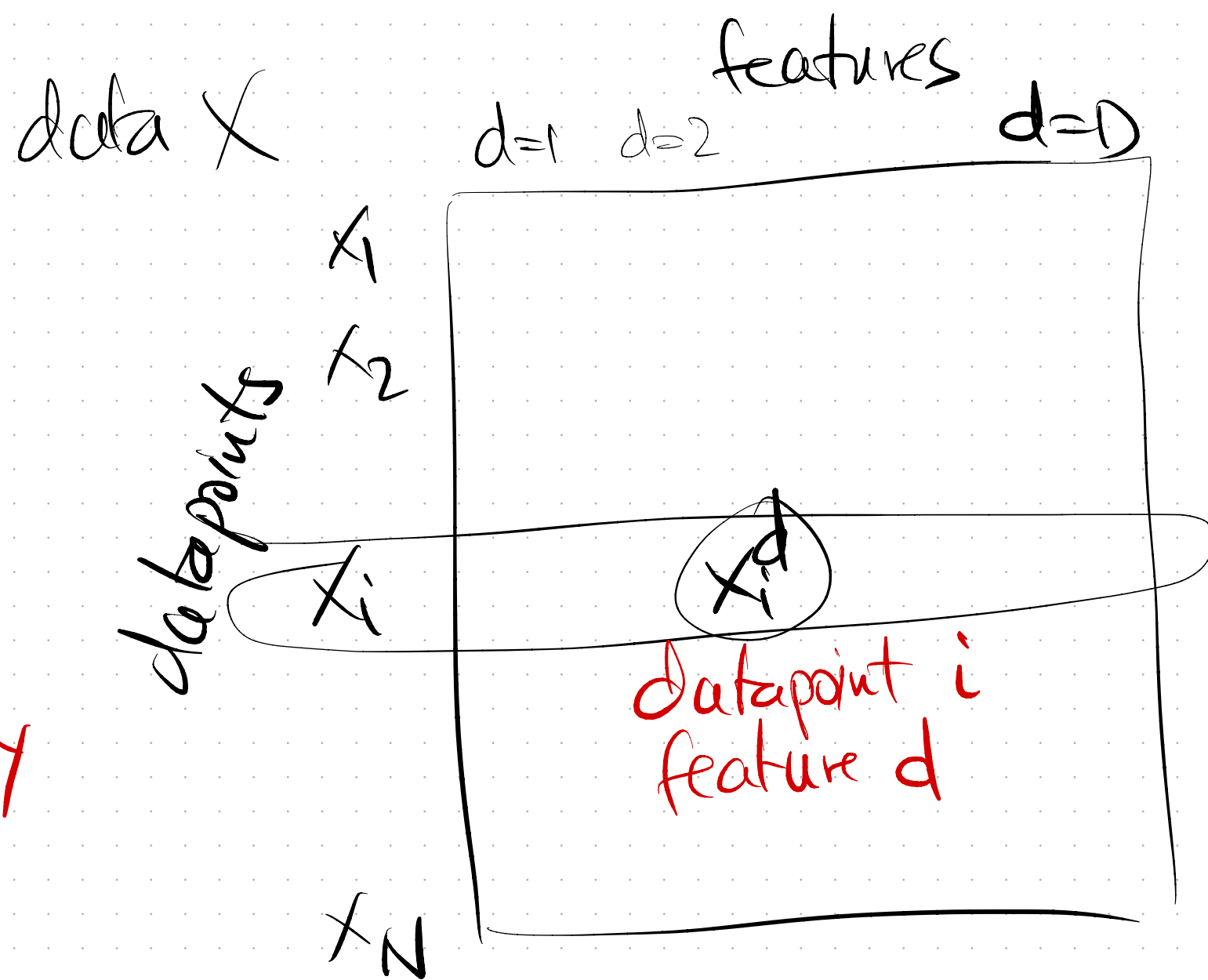


Lecture 5/21-23

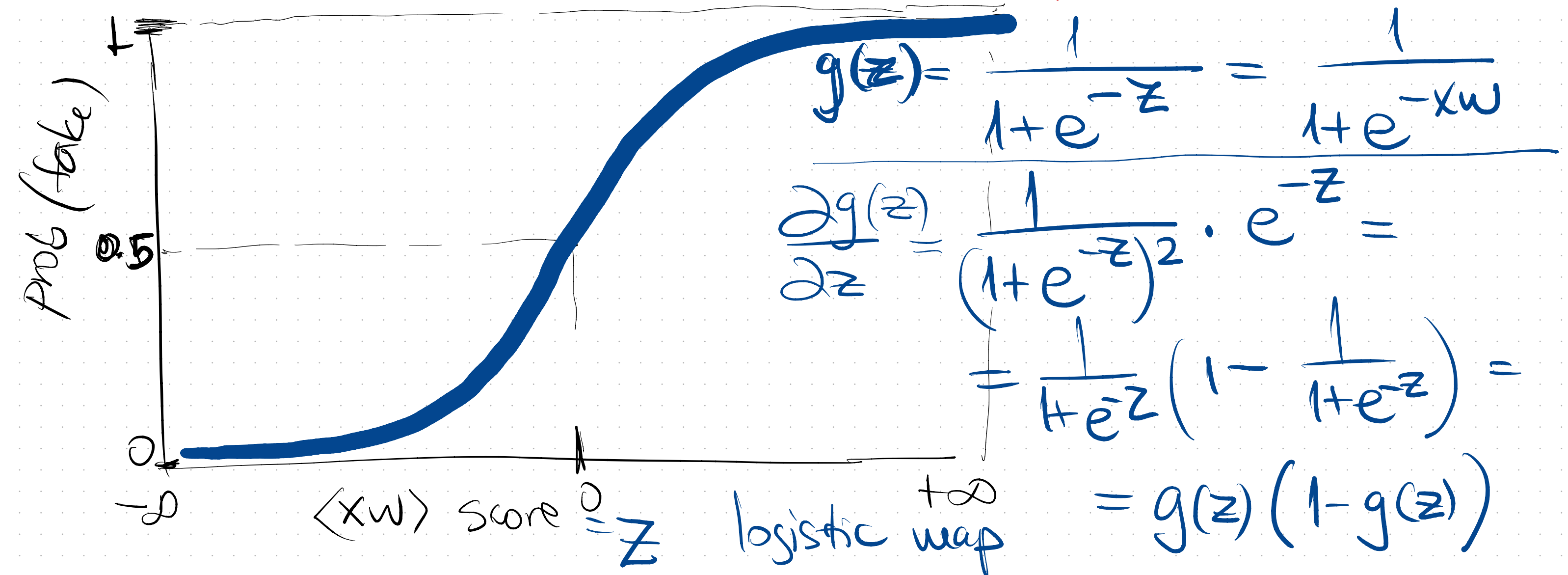
- Gradient Descent
- Linear Regression with GD
- Logistic Regression
- HW 2A (released later) due 6/3 today
- ROC curve $\approx$ AUC geometrically
- Perceptron w/GD or geometry



Logistic Regression : deal with classification classes $Y \rightarrow \begin{matrix} \text{yes} \\ \text{no} \end{matrix}$
 ex. Spam/ham dataset

Lin. Regression $h(x) = xw \approx Y = \text{quantity label (ex. house price)}$

Idea : $xw = \text{linear regression score}$
 • $h(x) = \text{map} \left\{ \begin{matrix} \text{linear score} \rightarrow \text{probability of } Y = \text{yes} \end{matrix} \right\}$ } $\text{prob}(\text{label} = \text{yes})$
 logistic map "g"
 $Z = \text{score} = xw$



- key property: overshooting does not penalize OBI

Lin Regression

want $xw \approx y$

if $xw \ll 0$: error $(xw - 0)^2$

if $xw \gg 1$ error $(xw - 1)^2$

$$y = \begin{cases} 0 \\ 1 \end{cases}$$

Logistic Regression

want $\frac{1}{1 + e^{-xw}} \approx y$

xw very small $\rightarrow -\infty$
 $\frac{1}{1 + e^{\infty}} \approx 0$

xw very large $\rightarrow +\infty$
 $\frac{1}{1 + e^{-\infty}} \approx 1$

Logistic Regression (coef w)

$$h(x) = \text{logistic}(\text{regscore}(x)) = \text{logistic}(xw)$$

$$= \frac{1}{1 + e^{-xw}}$$

interpret

probab (y = yes) $\rightarrow$ class

Logistic Regression uses log-likelihood OBJECTIVE (not sq error)

