

Linear regression

PRML 3, MLPP 7, UML 9

Domain set $X = \mathbb{R}^D$

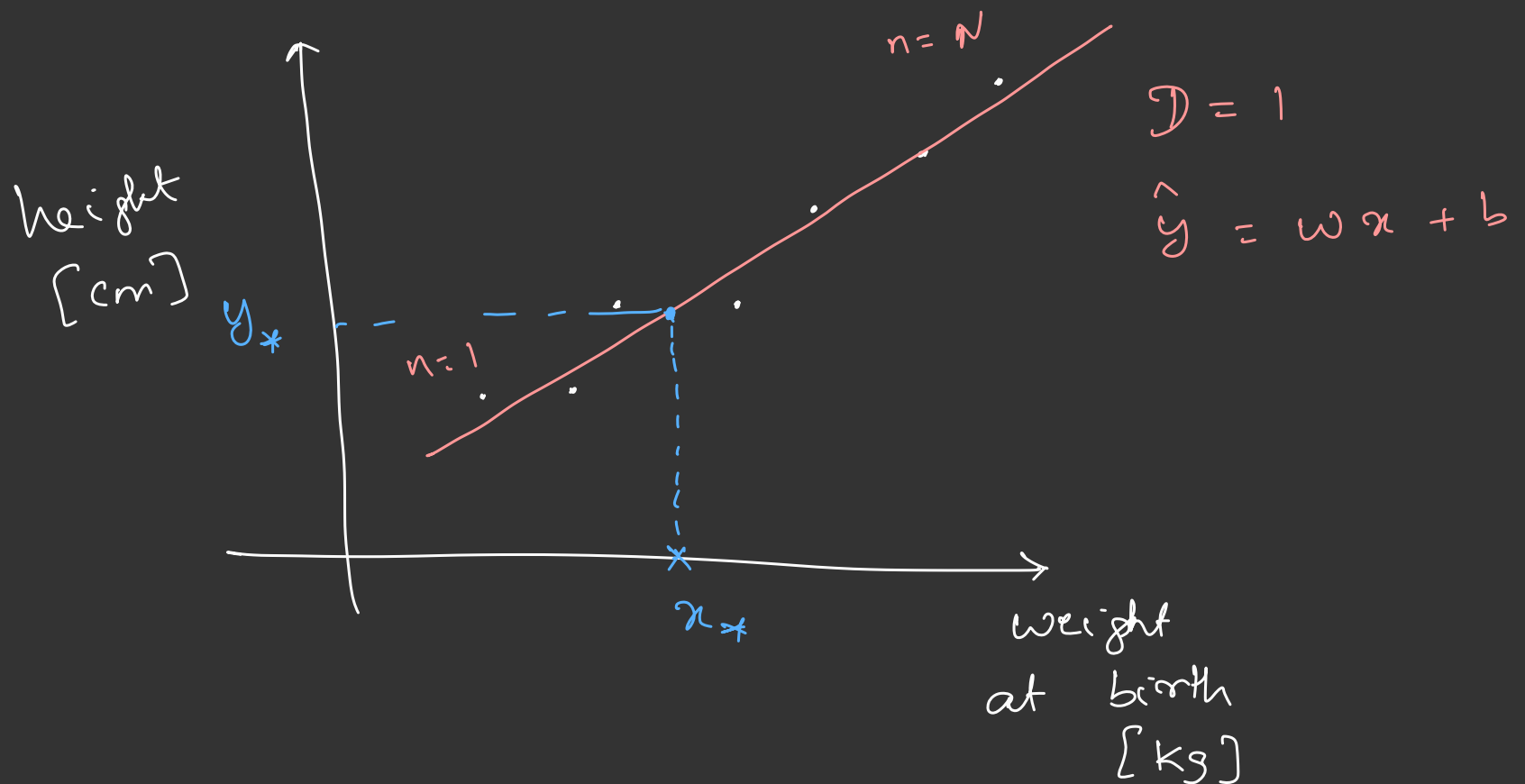
Label set $Y = \mathbb{R}$

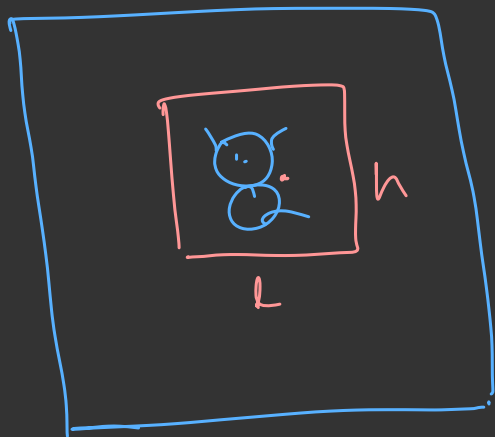
Goal: Learn a linear function $h: \mathbb{R}^D \rightarrow \mathbb{R}$
that best approximates the relationship
between our variables.

