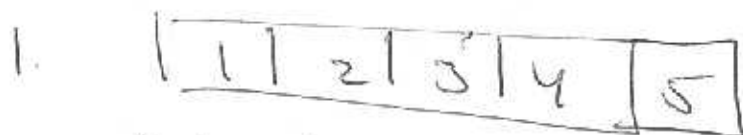


- 1a SV måste få tillgång till mycket stora datamängder så att hon i princip kan estimeras täthetsfunktioner som beskriva data. Med få exempel blir beslutsgränserna mycket osäkra eftersom samplingen av datarummet blir mycket liten.
- 1b Med endast 10 ~~stora~~ testexempel är prestandauppskattningen mycket osäker så det är ingen signifikant skillnad att göra 2 eller 3 fel. För att minska risken för överanpassning ska alltid den enklaste modellen väljas.
- 1c Bra: Bättre prestanda då alla tillgängliga exempel är med för design - säkrare beslutsgränser.
- Dåligt: Med så få exempel kommer prestanda på träningsdata att säga mycket lite om verkligt prestanda, trots att resultatet på träningsdata alltid är optimistiskt.

1d) Korsvalidering

Ex. 5-fold



Delar datamängden i 5 "icke överlappande" block.

2. Gör design m.h.a. 4 block, test med exemplar i det 5:e blocket

↓

$$\hat{e}_i \quad i=1,2,3,4,5$$

3.
$$\hat{e}_{cv} = \frac{1}{5} \sum_{i=1}^5 \hat{e}_i$$

1e) $k_t = 4 \quad N_t = 20 \quad k_t = 400 \quad N_t = 2000$

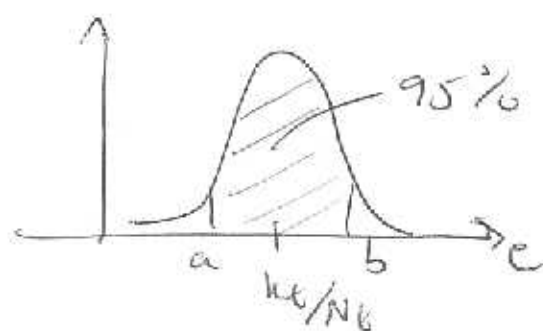
$$\hat{e} = \frac{4}{20} = \frac{400}{2000} = 20\%$$

Bayesianskt konfidensintervall

FE)

$$p(e | k_t, N_t) = \frac{P(k_t | e, N_t) P(e)}{P(k_t | N_t)} = \frac{\binom{N_t}{k_t} e^{k_t} (1-e)^{N_t-k_t}}{P(k_t | N_t)}$$

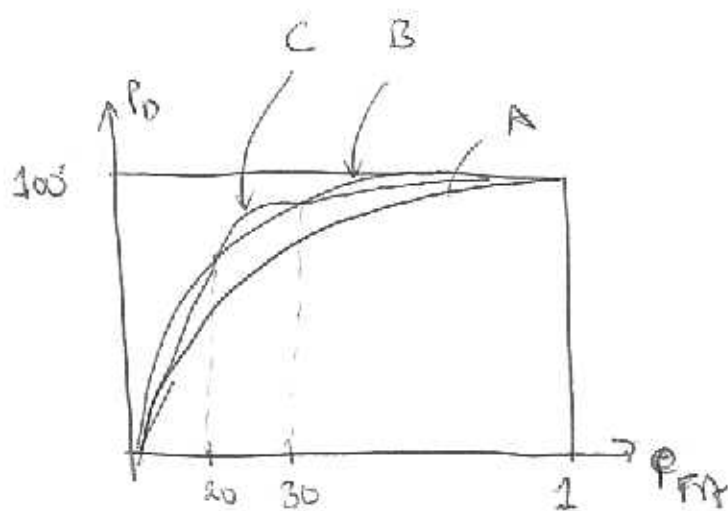
$$P(c | h_t, N_t) = \frac{e^{h_t} (1-e)^{N_t-h_t}}{\int_0^1 e^{h_t} (1-e)^{N_t-h_t} de}$$