WANT $h(x_i) = g(x_i w) = \frac{1}{1 + e^{-x_i w}} \approx \text{prob}(Y_i = 1)$ 2 classes

$$Y_i = \begin{cases} 1 & \text{yes} \\ 0 & \text{no} \end{cases}$$

$$\begin{aligned} P(Y_i = 1) &\approx h(x_i) \\ P(Y_i = 0) &\approx 1 - h(x_i) \end{aligned} \Rightarrow P(Y | x_i) = h(x_i) \cdot (1 - h(x_i))$$

correct pred

$$h(x_i) \cdot (1 - h(x_i))^{1 - y_i} = \begin{cases} h(x_i) & \text{if } y_i = 1 \\ 1 - h(x_i) & \text{if } y_i = 0 \end{cases}$$

y_i 0/1 filters

if $y_i = 1$ (yes) $\Rightarrow$ want $h(x_i) = \text{high}$; $1 - h(x_i)$ low
 $y_i = 0$ (no) $\Rightarrow$ want $h(x_i) = \text{low}$; $1 - h(x_i)$ high

log likelihood

"LL"

$$\text{OBS}(w) = \prod_{i=1}^N P(Y_i | x_i)$$

product of probs of correct prediction

Want MAX $LL(w)$

$\log$ $\downarrow$ all datapoints in training

$\rightarrow$ datapoints indep. of each other.

$$= \log \left(\prod_{i=1}^N h(x_i)^{y_i} (1-h(x_i))^{1-y_i} \right)$$

$$= \sum_{i=1}^N \log (h(x_i)^{y_i} (1-h(x_i))^{1-y_i})$$

$$= \sum_{i=1}^N (\log(h(x_i)^{y_i}) + \log((1-h(x_i))^{1-y_i}))$$

$$= \sum_{i=1}^N [y_i \log(h(x_i)) + (1-y_i) \log(1-h(x_i))]$$

all train data

gradient

$$\frac{\partial LL}{\partial w} = !$$

GD update

iterative

$$w_{new} = w + \lambda \frac{\partial LL}{\partial w}(w) \quad (\text{"ascent"})$$

$$\log(abc) = \log(a) + \log(b) + \log(c)$$

$$\log(a^n) = n \log(a)$$

$$LL = \sum_{i=1}^n [y_i \log(h(x_i)) + (1-y_i) \log(1-h(x_i))] = \sum_{i=1}^n [y_i \log(g(x_i; w)) + (1-y_i) \log(1-g(x_i; w))]$$

one datapoint x_i :

$$\frac{\partial LL}{\partial w_d} = \left[y_i \frac{1}{g(x_i; w)} - (1-y_i) \frac{1}{1-g(x_i; w)} \right] \frac{\partial}{\partial w_d} g(x_i; w)$$

$$= \left[y_i \frac{1}{g(x_i; w)} - (1-y_i) \frac{1}{1-g(x_i; w)} \right] \cdot \frac{\partial g}{\partial (x_i; w)} \cdot \frac{\partial x_i^d}{\partial w_d}$$

rate of compound
differentials

$\frac{d}{dx_i}$

$$= \left[y_i \frac{1}{g(x_i; w)} - (1-y_i) \frac{1}{1-g(x_i; w)} \right] \cdot [g(x_i; w)(1-g(x_i; w))] x_i^d$$

$$= [y(1-g(x_i; w)) - (1-y_i) \cdot g(x_i; w)] x_i^d$$

$$= [y - \cancel{y_i g(x_i; w)} - g(x_i; w) + \cancel{y_i g(x_i; w)}] x_i^d$$

$$= [y - g(x_i; w)] x_i^d = [y - h(x_i)] x_i^d$$

GD update for datap. x_i : $w_{new}^d = w^d + \lambda (y_i - h(x_i)) x_i^d$

GD update for all datap: $w_{new}^d = w^d + \lambda \sum_{i=1}^N (y_i - h(x_i)) x_i^d$

