

Lecture FYS-
STK3155/4155,
September 16, 2024

FYS-STK3155/4155, September 16

kaggle.com

Linear regression

$$y_i = f(x_i) + \varepsilon_i$$

$\varepsilon_i \sim N(0, \sigma^2)$
 $[x_a, x_b]$

$$x_i \in (-\infty, \infty)$$

$$y_i \in (-\infty, \infty) \quad [y_a, y_b]$$

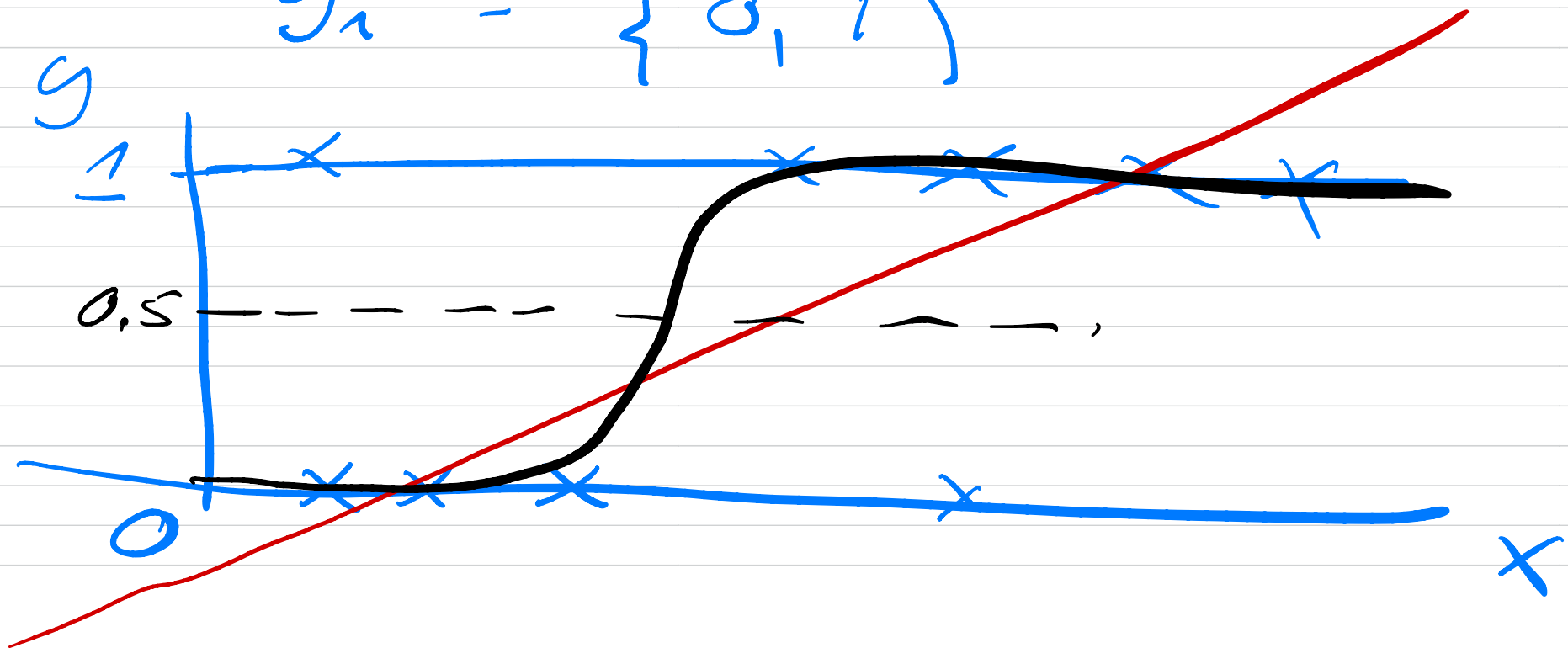
$$y_i \approx \underbrace{\sum_j x_{ij} \beta_j}_{\tilde{y}_i} + \varepsilon_i$$