$$\mathcal{D} = \left\{ (\underline{x}_i, y_i) \right\}_{i=1}^N$$

Hypothesis class of linear regression predictors
is simply a linear function:

$$H_{\text{reg}} = \mathcal{L}_{\mathcal{D}} = \{ \underline{x} \mapsto \langle \underline{w}, \underline{x} \rangle + b : \underline{w} \in \mathbb{R}^D, b \in \mathbb{R} \}$$





$$\underline{y} = \begin{bmatrix} l \\ h \\ p_x \\ p_y \end{bmatrix} \in \mathbb{R}^4$$

$$\{ \underline{x}_i, \underline{y}_i \}$$

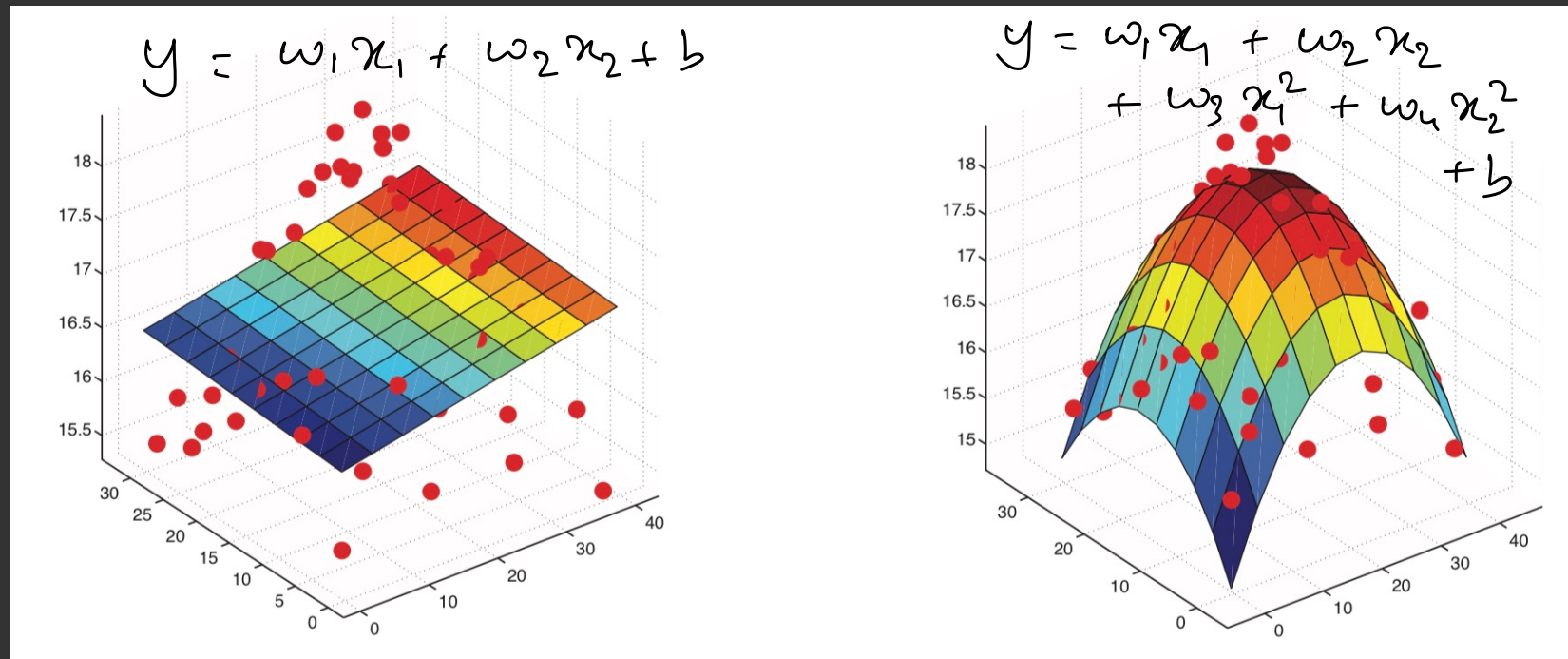
$$\underline{y}_i \approx h(\underline{x}_i)$$

Linear regression:

$$h(\underline{x}) = \underline{\omega}^T \underline{x}$$

$$= \omega_1 x_1 + \omega_2 x_2 + \dots + \omega_D x_D$$

Basis function models:



