

Ridge Regression  
MAP  
kernels

- Thanks for staying/auditing

- Ass 2 due Monday  
(Tutorial tomorrow)

- Graded Ass1 due Monday

- Notes on webpage:

- Gaussian MLE

- Solving least squares

- You do not need to know

linear programming,

but you do need to know the trick



Monday

## Converting to constrained problem

- non-smooth optimization is hard
- smooth constrained optimization often simpler

due Monday

Trick to convert (1 of many):

$$\max \{x, -x\} = |x|$$

If we have  $\min_x \max_i \{f_i(x)\}$

Equivalent

$\min_{x, v} v$  subject to  $v \geq \max_i \{f_i(x)\}$

$$v \geq f_i(x), \forall i$$

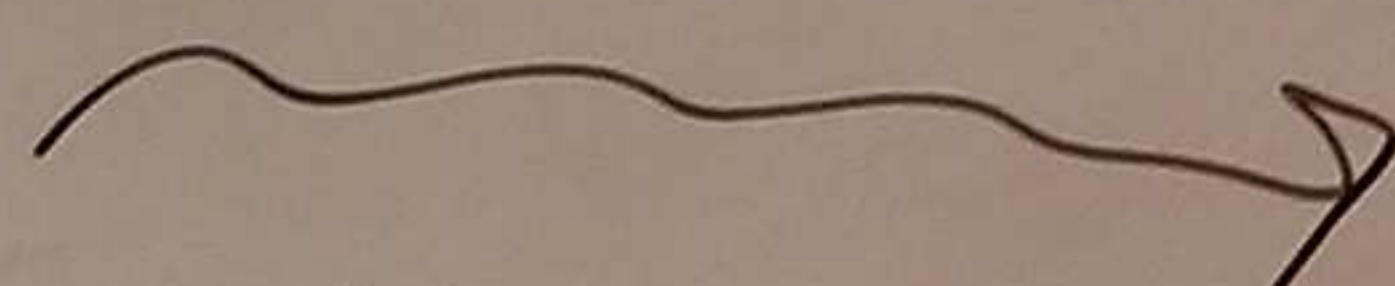


# Change of Basis

Original data:

$$X = \begin{bmatrix} 1 \\ .5 \\ 2 \\ \vdots \\ 10 \end{bmatrix}$$

make new X



" $\Phi(X)$ "

$\bar{x}_i \rightsquigarrow \phi(\bar{x}_i)$

template

$$[1 \quad x \quad x^2]$$

$$\begin{bmatrix} 1 & 1 & 1 \\ 1 & .5 & .25 \\ 1 & 2 & 4 \\ \vdots & \vdots & \vdots \\ 1 & 10 & 100 \end{bmatrix}$$

new design  
matrix

(then just business  
as usual)

$$= |X|$$



late  
 $x^2$  ]

1  
.25  
4  
:  
100 ]

design  
trix

## Cross-Validation

- large 'k': less chance of overestimating test error, but higher variance + running time.

- idea of repeating CV with different random partitions: lower variance.



Ridge Regression

MAP

kernels

Ridge Regression

$$f(\bar{w}) = \frac{1}{2} \|X\bar{w} - Y\|^2 + \frac{\lambda}{2} \|\bar{w}\|^2$$

$$\nabla f(\bar{w}) = X^T$$

What? Shrink  $\bar{w}_i$  toward 0

Why?

- solution is unique

- avoids overfitting due to huge values

- improves conditioning

• "Magic": do a huge basis expansion



ession

$$\frac{1}{2} \|\bar{w}\|^2$$

$$\nabla f(\bar{w}) = X^T X \bar{w} - X^T Y + \lambda \bar{w}$$

$$= 0 \quad (\text{at stationary point})$$

$$(X^T X + \lambda I) \bar{w} = X^T Y$$

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

$$\nabla^2 f(w) = X^T X + \lambda I \succ 0$$

"positive  
definite"