First-order polynomial

$$y_i \approx \beta_0 + \beta_1 x_i + \varepsilon_i$$

What about y_i given
by discrete values?

$$y_i = \{0, 1\}$$



$$y_i = f(x_i) + \varepsilon_i \rightarrow$$

$$y_i = p(x_i) + \varepsilon_i$$

$$\int_{x \in D} p(x) dx = 1 \quad \left\{ \begin{array}{l} \sum_{i \in D} p(x_i) = 1 \end{array} \right.$$

$$x \in [0, 1] \quad p(1) = 1$$

$$p(x_i) \leq p(x_j) \quad x_i \leq x_j$$

Example : Sigmoid function

$$p(x) = \frac{e^{\beta x}}{1 + e^{\beta x}}$$

$$x \in [0, a] \quad \beta = 1$$

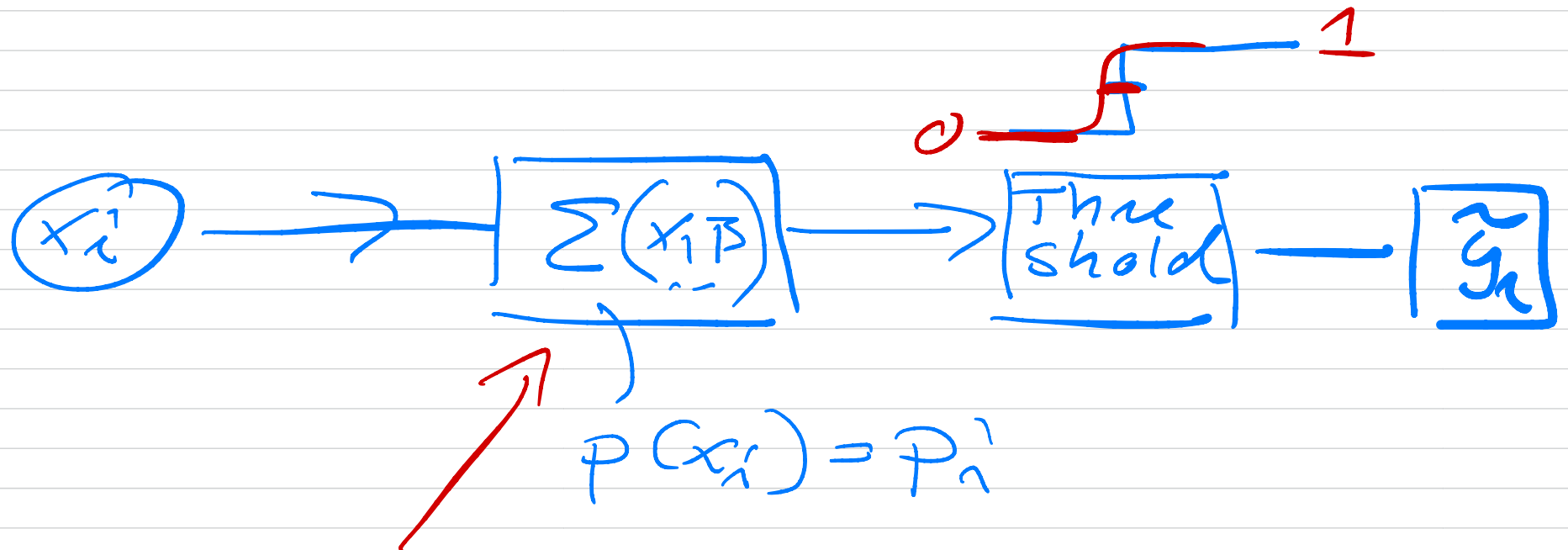
$$\int_0^a dx \frac{e^x}{1 + e^x} = \ln(1 + e^a)$$