GD update all comp(d)
all datap(i) $w_{new} = w + \lambda [y - h(x)] x$

Regularization Logistic Regression

$$OBJ = \text{Log Likelihood}(w) + L_2 ||w||^2 \rightarrow \text{easy}$$

$L_2: \sum_{d=1}^D (w_d)^2$

$\sum_{i=1}^N h(x_i)^{y_i} (1-h(x_i))^{1-y_i}$

$$OBJ = \text{Log Likelihood} + L_1 ||w||$$

$L_1: \sum_{d=1}^D |w_d| \rightarrow \text{not easy}$

ML training:

- $\frac{\partial OBJ}{\partial w} = \nabla_w$ (gradient)

- GD update $w_{\text{new}} = w + \lambda \cdot \nabla w$

Multiple classes (multinomial) $\{ \in \{1, 2, 3, \dots, k\} \}$ k classes.

- train a linear score (lin regression) for each class

$$f_k(x) = \langle x \cdot w_k \rangle = \sum_{d=1}^D x^d \cdot w_k^d \quad w_k = \text{vector of coef for class } k.$$

• ^{binary} logistic $\rightarrow$ ^{k classes} Softmax $f_k(x)$

$$\text{prob}(Y = y_k) = \frac{e^{f_k(x)}}{\sum_{t=1}^k \exp(f_t(x))}$$

↓
unnormalized

distribution over possibilities

$$Y = y_1 y_2 \dots y_k$$

Newton Update instead of GD

un'dim:

$$w_{\text{new}} = w - \frac{\nabla LL(w)}{\nabla^2 LL(w)}$$

1st diff

2nd diff

(twice differential)

$$LL = \text{obj}(w)$$

$$\nabla LL = \text{gradient } \frac{\partial LL}{\partial w}$$

$$\nabla^2 LL = \text{second grad } \frac{\partial^2 LL}{\partial w \partial w}$$

Intuition: $f(w) = LL(w)$. use Taylor series expansion for $f(w)$
at point w_n = current w at iteration n

$$f(w) = f(w_n) + \frac{\partial f(w_n)}{\partial w} \frac{(w - w_n)}{1!} + \frac{\partial^2 f(w_n)}{\partial w \partial w} \frac{(w - w_n)^2}{2!} + \frac{\partial^3 f(w_n)}{\partial w \partial w \partial w} \frac{(w - w_n)^3}{3!}$$

$$\approx f(w_n) + \frac{\partial f(w_n)}{\partial w} (w - w_n) + \frac{1}{2} \frac{\partial^2 f(w_n)}{\partial w \partial w} (w - w_n)^2 \quad (\text{two terms})$$

Solve for w
assume $f(w) \approx f(w_n)$
 $\partial = 0$

$$-\partial f = \partial^2 f (w - w_n)$$

$$w = w_n - \frac{\partial f}{\partial^2 f}(w_n)$$

multidim $w = (w^1 w^2 \dots w^D)$

first differential is $\frac{\nabla f}{\text{vector}} = \frac{\partial f}{\partial w} = \left(\frac{\partial f}{\partial w^1}, \frac{\partial f}{\partial w^2}, \dots, \frac{\partial f}{\partial w^D} \right)$

2nd differential is Hessian Matrix partial second deriv.

$$D^2 = \begin{bmatrix} \frac{\partial^2 f}{\partial w^1 \partial w^1} & \dots & \frac{\partial^2 f}{\partial w^1 \partial w^D} \\ \vdots & \ddots & \vdots \\ \frac{\partial^2 f}{\partial w^D \partial w^1} & \dots & \frac{\partial^2 f}{\partial w^D \partial w^D} \end{bmatrix}$$

GD update: $w_{\text{new}} = w - \text{learn rate} \cdot H^{-1} \nabla f(w)$

intuition $\rightarrow \frac{\partial f}{\partial w}$ 1st dif

$\rightarrow \frac{\partial^2 f}{\partial^2 w}$ 2nd dif

Newton: Hessian for $LL=OBJ$ of Logistic Regression

$$H = X^T \underline{DIAG} \cdot X$$

Diag:

