + \underbrace{\text{optimization error}}_{\text{approximation error}} = 2 \sup_{g \in F} |\bar{R}(g) - R(g)| + \inf_{g \in F} R(g) - R^*$

DR(SN)

consider $\sigma(u) = \text{ReLU}(u) = (u)_+ = \max\{0, u\}$.
homogeneously $\sigma(\lambda u) = \lambda \sigma(u)$ for $\lambda \geq 0$

$$f(w) = \sum_{j=1}^m \alpha_j (w_j^T x + b_j)_+ \quad \text{if } \alpha_j > 0$$

penalty on weights: $\ell_2\text{-norm}^2 = \frac{1}{2} \sum_{j=1}^m \alpha_j^2 + \frac{\|w\|_2^2 + b^2}{\alpha_j^2}$

optimal: $|q_j| = \sqrt{\alpha_j^2 + b_j^2}$

$$\Rightarrow \sum_{j=1}^m |q_j| \sqrt{\alpha_j^2 + b_j^2}$$

$$F = \left\{ \sum_{j=1}^m \alpha_j (w_j^T x + b_j)_+ \mid \|w\|_2^2 + b^2 = 1 \text{ and } \|q\|_2 \leq D \right\}$$

($\Rightarrow$ weight decay)
($\Rightarrow$ "implicit bias")
of (SIG)
see Ch. 12

$$\mathcal{F} = \left\{ \sum_{j=1}^m a_j (w_j^T x + b_j)_+ \text{ s.t. } \|w_j\|_2^2 + b_j^2 = 1 \text{ and } \|w\|_2 \leq D \right\}$$

assume $\|x\|_2 \leq R$.

$$\text{if } \hat{C} = \arg \min_{C \in \mathcal{F}} \hat{R}(C)$$

$$\underbrace{R(\hat{C}) - \inf_{C \in \mathcal{F}} R(C)}_{\text{estimation error}} \leq \sup_{C \in \mathcal{F}} \hat{R}(C) - R(C) + \sup_{C \in \mathcal{F}} R(C) - \bar{R}(C) \leq 2 \sup_{C \in \mathcal{F}} |R(C) - \bar{R}(C)|$$

$$\mathbb{E} \sup_{C \in \mathcal{F}} \hat{R}(C) - R(C) \leq 2 \mathbb{E} \left[\sup_{C \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \varepsilon_i \ell(w_i, C(x_i)) \right]$$

Rademacher complexity average

$$\leq 2G \mathbb{E} \sup_{C \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \varepsilon_i f(x_i)$$

if f is G -Lipschitz

$$\left(\text{typically } \frac{D}{\sqrt{n}} \right)$$

$$\leq 2G \times \frac{4eDL}{\sqrt{n}}$$

$$F = \int \sum_{j=1}^m q_j (w_j^T x + b_j)_+ \text{ s.t. } \|w\|_2^2 + \frac{b^2}{R^2} = 1$$

and $\|q\|_2 \leq \frac{1}{R^2}$

same $\|x\|_2 \in R$

$$\|q\|_1 = \sum_{j=1}^m |q_j|$$

$$\mathbb{E}_{\substack{\varepsilon \\ \text{Data}}} \sup_{f \in F} \frac{1}{n} \sum_{i=1}^n \varepsilon_i f(x_i)$$

$$= \mathbb{E} \sup_{q, w, b} \frac{1}{n} \sum_{i=1}^n \varepsilon_i \sum_{j=1}^m q_j (w_j^T x_i + b_j)_+$$

$$= \mathbb{E} \sup_{w, b} \sup_{\|q\|_2 \leq 1} \sum_{j=1}^m q_j \left[\frac{1}{n} \sum_{i=1}^n \varepsilon_i (w_j^T x_i + b_j)_+ \right]$$

$$\begin{aligned} \sup_{\|q\|_1 \leq 1} q^T z &= \|z\|_\infty \\ q^T z &= \sum_{j=1}^m q_j z_j \\ &\leq \sum_{j=1}^m |q_j| |z_j| \\ &\leq \|q\|_1 \sum_{j=1}^m |z_j| \\ &\leq \|z\|_\infty \sum_{j=1}^m |z_j| \\ &\leq \|z\|_\infty \cdot 1 \end{aligned}$$

