

datapoint x_i Regression line (classifier) by coef $W = (w^1 w^2 \dots w^D)$

$$h(x_i) = \langle x_i \cdot W \rangle = \sum_{d=1}^D x_i^d \cdot w^d$$

 want $\underbrace{h(x_i)}_{\text{predict}} \approx \underbrace{y_i}_{\text{label}}$

$$\begin{aligned}
 \text{sq error}_i &= (h(x_i) - y_i)^2 \\
 &= (\langle x_i W \rangle - y_i)^2
 \end{aligned}$$

$$x_i = (x_i^1 x_i^2 \dots x_i^D)$$

All datapoints $\Rightarrow$ matrix $X_{N \times D}$

	1	2	D
x_1	x_1^1	x_1^2	x_1^D
x_2	x_2^1	x_2^2	x_2^D
X			
x_N	x_N^1	x_N^2	x_N^D

Y
y_1
y_2
y_N

$$\begin{array}{|c|} \hline x_i \\ \hline x_1 \dots x_D \\ \hline N \\ \hline \end{array} \cdot \begin{array}{|c|} \hline w \\ \hline w_1 \\ \hline \vdots \\ \hline w_D \\ \hline D \times 1 \\ \hline \end{array} = \begin{array}{|c|} \hline x_1 w \\ \hline x_2 w \\ \hline \vdots \\ \hline x_i w \\ \hline \vdots \\ \hline x_N w \\ \hline \end{array} \text{ VS labels (diff) } \begin{array}{|c|} \hline y_1 \\ \hline y_2 \\ \hline \vdots \\ \hline y_N \\ \hline \end{array}$$

def

$$\boxed{X \cdot W - Y = E} \quad \begin{array}{|c|} \hline x_1 w - y_1 \\ \hline x_2 w - y_2 \\ \hline \vdots \\ \hline x_i w - y_i \\ \hline \vdots \\ \hline x_N w - y_N \\ \hline N \times 1 \\ \hline \end{array}; E^T = [x_1 w - y_1, x_2 w - y_2, \dots, x_N w - y_N]$$

want sq error: $J(w) = \frac{1}{2} \sum_{i=1}^N (h(x_i) - y_i)^2 = \frac{1}{2} E^T \cdot E = \frac{1}{2} (Xw - Y)^T (Xw - Y)$

$$\langle E^T \cdot E \rangle = \sum_{i=1}^N e_i \cdot e_i = \sum_{i=1}^N e_i^2 = \sum_{i=1}^N (x_i w - y_i)^2$$

want: to find w vector $w = (w^1 w^2 \dots w^D)$ to minimise the sq error $\frac{1}{2} (Xw - Y)^T \cdot (Xw - Y)$

CONVEX: $\frac{\partial^2 J}{\partial w} > 0$ positive

intuition $J(w) = \frac{1}{2}(Xw - y)^T(Xw - y)$ convex in w

$\Rightarrow \frac{\partial J}{\partial w} = 0$ when $J(w)$ is minimum.

Q?

$$\frac{\partial J}{\partial w} = \frac{\partial}{\partial w} (Xw - y)^T(Xw - y)$$

$$= \frac{\partial}{\partial w} (w^T X^T X w - w^T X^T y - y^T X w + y^T y)$$

1x1 dim (1xD)(DxD)(NxD)
same transposed

$$w^T X^T y$$

$$y^T X w$$

$$(AB)^T = B^T A^T$$

$$+ \frac{\partial}{\partial w} w^T X^T y + \frac{\partial}{\partial w} y^T X w$$

$\hookrightarrow$ L_2 reg term $J(w)$
convex

$$= \frac{\partial}{\partial w} (w^T X^T X w - 2 w^T X^T y - y^T y)$$

$$+ \frac{\partial}{\partial w} w^T w \cdot \lambda$$

pretend 1-dim

$$= \frac{\partial}{\partial w} w^T X^T X w - \frac{\partial}{\partial w} w^T X^T y$$

pretend is 1-dim(w)

pretend 1-dim

$$+ \lambda I \cdot w$$

Id. matrix

$$= \frac{\partial}{\partial w} 2x^T X w - x^T y + \lambda I \cdot w$$

$$= x^T X w - x^T y + \lambda I w$$

want $= 0 \rightarrow$ sq DxD, pos definite (pos eig val)

$$x^T X w = x^T y$$

invert-left by $(x^T X)^{-1}$

$$(x^T X + \lambda I) w = x^T y$$

$$w = (x^T X + \lambda I)^{-1} x^T y$$

$w = (x^T X)^{-1} x^T y$ NORMAL EQ (closed form)
Best (Th) w for min sq errors of reg. line $\hookrightarrow$ L_2 reg

1-dim w differential

$$\frac{\partial}{\partial w} (\underline{w}^T \cdot C \cdot \underline{w}) = \frac{\partial}{\partial w} (w^2 \cdot c) = \underline{2wc}$$