$$y = 1 ; \quad p(x) = \frac{e^{\beta x}}{1 + e^{\beta x}}$$

↓

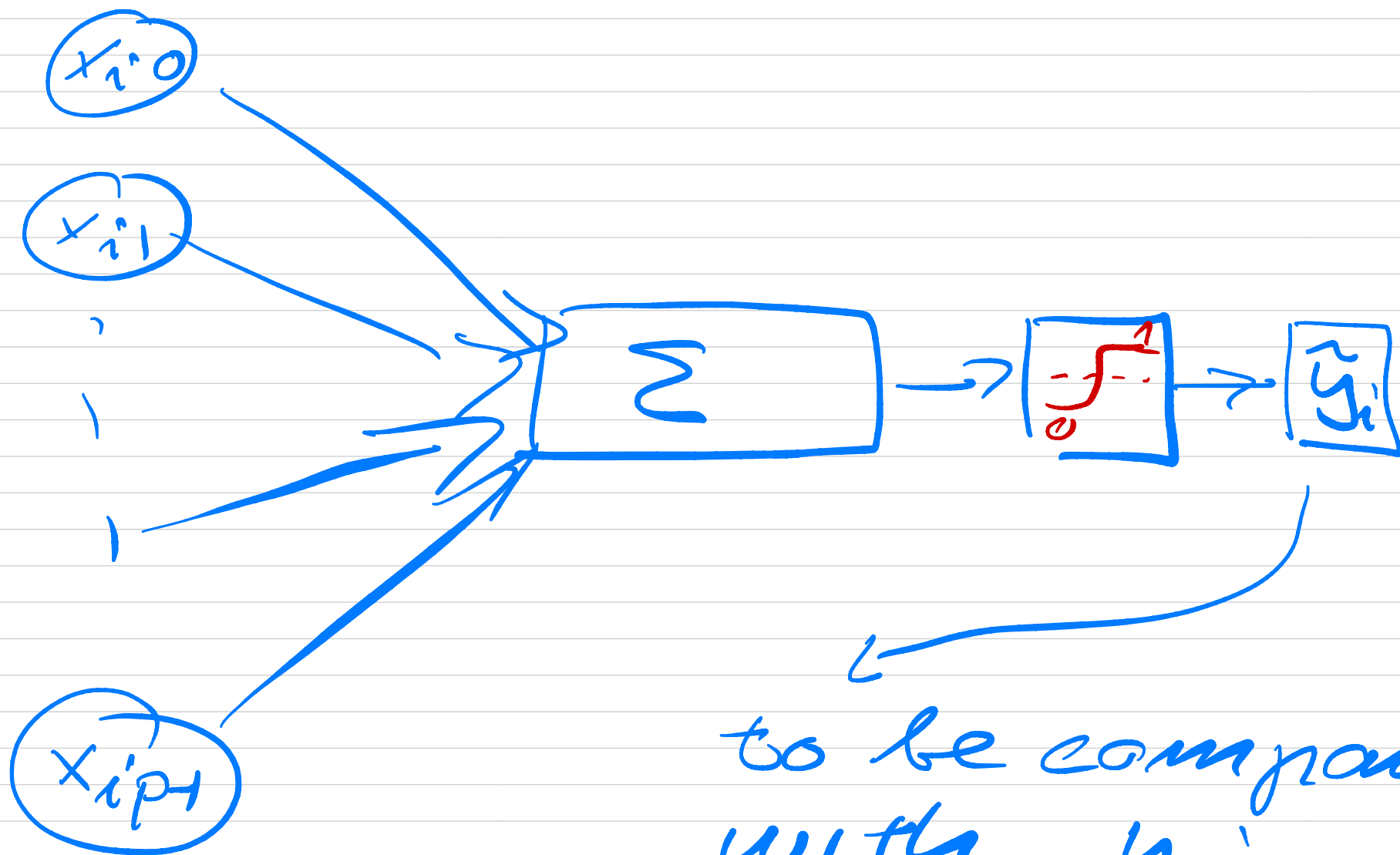
$$y = 0 ; \quad 1 - p(x) = \frac{1}{1 + e^{\beta x}}$$