Linear regression can be used to model non-linear relationships by replacing $\underline{x}$ with $\phi(\underline{x})$

$$h(\underline{x}) = \underline{\omega}^T \phi(\underline{x})$$

$$\phi(\underline{x}) = [1, x, x^2, \dots, x^D]$$

Loss functions :-

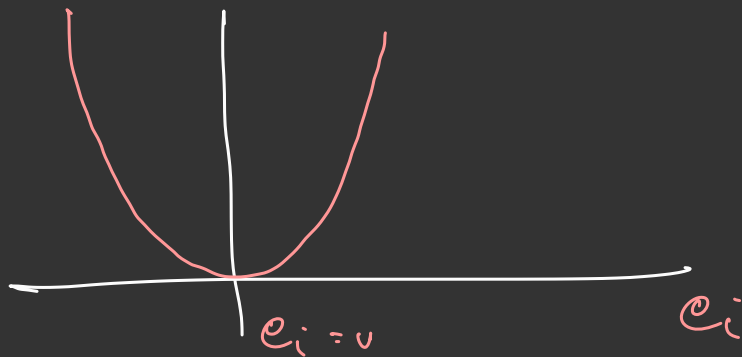
Discrepancy between $h(\underline{x})$ and y

$$l(h(\underline{x}), y) = h(\underline{x}) - y$$

ERM:

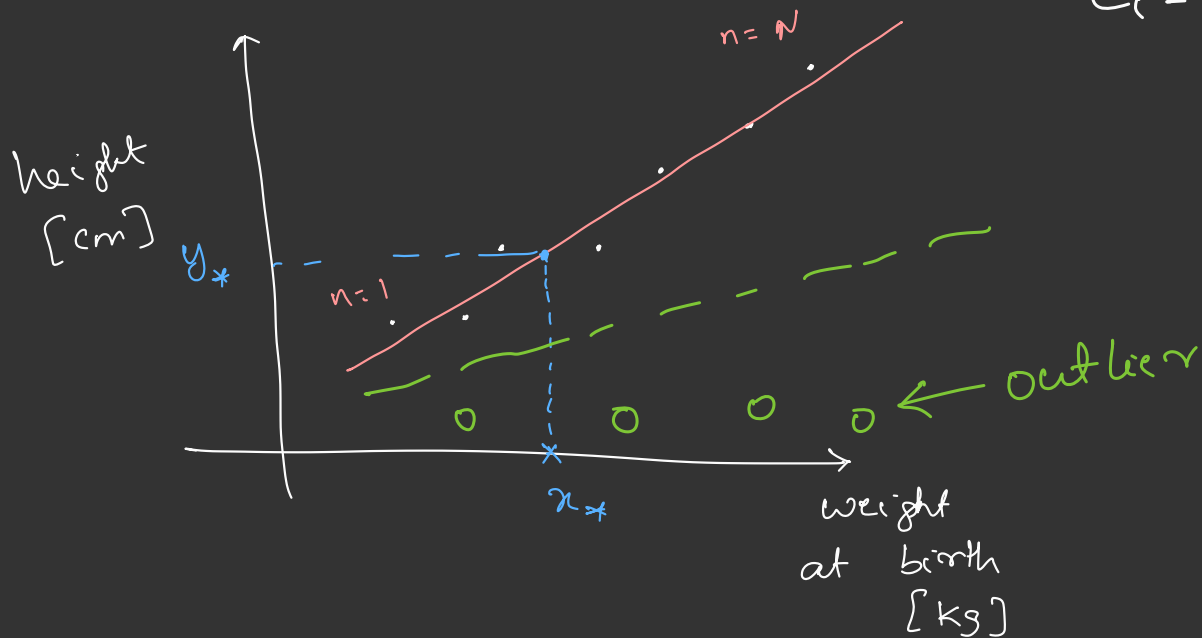
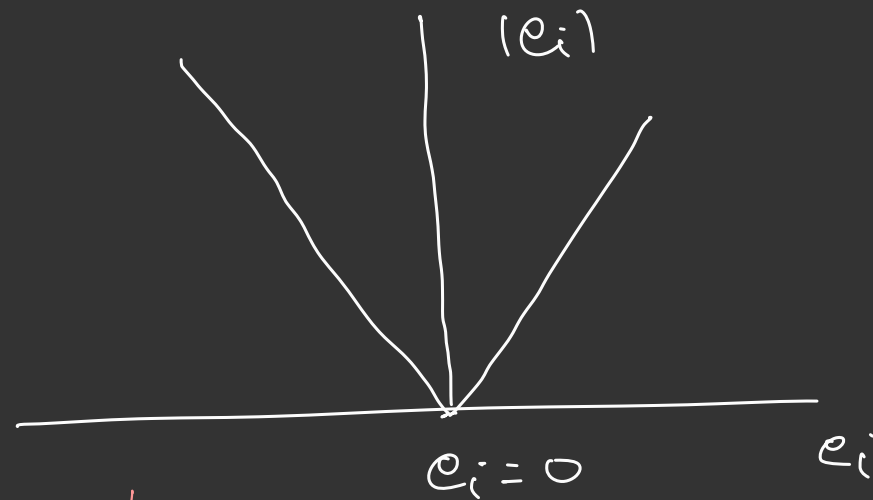
(A) Mean squared error:

$$L_D(\underline{w}) = \frac{1}{N} \sum_{i=1}^N \underbrace{(h(\underline{x}_i) - y_i)^2}_{e_i}$$



(B) Mean absolute error:

$$L_D(\underline{w}) := \frac{1}{N} \sum_{i=1}^N |h(\underline{x}_i) - y_i|$$



Least Squares:

$$\arg \min_{\underline{w}} \quad \mathcal{L}_D(\underline{w}) = \frac{1}{N} \sum_{i=1}^N (\underline{w}^T \underline{x}_i - y_i)^2$$

$$\frac{\partial \mathcal{L}_D}{\partial \underline{w}} = 0$$

$$\frac{2}{N} \sum_{i=1}^N \underline{x}_i (\underline{w}^T \underline{x}_i - y_i) = 0$$

$$\left(\sum_{i=1}^N \underline{x}_i \underline{x}_i^T \right) \underline{w} = \sum_{i=1}^N y_i \underline{x}_i$$

$$\underline{w} = \left(\sum_{i=1}^N \underline{x}_i \underline{x}_i^T \right)^{-1} \sum_{i=1}^N y_i \underline{x}_i$$

$$\underline{w} = (X^T X)^{-1} X^T y$$

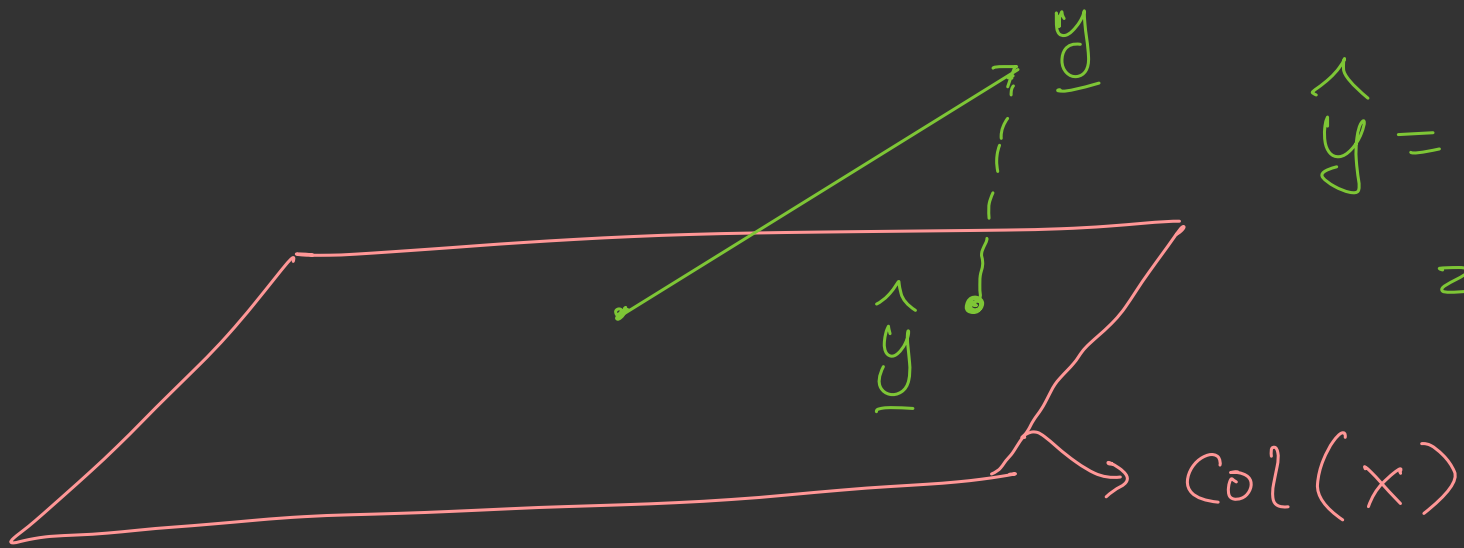
$$X^T = [\underline{x}_1, \underline{x}_2, \dots, \underline{x}_N] : D \times N$$