$\frac{\partial}{\partial w} = ? \Rightarrow$ let math people worry about this!
 $\frac{\partial}{\partial w} \downarrow$ vector D -dim $\Rightarrow$ if interested, talk about OI (not repure for course)

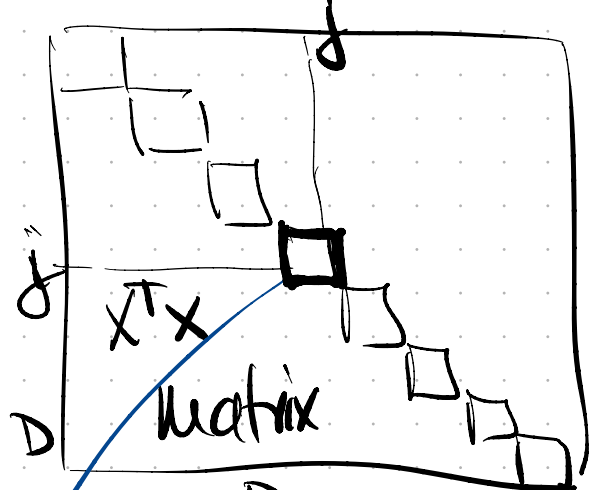
Naive way: we want $X \cdot W \approx Y$
why not (if $X \approx$ sq matrix) make it square?
 $N \times D \quad D \times 1 \quad N \times 1$

$W = X^{-1} Y$?
unstable,
 W not min sq err.
 X^{-1} terrible
to invert

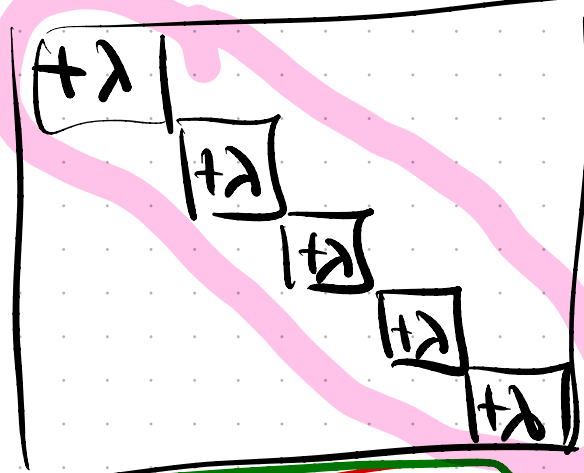
Correct way

$X^T X$ stable matrix, pos def $\rightarrow$
 $D \times D \quad \Rightarrow$ invert this matrix

$X = \text{centered}$ (all column mean $\rightarrow 0$) $\Rightarrow X^T X = N \cdot \text{COVAR}(X)$

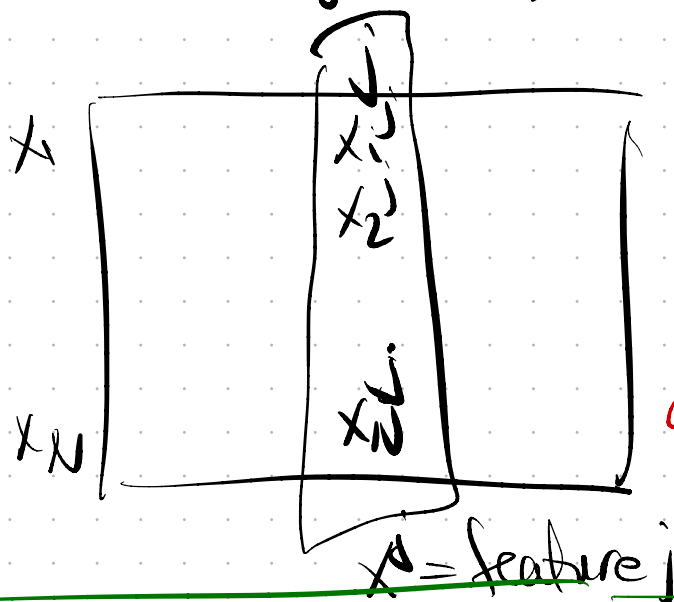


add stability
 (for inversion)
 $X^T X^{-1}$



Increase
 main
 diagonal

diagonal vals (pos j)
 $\langle \text{column } j, \text{column } j \rangle$
 $m \times$
 $\text{var}(j \text{ feature}) = \text{var}(x^j)$



this op $X^T X + \lambda I$ actually
 corresponds to regularized L_2 error

$$J(w) = \underbrace{\frac{1}{2} (Xw - y)^T (Xw - y)}_{\text{error}} + \underbrace{\frac{1}{2} \lambda \sum_{d=1}^D w_d^2}_{\text{OB}}$$