Graphical representation



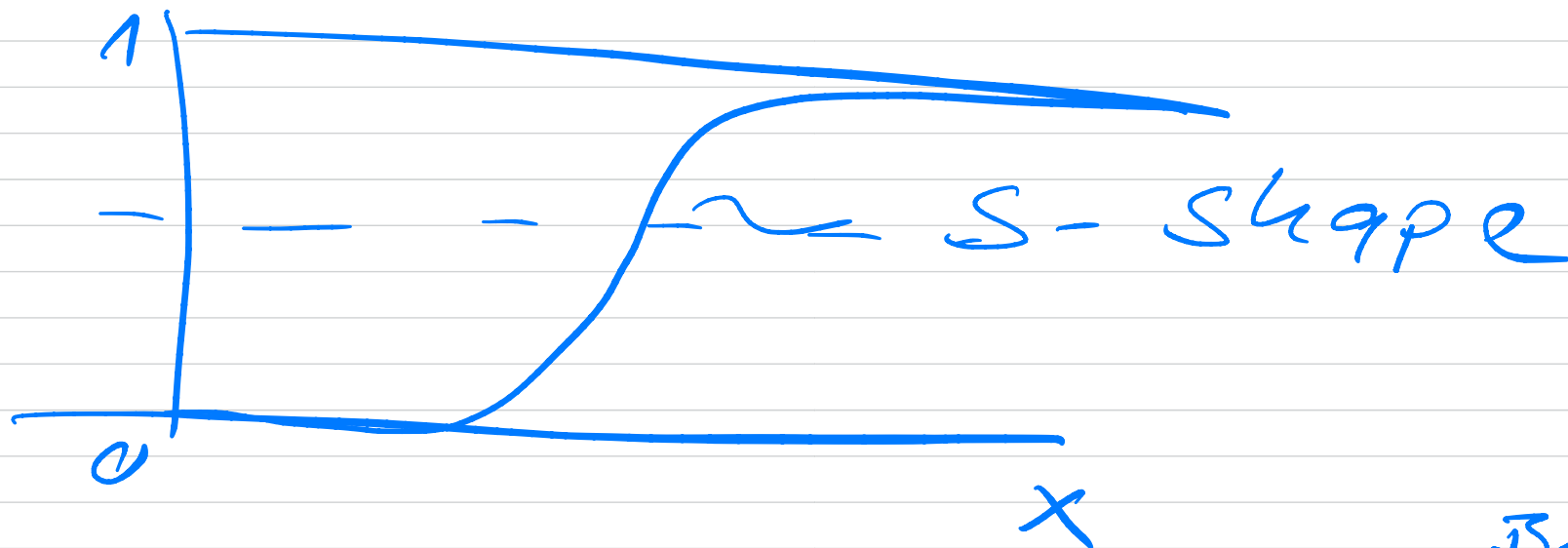
Neural networks
activation
function

$$X = \begin{bmatrix} x_{00} & x_{01} & - & - & x_{0p-1} \\ x_{10} & & & & \\ x_{20} & & & & \\ \vdots & & & & \\ x_{i0} & x_{i1} & x_{i2} & - & x_{ip-1} \\ \vdots & & & & \\ x_{n-10} & - & - & - & x_{n-1p-1} \end{bmatrix}$$



to be compared
with y_i

$$P(x_i) = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}$$



$$y_i = 1 ; P(x_i) = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}$$

$$y_i = P_i + \varepsilon_i = P_i$$

what probability does ε_i' follow?

$$y_i = 1 = p_i' + \varepsilon_i'$$

$$y_i = 0 = p_i' + \varepsilon_i'$$

$$\varepsilon_i' = 1 - p_i' \quad \text{when } y_i = 1$$

$$\varepsilon_i' = -p_i' \quad \text{when } y_i = 0$$

$$E[\varepsilon_i] = p_i(1 - p_i) + (-p_i)(1 - p_i)$$

$$\left(E[X] = \sum_{i=0}^{n-1} x_i p(x_i) \right)$$

$$= 0$$

$$\begin{aligned} \text{Var}[\varepsilon_i] &= (1 - p_i)^2 p_i + (-p_i)^2 (1 - p_i) \\ &= p_i(1 - p_i) \end{aligned}$$

ε_i follows a binomial distribution.