$$\underline{y} = [y_1, y_2, \dots, y_N]^T : N \times 1$$

$$X^+ = (X^T X)^{-1} X^T$$

$$X^+ X = (X^T X)^{-1} (X^T X) = \underline{I}$$

$$X X^+ = P_X$$



$$\hat{y} = P_X y$$

$$= X(X^T X)^{-1} X^T y$$

$$e = y - \hat{y} = [I - X(X^T X)^{-1} X^T] y$$

$$X^T e = 0$$

Probabilistic discriminative model:-

$$y = h(x) + e$$

$$e \sim \mathcal{N}(0, \sigma^2)$$

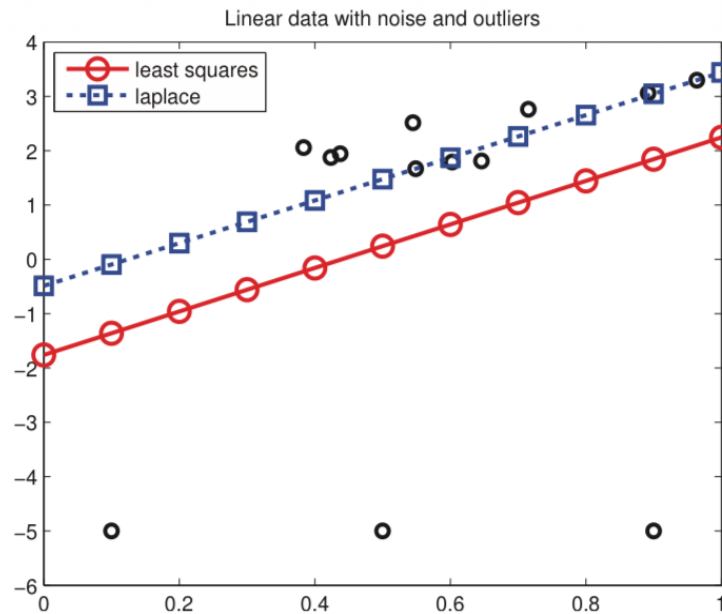
$$P(\underline{y} | \underline{x}; \underline{w}) = \mathcal{N}(\underline{w}^T \underline{x}, \sigma^2)$$

Log-likelihood function:

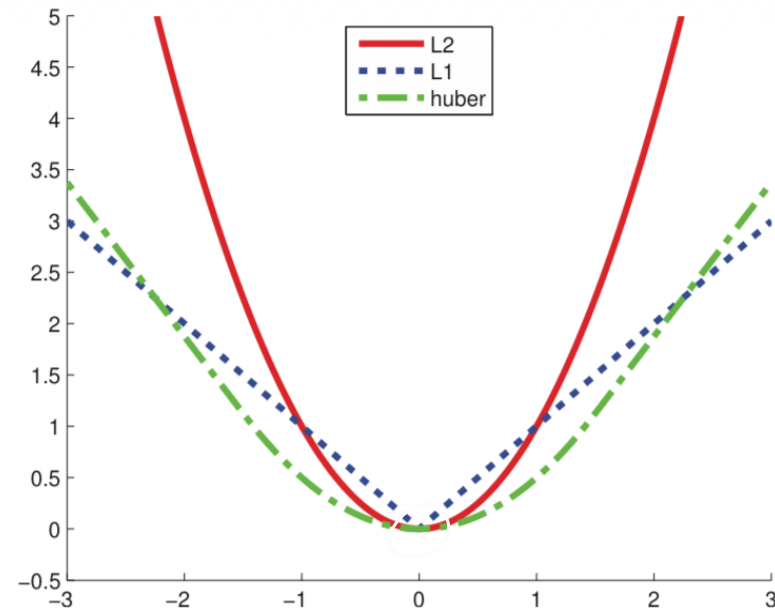
$$\mathcal{L}(\underline{w}) = \sum_{i=1}^N \log \left(\left(\frac{1}{2\pi\sigma^2} \right)^{\frac{1}{2}} \exp \left(-\frac{1}{2\sigma^2} (y_i - \underline{w}^T \underline{x}_i)^2 \right) \right)$$

$$= -\frac{1}{2\sigma^2} \sum_{i=1}^N (y_i - \underline{w}^T \underline{x}_i)^2 - \frac{N}{2} \log(2\pi\sigma^2)$$

Robust linear regression:



(a)



(b)

To account for outliers, use distributions with heavy tail:

$$p(\underline{y} | \underline{x}; \underline{w}, b) = \text{Lap}(\underline{y} | \underline{w}^T \underline{x}, b) \\ \propto \exp\left(-\frac{1}{b} |y - \underline{w}^T \underline{x}|\right)$$

Assume b is fixed.