$$= D \mathbb{E} \sup_{\substack{\|w\|_2^2 + \frac{b^2}{R^2} = 1 \\ j \in \{1, \dots, m\}}} \sup_{\substack{\varepsilon \in \{-1, 1\}^n \\ \varepsilon \in \{1, \dots, m\}}} \left| \frac{1}{n} \sum_{i=1}^n \varepsilon_i (w_j^T x_i + b_j)_+ \right| \text{ using lemma}$$

$$= D \mathbb{E} \sup_{\|w\|_2^2 + \frac{b^2}{R^2} = 1} \left| \frac{1}{n} \sum_{i=1}^n \varepsilon_i (w^T x_i + b)_+ \right|$$

$$\leq 2D \mathbb{E} \sup_{\|w\|_2^2 + \frac{b^2}{R^2} = 1} \left| \frac{1}{n} \sum_{i=1}^n \varepsilon_i (w^T x_i + b) \right|$$

using the contraction principle with 1-Lipschitz for

$$\left\| \begin{pmatrix} w \\ b/R \end{pmatrix} \right\|_2 = 1 \quad \left| \frac{1}{n} \sum_{i=1}^n \varepsilon_i \begin{pmatrix} w \\ b/R \end{pmatrix}^T \begin{pmatrix} x_i \\ 1 \end{pmatrix} \right|$$

$\in \mathbb{R}^{d+1}$

$$= 2D \mathbb{E} \left\| \frac{1}{n} \sum_{i=1}^n \varepsilon_i \begin{pmatrix} x_i \\ 1 \end{pmatrix} \right\|_2 \leq 4D \frac{R}{\sqrt{n}}$$