Solution is unique

$$\bar{w} = X^T (X X^T)^{-1} Y$$

$$\frac{1}{N} X^T Y$$

$$X X^T$$

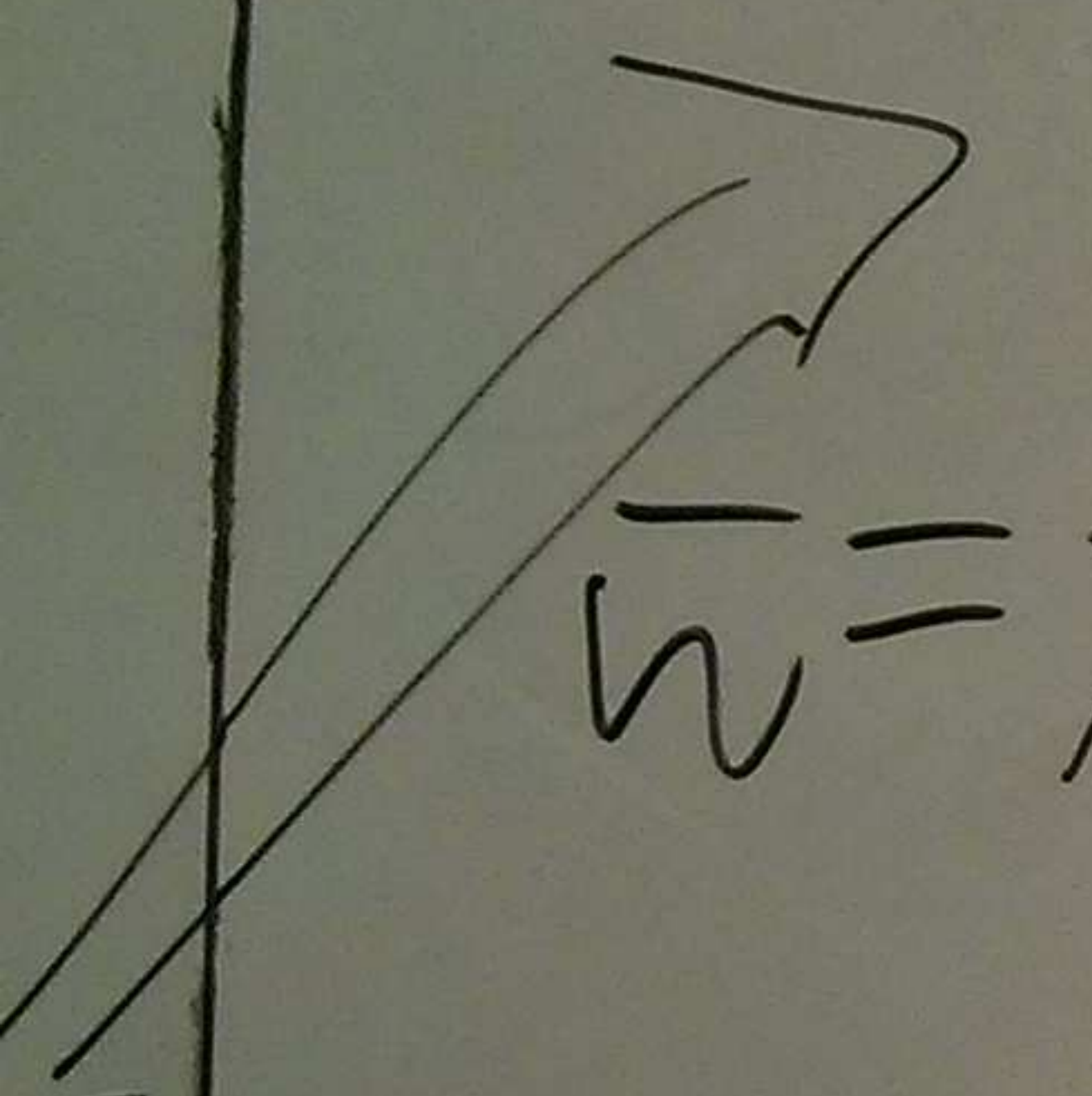
sion



# MAP estimation

MLE:  $\operatorname{argmax}_{h \in H}$

MAP:  $\operatorname{argmax}_{h \in H}$


$$\bar{w} = X^T (XX^T + \lambda I)^{-1} Y$$

$X^T X$  "scatter matrix"

$\frac{1}{N} X^T X$  "covariance MLE"

$XX^T$  "Gram matrix"

$p(D|A)$   
(posterior)

$p(h|D)$

positive  
finite"



ation

max  
H:

max  
H

$$p(D|h)$$

(posterior)

$$p(h|D)$$

$$f(x) \propto x$$

$$= cx$$

$$f(x) \propto \frac{y^2 + \alpha^3}{\gamma} x$$

(likelihood) (prior)

"Is model h reasonable"

$$\sum_x p(x) = 1$$

"exists constant, not dependent h,  
s.t.  $\sum_h p(h|D) = 1$ "



Ridge Regression  
MAP  
kernels

Ridge Regression

$$\underbrace{\frac{1}{2} \|X\bar{w} - Y\|^2}_{\text{log-likelihood}} + \underbrace{\frac{\lambda}{2} \|\bar{w}\|^2}_{\text{prior}}$$

$$\begin{aligned} -\log p(h|D) &= -\log p(D|h) - \log p(h) + \text{const.} \\ &= \sum_{i=1}^N \underbrace{\frac{1}{2} (y_i - \bar{x}_i^T \bar{w})^2}_{\text{"precision"} \quad \lambda = \frac{1}{\sigma^2}} + \underbrace{\frac{1}{2\sigma^2} \bar{w}^T \bar{w}}_{\frac{\lambda}{2} \|\bar{w}\|^2} \end{aligned}$$

$$y_i | \bar{x}_i, w_i | \theta$$

$$p(w_i) =$$



$$y_i | \bar{x}_i, \bar{w} \sim N(\bar{w}^T \bar{x}_i, 1)$$

$$w_i | \theta^2 \sim N(0, \theta^2)$$

$$p(w_i) = \frac{1}{\theta \sqrt{2\pi}} \exp\left(-\frac{1}{2} \frac{(w_i - 0)^2}{\theta^2}\right)$$

$$\log p(w_i) = -\frac{1}{2\theta^2} w_i^2 - \underbrace{\log(\theta) - \log(\sqrt{2\pi})}_{\text{const.}}$$

$$\log(ab) = \log a + \log b$$

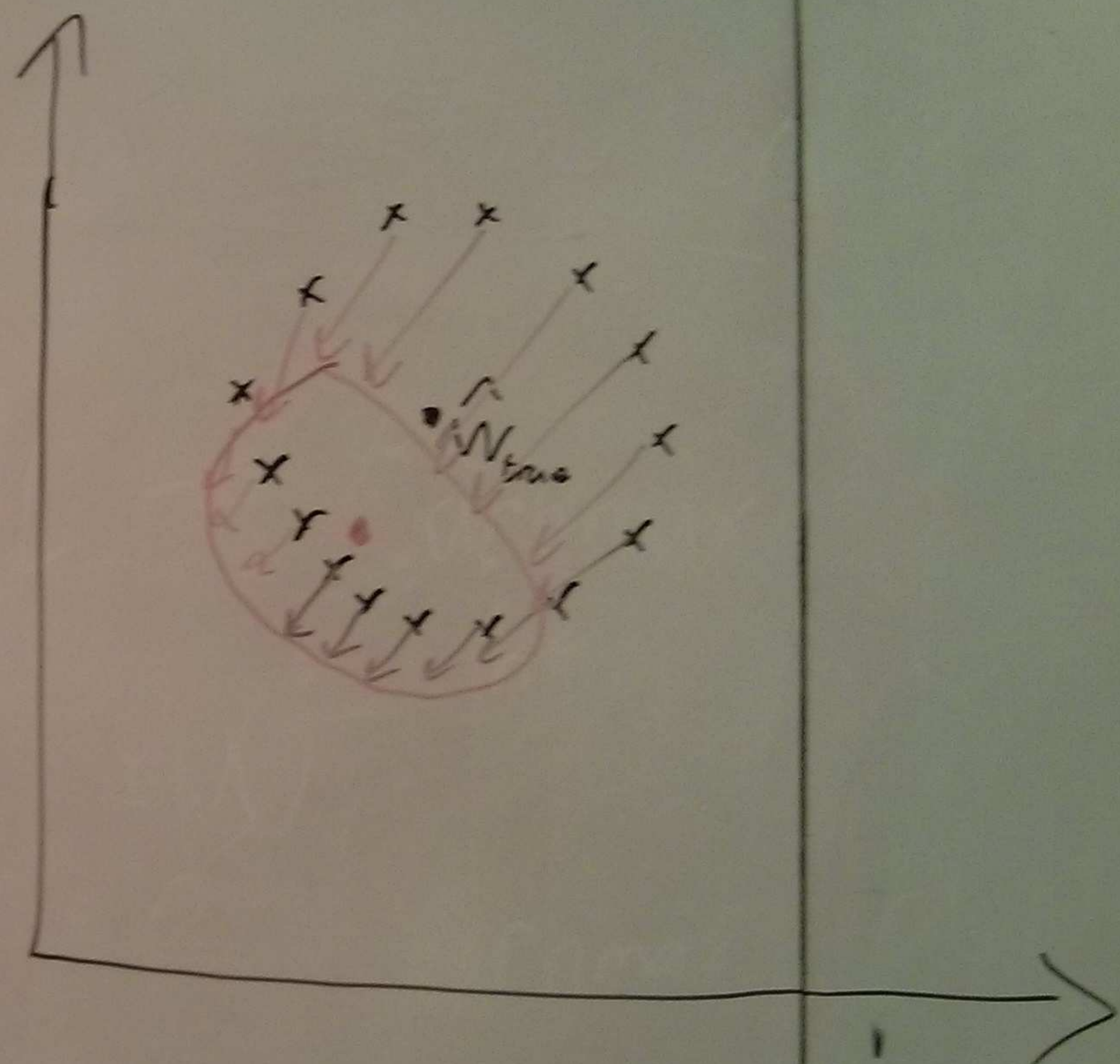
$$\log\left(\frac{a}{b}\right) = \log a - \log b$$