$$(P): \min_{\underline{\omega}} \quad \ell(\underline{\omega}) = \sum_{i=1}^N |\gamma_i(\underline{\omega})|$$

$$\gamma_i = \gamma_i^+ - \gamma_i^- \quad \gamma_i^+ \geq 0$$

$$\gamma_i^- \geq 0$$

$$|\gamma_i| = \gamma_i^+ + \gamma_i^-$$

$$\underline{\omega}^T x_i + \gamma_i^+ - \gamma_i^- = y_i$$

$\equiv (P)$

$$\min_{\underline{\omega}, \underline{\gamma}^+, \underline{\gamma}^-}$$

$$\sum_{i=1}^N \gamma_i^+ + \gamma_i^-$$

Linear program:

s.t.

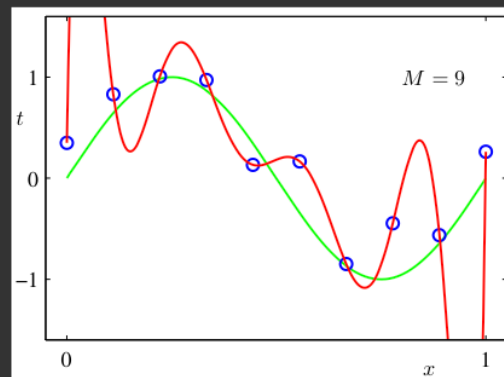
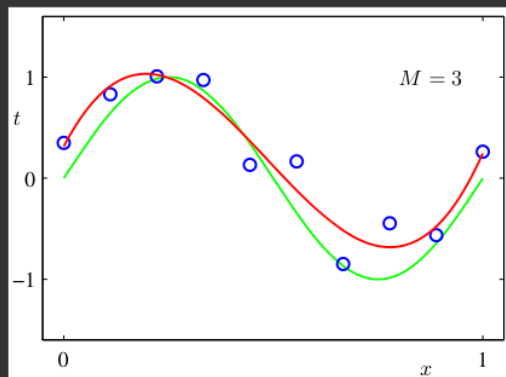
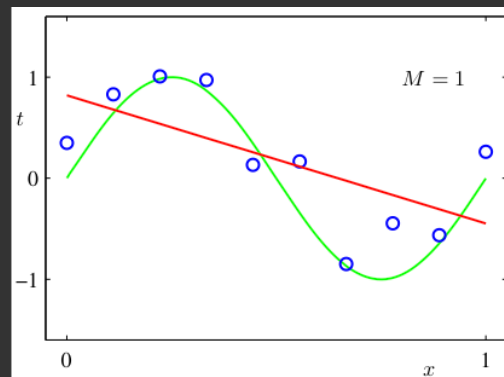
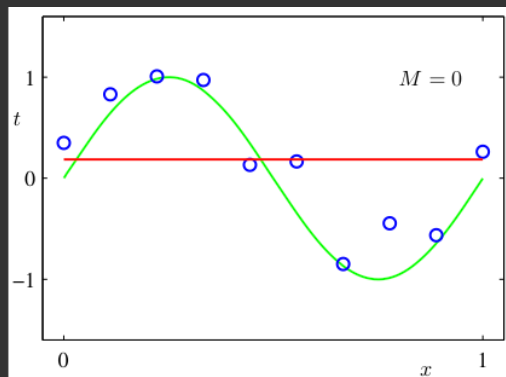
$$\gamma_i^+ \geq 0$$

$$\gamma_i^- \geq 0$$

$$\underline{\omega}^T x_i + \gamma_i^+ - \gamma_i^- = y_i \quad , \quad i=1, \dots, N$$

Ridge regression:

$$h_w(x) = \sum_{j=0}^M w_j x^j$$



Add a regularizer to the error to control over-fitting:

$$E_D(\underline{w}) + \lambda E_w(\underline{w})$$

$\underbrace{E_D(\underline{w})}_{\text{data-dependent}}$ $\underbrace{E_w(\underline{w})}_{\text{regularizer}}$

ℓ_2 -norm regularizer:

$$E_w(\underline{w}) = \frac{1}{2} \underline{w}^T \underline{w} = \frac{1}{2} \|\underline{w}\|_2^2$$