$$= \frac{1}{2} (Xw - y)^T (Xw - y) + \frac{1}{2} \lambda w^T w$$

prevents large w values $\Leftarrow L_2$ regularization

$\lambda = \text{hyperparam of } L_2 \text{ reg}$

Regression Visual intuition

same plane
 $+b=bias$

NEG

high
 $x \cdot w$

$x \cdot w$

$w =$
normal
to hyperplane

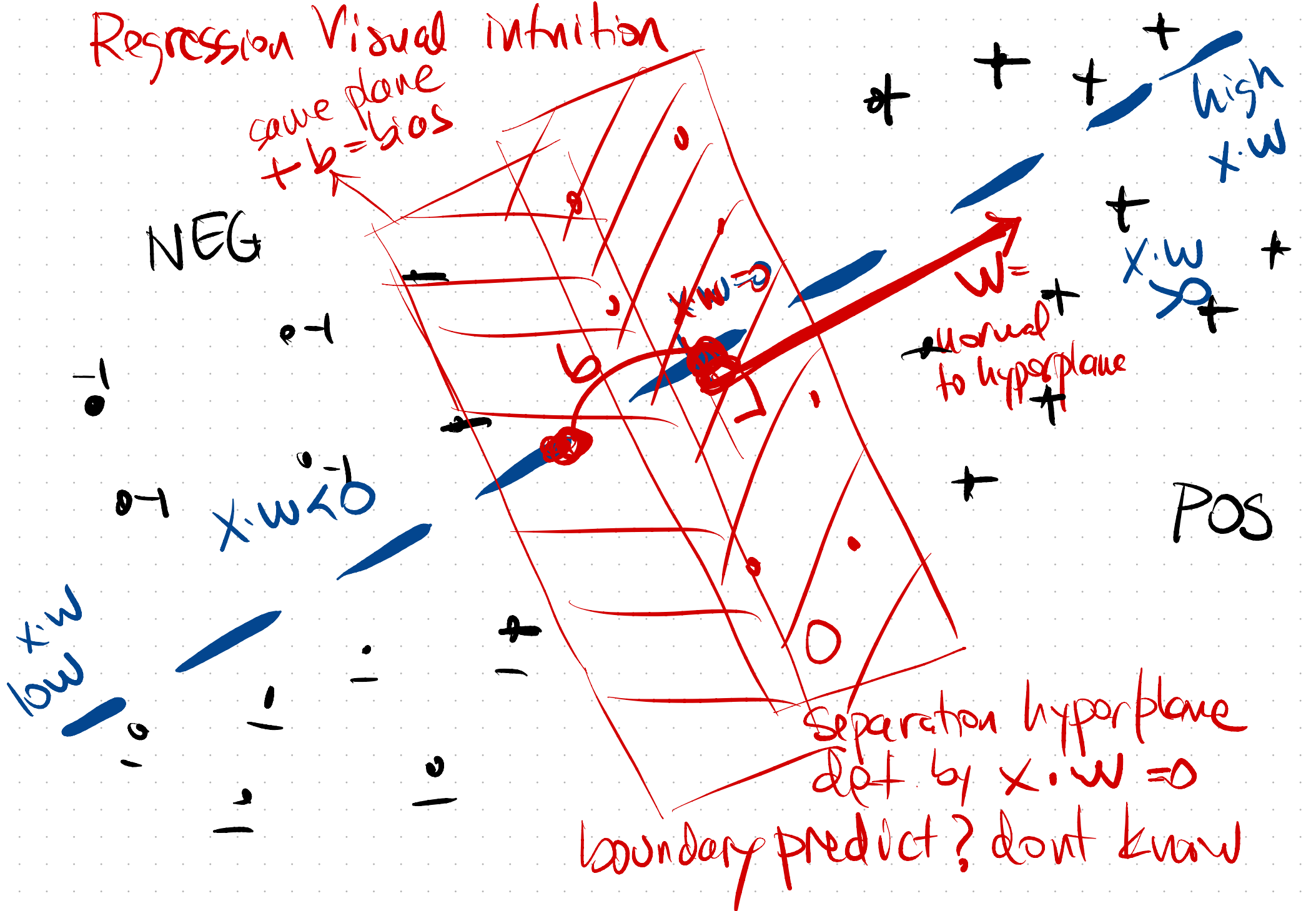
POS

$x \cdot w < 0$

low
 $x \cdot w$

Separation hyperplane
def by $x \cdot w = 0$

boundary predict? dont know



HW 1 BB 3: 1-dim $w = a$ bias b

Reg $h(x_i) = x_i a + b$ $\sum_{i=1}^n y_i$ $\forall i=1:n$

1-dim "1" vector

$$\begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_N \end{bmatrix} \cdot \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix}$$

$$\cdot \begin{bmatrix} a & b \end{bmatrix} =$$

$$\begin{bmatrix} x_1 a + b \\ x_2 a + b \\ \vdots \\ x_N a + b \end{bmatrix}$$

vs

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix}$$