$$\log(ab) = \log(a) + \log(b)$$

$$\log\left(\frac{a}{b}\right) = \log(a) - \log(b)$$

(124)





Ridge Regression

MAP

kernels

Kernels

We assume  $\bar{x}_i \in \mathbb{R}^D$

What if don't know representation?

E.g. "the cat is back"

"the black cat is back"

LCS

Edit distance

"kernel

$k(\bar{x}_i, \bar{x}_j)$

$\bar{x}_i \in \mathbb{R}^D$

"the c

Typical



"kernel function"

$k(\bar{x}_i, \bar{x}_j) \mapsto \mathbb{R}$ , often a similarity.

$\bar{x}_i \in X$

"the cat is back"  $\in X$

Typically,  $k(\bar{x}_i, \bar{x}_j) \geq 0$

$$k(\bar{x}_i, \bar{x}_j) = k(\bar{x}_j, \bar{x}_i)$$

Example

vector

"Linear"

"Poly"

"RBF"



How to use?

$$\phi(\bar{x}_i) = [k(\bar{x}_i, \bar{x}_1) \quad k(\bar{x}_i, \bar{x}_2) \quad \dots \quad k(\bar{x}_i, \bar{x}_N)]$$

$$[k(\bar{x}_1, \bar{z}_1) \quad k(\bar{x}_1, \bar{z}_2) \quad \dots \quad k(\bar{x}_1, \bar{z}_m)]$$

$$\Phi(X) = \begin{bmatrix} k(\bar{x}_1, \bar{x}_1) & k(\bar{x}_1, \bar{x}_2) & \dots & \dots \\ k(\bar{x}_2, \bar{x}_1) & \dots & \dots & \dots \\ \vdots & \ddots & \ddots & \ddots \\ \vdots & \vdots & \vdots & k(\bar{x}_N, \bar{x}_N) \end{bmatrix}$$

→ Gram matrix

linear:  $\Phi(X) = XX^T$



## Examples

vectors:

$$\text{"Linear"} = \bar{x}_i^T \bar{x}_j$$

$$\text{"Poly"} = (\bar{x}_i^T \bar{x}_j + \alpha)^m$$

$$\text{"RBF"} = \exp\left(-\frac{\|\bar{x}_i - \bar{x}_j\|^2}{2\sigma^2}\right)$$

Section 14.2:

- string kernels
- pyramid match kernels



Ridge reg

MAP

kernels

Kernel Trick

Some kernels have explicit feature map  
(inner product)

$$\begin{aligned}k(x, z) &= (x^T z)^2 \\&= (x_1 z_1 + x_2 z_2)^2 \\&= x_1^2 z_1^2 + 2x_1 x_2 z_1 z_2 + x_2^2 z_2^2 \\&= (x_1^2, \sqrt{2} x_1 x_2, x_2^2)^T (z_1^2, \sqrt{2} z_1 z_2, z_2^2) \\&= \phi(x)^T \phi(z)\end{aligned}$$

$$\overline{W}_{\text{MAP}} = \lambda$$

Test time

$$\hat{y} = x^T \hat{w}$$

$$= x^T \hat{w}$$

$$= K(x, z)$$

$$= K$$



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

Test time

$$\hat{y} = \hat{X} \bar{w}_{MAP}$$

$$= \hat{X} X^T (X X^T + \lambda I)^{-1} Y$$

$$K(\hat{X}, X) \overset{\text{Gram matrix}}{K(X, X)}$$

$$= K(\hat{X}, X) (K(X, X) + \lambda I)^{-1} Y$$

"Gaussian Process"

$$f(\underline{x}) + \lambda \| \underline{w} \|^2$$



# Feature Selection

What if we don't know which  $x_i^j$  are relevant?

$$X = \begin{bmatrix} \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \end{bmatrix}$$

Search over features:

- NP-hard

- fine for 16 vars, 1 million vars