Loss:

$$L(\underline{w}) = \frac{1}{2} \sum_{i=1}^N (y_i - \underline{w}^T \underline{x}_i)^2 + \frac{\lambda}{2} \underline{w}^T \underline{w}$$

minimize $L(\underline{w})$
 $\underline{w}$

Ridge regression
or
weight decay
or
parameter
shrinkage

$\equiv$

minimize $L_D(\underline{w})$
 $\underline{w}$

s.t. $\|\underline{w}\|_2^2 \leq \alpha$

sol:

$$\frac{\partial L}{\partial \underline{w}} = 0$$

$$\frac{1}{2} \cdot 2 \sum_{i=1}^N (-x_i)(y_i - \underline{w}^T x_i) + \frac{\lambda}{2} \cdot 2 \underline{w} = 0$$

$$\left(\sum_{i=1}^N x_i x_i^T + \lambda I \right) \underline{w} = \sum_{i=1}^N y_i x_i$$

$$\underline{w}_{\text{ridge}} = (\lambda I + X^T X)^{-1} X^T \underline{y}$$

Bayesian view (Maximum a posteriori estimate.)

$$P(\underline{w} | \underline{x}, \underline{y}) \propto P(\underline{y} | \underline{x}, \underline{w}) P(\underline{w})$$

Suppose

$$P(\underline{y} | \underline{x}, \underline{w}, \beta) = \prod_{n=1}^N \mathcal{N}(y_n | \underline{w}^T \underline{x}_n, \beta^{-1})$$

$$P(\underline{w} | \alpha) = \mathcal{N}(\underline{w} | 0, \alpha^{-1} \mathbf{I})$$

$$= \left(\frac{\alpha}{2\pi} \right)^{\frac{(D+1)}{2}} \exp \left(-\frac{\alpha}{2} \underline{w}^T \underline{w} \right)$$

$$\text{MAP: } \max_{\underline{w}} P(\underline{w} | \underline{x}, \underline{y}, \alpha, \beta)$$

$$= \min_{\underline{w}} \frac{\beta}{2} \sum_{i=1}^N (y_i - \underline{w}^T \underline{x}_i)^2 + \frac{\alpha}{2} \underline{w}^T \underline{w}$$

$$\lambda = \alpha / \beta$$

Bias - variance trade-off:

(model complexity)

How to determine λ ?

$\lambda = 0$ over fitting

$\lambda \uparrow$ under fit.