by Cauchy-Schwarz

same as chap. 4

3.36

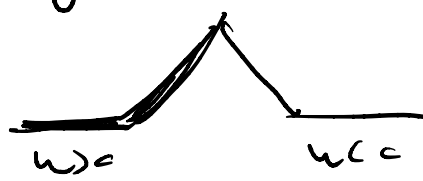
3.56

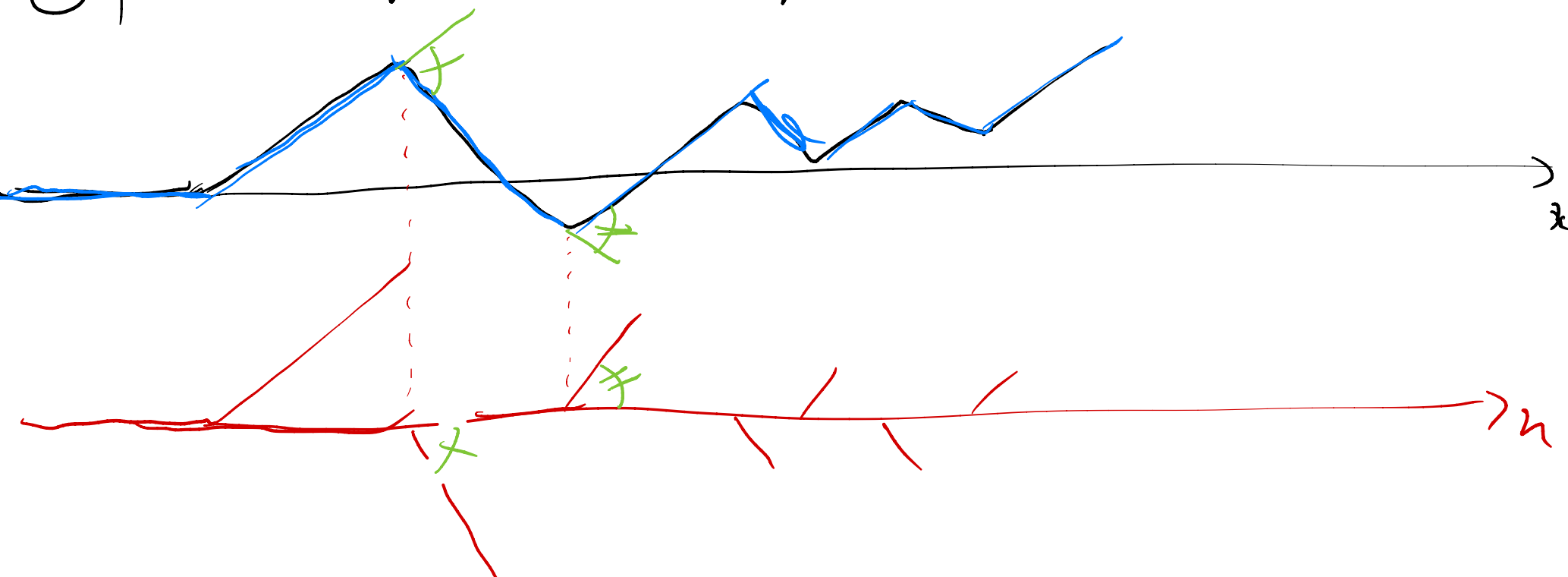
Approximation properties

- ① Universality:
 - ② if $\|g\|_2 \in \mathcal{D}$ which functions can I approximate
 - ③ how big m should be?
- $m = m$ is enough
- $m = \infty$

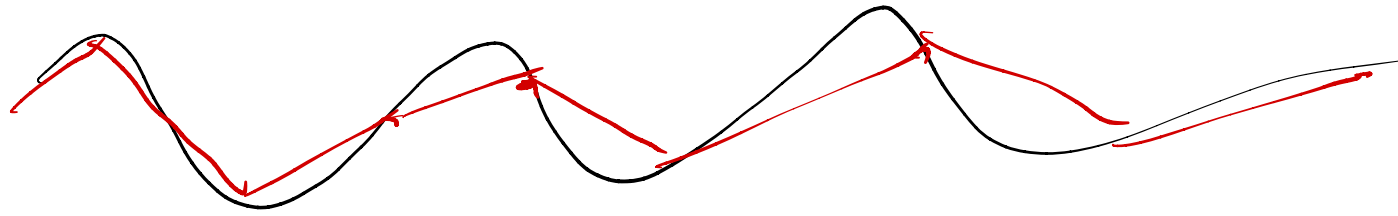
Universality: if $f: \mathbb{R}^d \rightarrow \mathbb{R}$, can I find $\theta = (w, b)$ s.t.
 $f_\theta \approx f$

① in 1D proof by drawing

② piecewise affine continuous functions (units) + 



③ for any "reasonable" function, there is a piecewise affine
calculus for which is a
"close" as



$\Rightarrow$ all "reasonable" fcts can be approx. by a neural ^{desired} net

Formulation through measures

$(w_j, b_j) \in K$ compact set

$$f(x) = \sum_{j=1}^m \eta_j (w_j^T x + b_j)_+ = \int_K (w^T x + b)_+ d\eta(w, b)$$

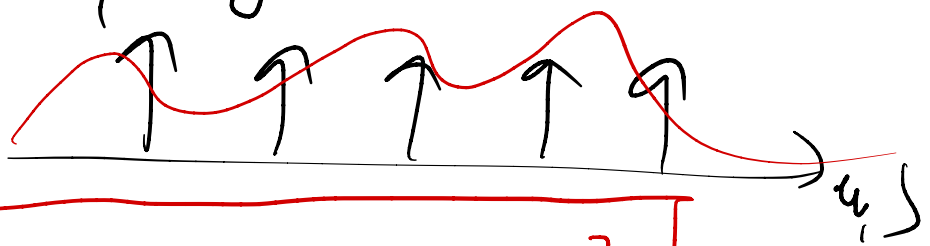
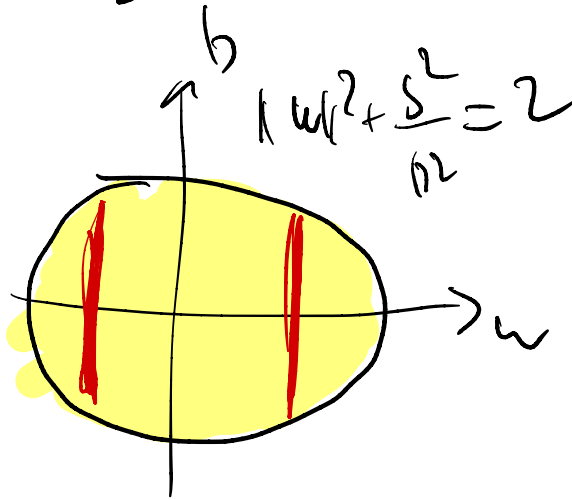
penalty is $\|\eta\| = \sum_j |\eta_j|$