How do we find β ?

y_i are i.i.d.

$$D = \{ (x_0, y_0), (x_1, y_1), \dots, (x_{n-1}, y_{n-1}) \}$$

$$P(D | \beta) = \prod_{i=0}^{n-1} P_i^{y_i} (1 - P_i)^{1-y_i}$$

$\nwarrow$ $\beta_0 + \beta_1 x_i$

$$y_i = \{0, 1\}$$

$$P_i = \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}$$

$$\hat{\beta} = \arg \max_{\beta \in \mathbb{R}^D} \underbrace{P(D/\beta)}_{C(\beta)}$$

$$\hat{\beta} = \arg \min_{\beta \in \mathbb{R}^D} \underbrace{(-\log P(D/\beta))}_{C(\beta)}$$

$$\Rightarrow C(\beta) = - \sum_{i=0}^{n-1} \left[y_i \log p_i + (1-y_i) \log (1-p_i) \right]$$

$$\log \left[\frac{e^{\beta_0 + \beta_1 x_1'}}{1 + e^{\beta_0 + \beta_1 x_1'}} \right]$$

$$= (\beta_0 + \beta_1 x_1') - \log(1 + e^{\beta_0 + \beta_1 x_1'})$$

$$\Rightarrow C(\beta) = - \sum_{i=0}^{n-1} \left[y_i' (\beta_0 + \beta_1 x_{1i}') - \log(1 + e^{\beta_0 + \beta_1 x_{1i}'}) \right]$$

$$\frac{\partial C}{\partial \beta_0} = 0 \quad \wedge \quad \frac{\partial C}{\partial \beta_1} = 0$$

$$\frac{\partial C}{\partial \beta_0} = 0 = - \sum_{i=0}^{n-1} (y_i' - p_i')$$

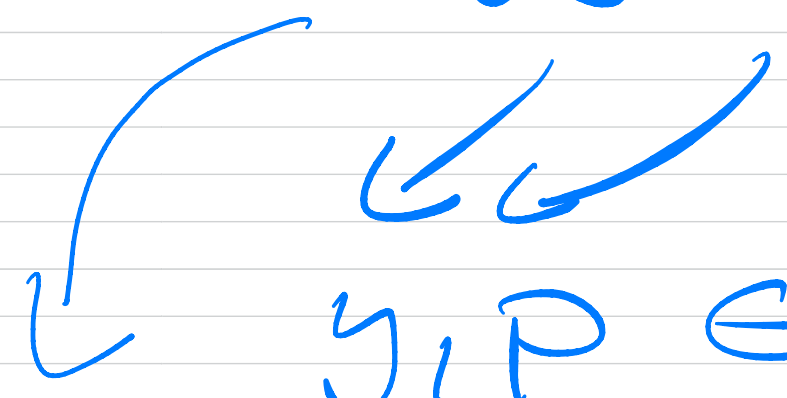
$$p_i' = \frac{e^{\beta_0 + \beta_1 x_i'}}{1 + e^{\beta_0 + \beta_1 x_i'}}$$

$$\frac{\partial C}{\partial \beta_1} = 0 = - \sum_{i=0}^{n-1} x_i' (y_i' - p_i')$$

in general for the

gradient $g = \frac{\partial C}{\partial \beta} =$

$$\frac{\partial \mathcal{L}}{\partial \beta} = -X^T (y - p)$$



$$X \in \mathbb{R}^{n \times p} \quad y, p \in \mathbb{R}^n \quad \Rightarrow$$

$$g = \frac{\partial \mathcal{L}}{\partial \beta} \in \mathbb{R}^p$$

$\hat{\beta}$ is included in $\mathcal{D}$!

in linear regression

$$\frac{\partial^2 C}{\partial \beta \partial \beta^T} = \frac{2}{n} X^T X$$

Hessian matrix

independent of β

$$\frac{\partial^2 C}{\partial \beta \partial \beta^T} = X^T W X \text{ depends on } \beta,$$

Log Reg

$$W_{ii} = P_i (1 - P_i)$$
$$W_{ij} (i \neq j) = 0$$

How do we solve the eqs for $\hat{\beta}$?