$$y = f(x) + \varepsilon$$

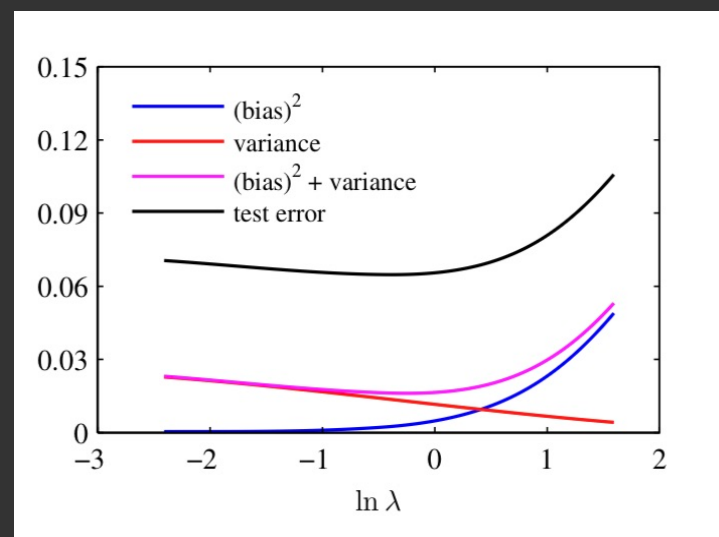
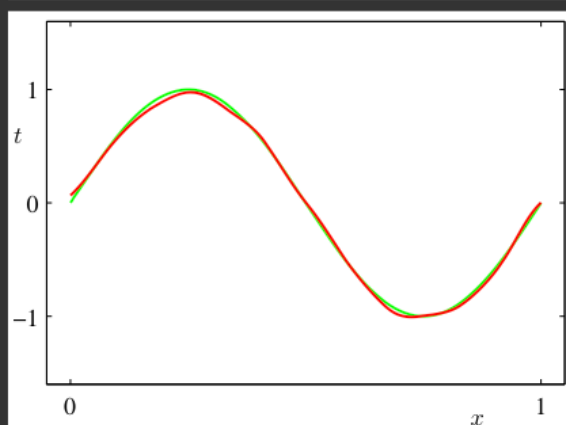
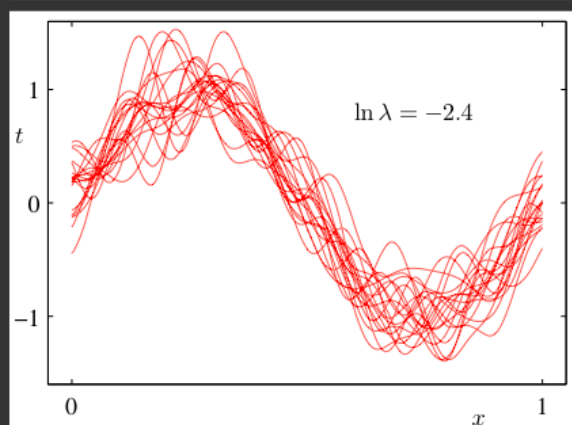
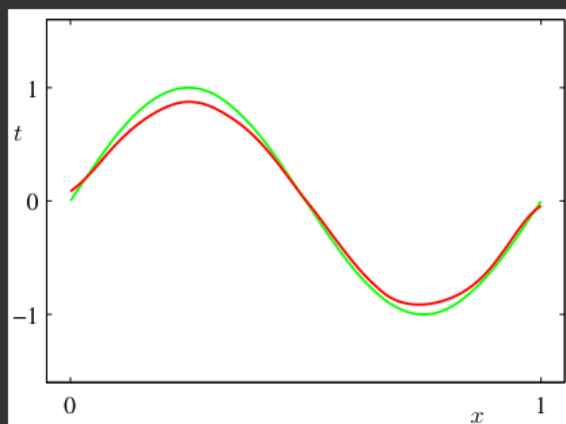
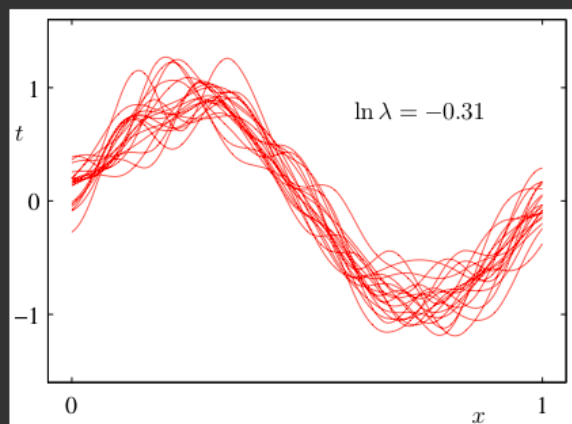
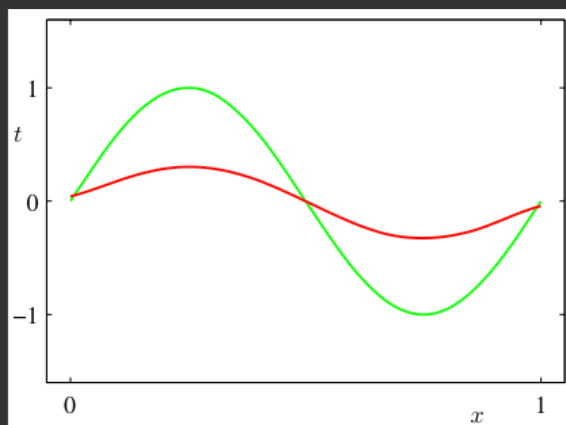
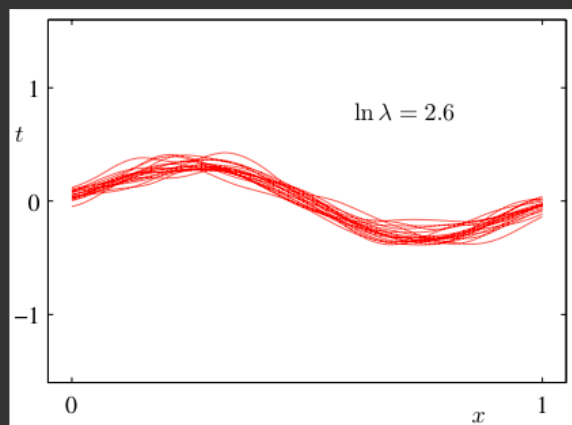
$$\text{var}(y) = \sigma^2$$

$$E[(y - \hat{f})^2] = E[(f(x) + \varepsilon - \hat{f})^2]$$

$$= E[y^2] + E[\hat{f}^2] - E[2y\hat{f}]$$

$$= \text{var}[y] + E[y]^2 + \text{var}[\hat{f}] + E[\hat{f}]^2 - 2f E[\hat{f}]$$

$$= \underbrace{\sigma^2}_{\text{irreducible error}} + \underbrace{\text{var}(\hat{f})}_{\text{variance}} + \underbrace{(\hat{f} - E[\hat{f}])^2}_{(\text{bias})^2}$$



Other regularization:

$$E_{\omega}(\underline{\omega}) = \frac{\lambda}{2} \sum_{j=1}^D |\omega_j|^q$$

$$q = 1$$

LASSO:

$$L(\underline{\omega}) = E_D(\omega) + \lambda \|\underline{\omega}\|_1$$

$$q = 0.5$$