$[a, b]$: Bayesianscher
konf. Intervall

1P)

ROC



2a) SVM har större marginal eftersom
Perceptronalgoritmen stannar så snart
alla exempel korrekt klassade

2b) Med få exempel blir det lätt
"överanpassning" av flexibla klassifikatorer

2c) Fisher: Vill projicera alla $\vec{x}_n$ på
en linje där de får värden

$$y_n = \hat{w}^T \vec{x}_n$$

strategin är att välja projektionsvektorn
 $\hat{w}$ så att

$$J(\hat{w}) = \frac{(m_1 - m_2)^2}{\sigma_1^2 + \sigma_2^2}$$

maximeras
~~minimeras~~ där $m_1 = \frac{1}{N} \sum_{y_n \in \text{Klass 1}} y_n$ osv.

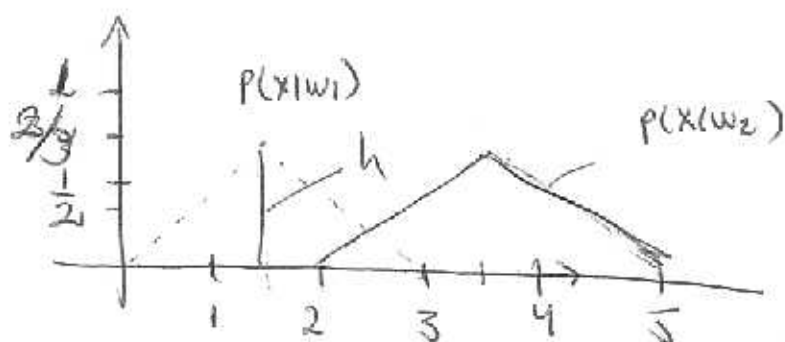
Med denna maximeras separerar
klasserna maximalt i termer av
projektionsvärdena y_n

2d)

Man kan t.ex. använda ridge regression eller PLS, som ju är designade för att hantera stora korrelationer

~~are~~

3



$$\frac{3 \cdot h}{2} = 1 \quad \boxed{h = \frac{2}{3}}$$

(a) Beslutsgräns: $\{x \mid p(x|w_1) = p(x|w_2)\}$

$$x = 2\frac{1}{2}$$

(b) Flytta beslutsgräns till $x=3$. Då blir inga w_1 -exempel felklassade

(c) I. Hållt anta/rek modell

II. Omöjligt att skatta parametrar perfekt med få data.

4a) Teor för PCA $\Rightarrow \hat{x} = \bar{m} + \sum_{i=1}^N t_i \cdot \bar{e}_i$
 ger minsta fel i minsta kvadratsmening
 (värde)

4b)

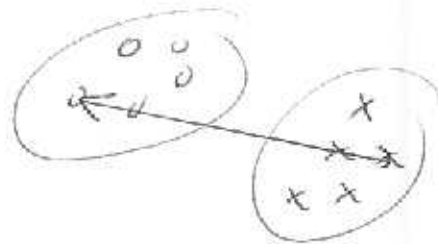
$$C : 10 \times 10$$

$\bar{x} \in \mathbb{R}^{10 \times 1} \Rightarrow$ Alla exempel ligger i
 ett underrum av dimension 2.

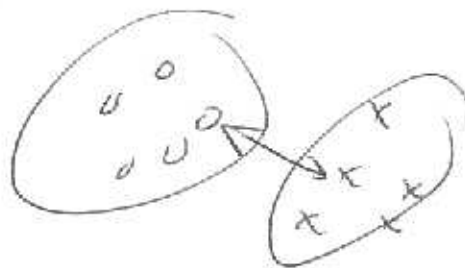
$\Rightarrow$ Alla exempel kan enkelt
 visualiseras i en 2D-plot.

4c)

Fulltast neighbor
 complete linkage



Nearest neighbor
 single linkage



4d)

1. Kluster exemplen

2. För alla exemplen i ett kluster, utför PCA.
 Om det räcker med tex 3 principalkomponenter för
 alla kluster kan alla exemplen beskrivas med ett
 heltal och tre koordinater.