Newton-Raphson's method,

β - scalar, $\frac{\partial C}{\partial \beta} \Rightarrow \frac{dC}{d\beta}$

Taylor - expansion around
the optimal $\hat{\beta}$

$$\hat{\beta} - \beta^{(n)}$$

$$C(\hat{\beta}) = C(\beta^{(n)}) +$$

$$\underbrace{g(\beta^{(n)})}_{\left(\frac{dC}{d\beta}\right)\bigg|_{\beta=\beta^{(n)}}} (\hat{\beta} - \beta^{(n)})$$

$$+ \frac{1}{2} \underbrace{H(\beta^{(n)})}_{\frac{d^2 C}{d\beta^2}\bigg|_{\beta=\beta^{(n)}}} (\hat{\beta} - \beta^{(n)})^2 + \dots$$

$$b = \hat{\beta} - \beta^{(n)}$$

$$C(\hat{\beta}) \approx C(\beta^{(n)}) + \underbrace{g^{(n)}}_{g(\beta^{(n)})} b$$

$$+ \frac{1}{2} b^2 \underbrace{H^{(n)}}_{H(\beta^{(n)})}$$

$$\frac{dC}{db} = g^{(n)} + H^{(n)} b = 0$$

$$b = \hat{\beta} - \beta^{(n)} = -\left(H^{(n)}\right)^{-1} g^{(n)}$$

$$\hat{\beta} = \beta^{(n+1)}$$

$$\beta^{(n+1)} = \beta^{(n)} - (H^{(n)})^{-1} g^{(n)}$$

$$n = 0, 1, \dots$$

- (i) start with a guess $\beta^{(0)}$
- (ii) iterate over n till

$$|\beta^{(n+1)} - \beta^{(n)}| \leq \epsilon \sim 10^{-8}$$

in general, $H(\beta^{(n)})$ is
a matrix, $g(\beta^{(n)})$ is a
vector

$$g \in \mathbb{R}^P$$

$$H \in \mathbb{R}^{P \times P}$$

$$\beta \in \mathbb{R}^P$$

$$\beta^{(n+1)} = \beta^{(n)} - (H(\beta^{(n)})^{-1}) g(\beta^{(n)})$$

Replace $H(\beta^{(n)})$ with
a parameter $\gamma^{(n)}$

$$\beta^{(n+1)} = \beta^{(n)} - \boxed{\gamma^{(n)}} g^{(n)}$$

learning rate

γ

Can we find an optimal

γ ?

Expand $C(p)$ around

$$p^{(n)} - \gamma g^{(n)}$$

$$C(p^{(n)} - \gamma g^{(n)}) = C(p^{(n)})$$

$$- \gamma (g^{(n)})^T (g^{(n)})$$

$$+ \frac{1}{2} \gamma^2 [g^{(n)}]^T H^{(n)} g^{(n)} + \dots$$

$$g^{(n)} \rightarrow g \quad H^{(n)} \rightarrow H$$

Derivative wrt γ

$$\frac{\partial c}{\partial \gamma} = 0 \Rightarrow$$

$$\gamma = \frac{g^T g}{g^T H g} \quad (\gamma^{(n)})$$

Assume

$$H g = \lambda g$$

$$\gamma = \frac{g^T g}{g^T \lambda g} = \frac{1}{\lambda} \Rightarrow$$

$$\lambda_{\min} = \frac{1}{\lambda_{\max}}$$

$$\lambda_{\max} = \frac{1}{\lambda_{\min}}$$

if you can diagonalize $H^{(n)}$

$$\lambda < \frac{2}{\lambda_{\max}}$$