$$= \int |d\eta(w, b)| = \text{total variation}$$

with $d\eta(w, b) = \sum_{j=1}^m \eta_j \delta_{w_j, b_j}$ Dirac $\nearrow$

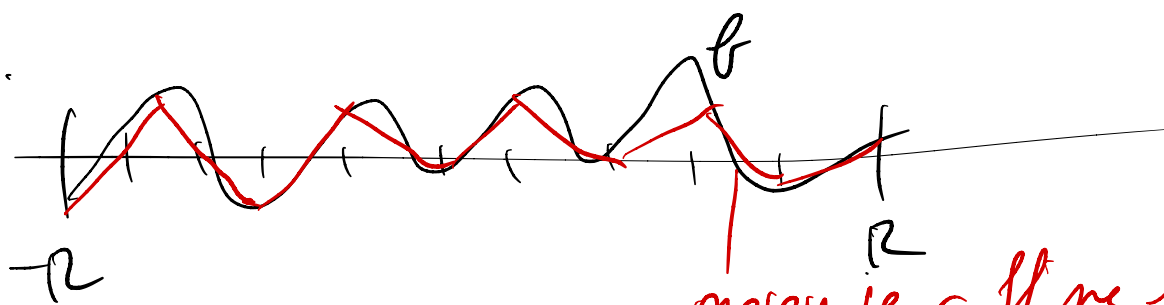
limit when $m \rightarrow +\infty$

$f(x) = \int (w^T x + b)_+ d\eta(w, b)$ with penalty $\int |d\eta(w, b)|$
for any finite measure



$$K = \{ \|v\|_2 = 1, b \in \mathbb{R} \}$$

in 1D: optan 1.

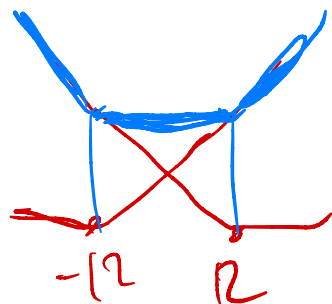


piecewise = f line interpolant

optan 2: Taylor's formula with integral remainder

$$\forall n \in \mathbb{N}; f(x) = f(-R) + \int_{-R}^x f'(s) ds$$

$$= f(-R) + f'(-R)(x+R) + \int_{-R}^x f''(s)(x-s) ds$$



$$\frac{f(-R)(x+R)}{2R - (-R)} = \frac{f(-R)(x+R)}{3R}$$

$$= \int_{-R}^x (x-s)_+ dy(w, s)$$

$$[-1, 1] \times [-R, R]$$

$$\text{with } dy(1, s) = \left(\frac{f(-R)}{2R} + f'(-R) \right) \delta_{-R} + \frac{f''(s) ds}{2R}$$

$$dy(-1, s) = \frac{f(-R)}{2R} \delta_{-R}$$

$$\int_{-R}^R f''(s)(x-s)_+ ds$$

$$\int |dy(w, s)| = \int |f''(s)| ds$$

see book



From 1D to any dimension

$$f: \mathbb{R}^d \rightarrow \mathbb{R}$$

$$f(x) = \frac{1}{(2\pi)^d} \int \hat{f}(u) \underbrace{e^{i u^T x}}_{e^{i u}} du$$

inverse Fourier formula

$$e^{i u} = \int (\alpha u + \beta)_+ d\gamma(u, \beta)$$

$$= \frac{1}{(2\pi)^d} \iint \hat{f}(u) (\alpha u + \beta)_+ du d\gamma(u, \beta) \Rightarrow \underline{\text{universality}}$$

with an explicit norm on $\mathcal{G}$

the norm is "better" than for kernels.

why?

if $f(x) = h(w^T x)$ single index model

ex: $h(x_j)$

the norm will be controlled

derivatives
dependence on unknown
projector

[See Chap 9.3]