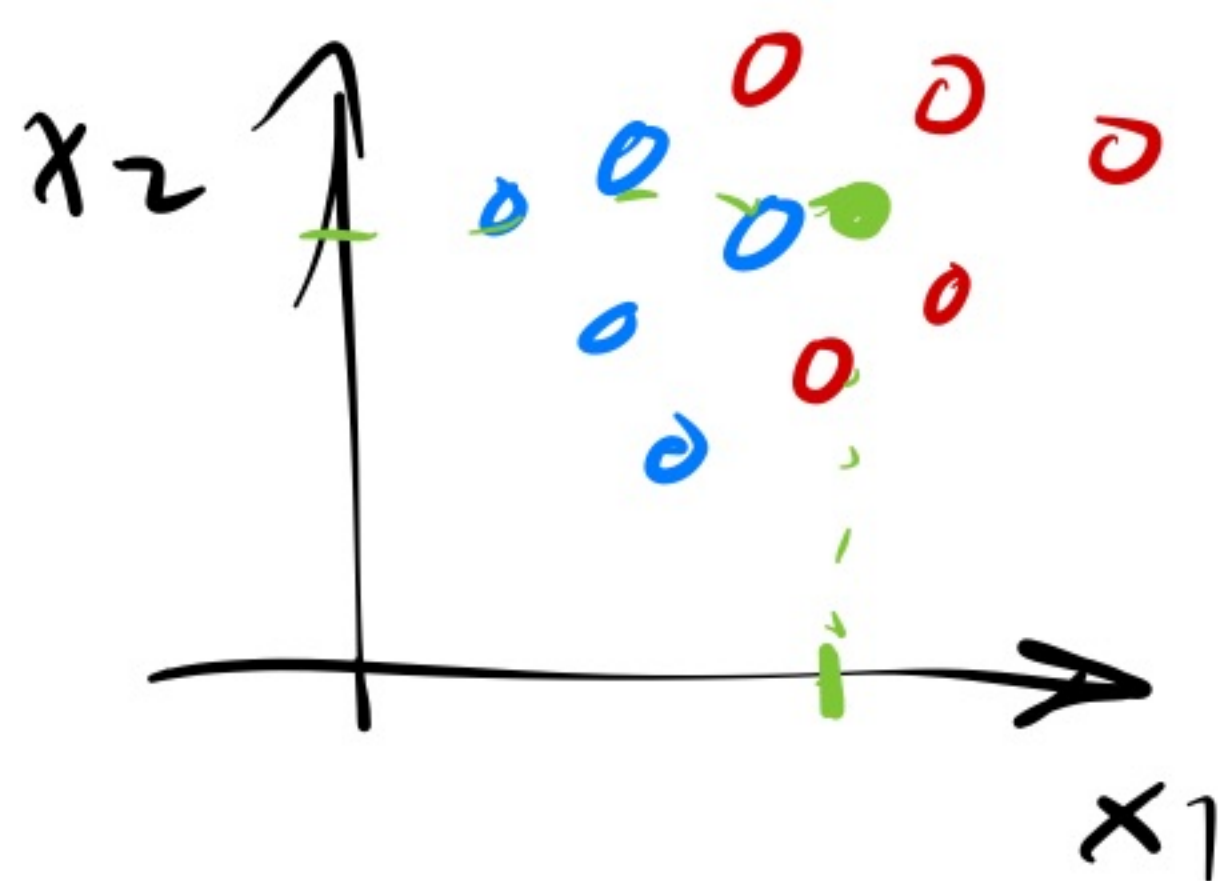


# Classification

Task



Give prediction  
for the class  
of  $\vec{x}_{\text{new}}$

3 ways to do this:

1) Discriminative functions

↳ just caring about the decision boundary

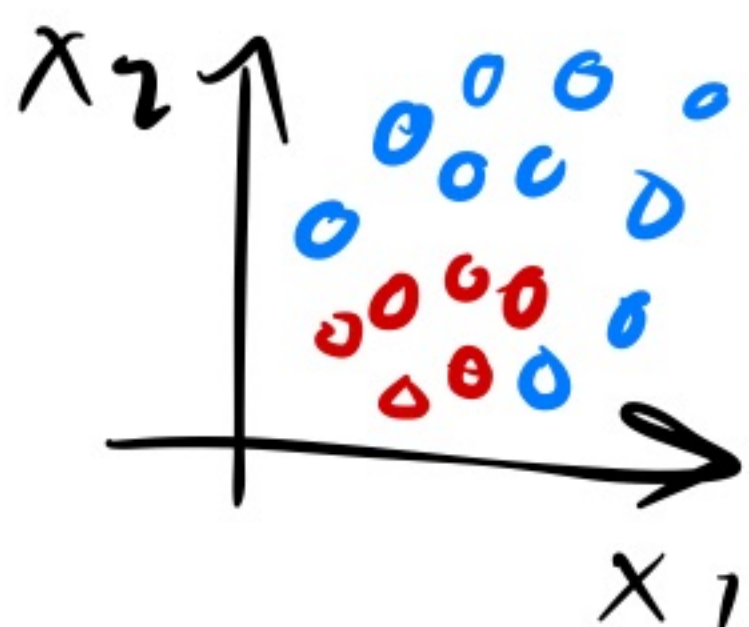
2) Probabilistic Discriminative models

3) " generative models

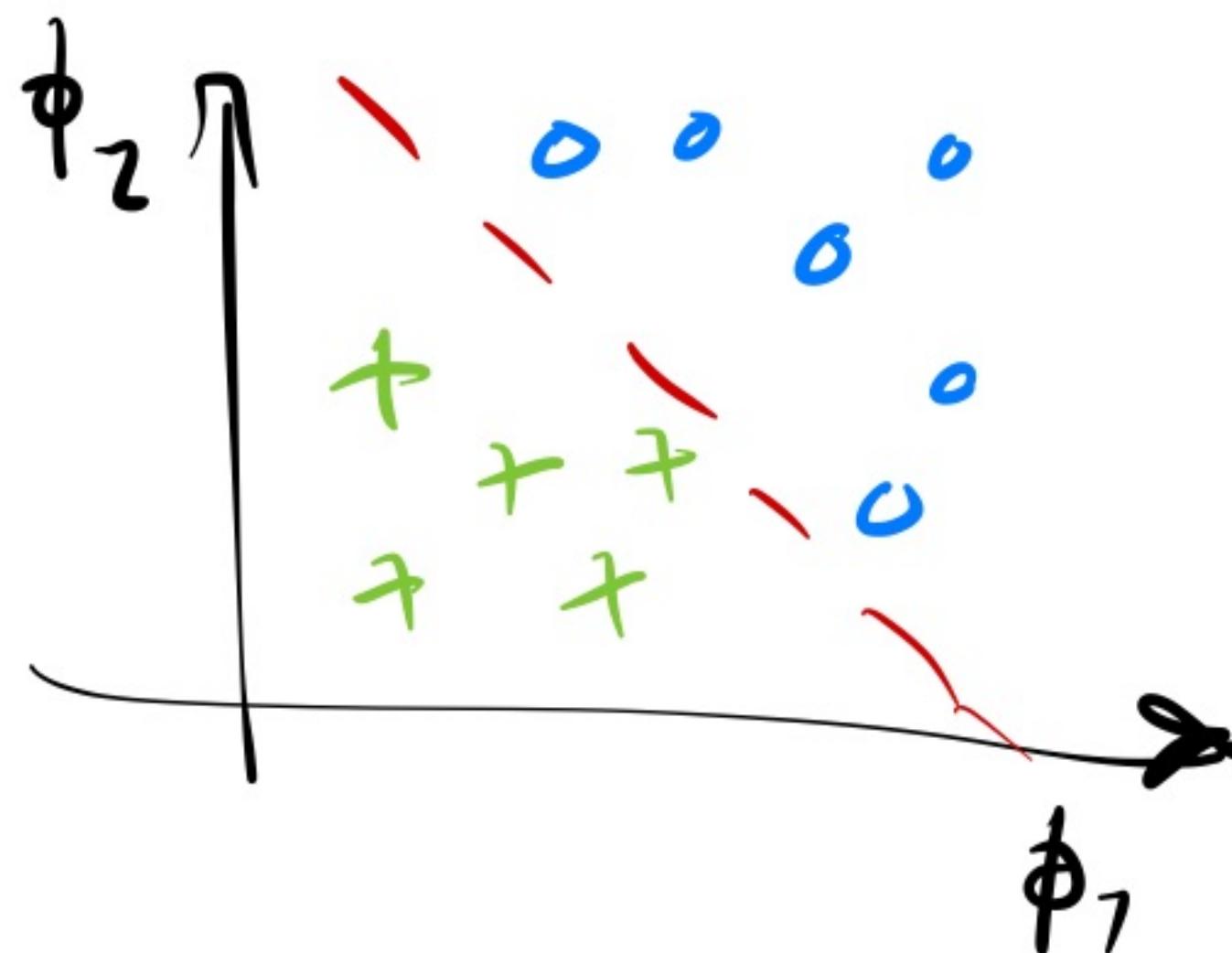
Inference on  
the prob. of  
a given class  
given some  $\vec{x}$

1) Discriminative functions

$$f(\vec{x}, \vec{w}) = \vec{w}^T \cdot \vec{\phi}(\vec{x})$$



doesn't look  
like linearly  
separable



Then decide for dm class with a rule like

$$C(\vec{x}) = \begin{cases} C_1 & \text{if } f(\vec{x}, \vec{w}) \geq 0 \\ C_2 & \text{otherwise} \end{cases}$$

If more than 2 classes

$$f_k(\vec{x}, \vec{w}_k) = \vec{w}_k^T \cdot \vec{x} + w_{k0}$$

for class  $k = 1, \dots, K$

s.t.

$$C(\vec{x}) = C_k \quad \text{if } f_k \geq f_j \quad \forall j \neq k$$

How to optimise the weights?  $\vec{w}$ ?

↙ least squares fitting?

Cost function

$$E = \frac{1}{2} \sum_{i=1}^N (\vec{t}_i - \vec{f}_i)^T (\vec{t}_i - \vec{f}_i)$$

$\vec{t}_i$  is the 1-to-K scheme for point  $i$

label e.g.  $(0, 1, 0)$  if point is class=2

and  $K=3$

$$\vec{f} = W^T \cdot \vec{\phi}(\vec{x})$$

$$K \times 1 = [(D+1) \times K] \times [(D+1) \times 1]$$

by minimizing  $E$  wrt  $W$ :

$$W^* = (\Phi^T \Phi)^{-1} \Phi^T T$$

matrix of labels  
 $N \times K$

in practice not giving  
good results in general

MSE is  
assuming that  
the output  
is Gaussian

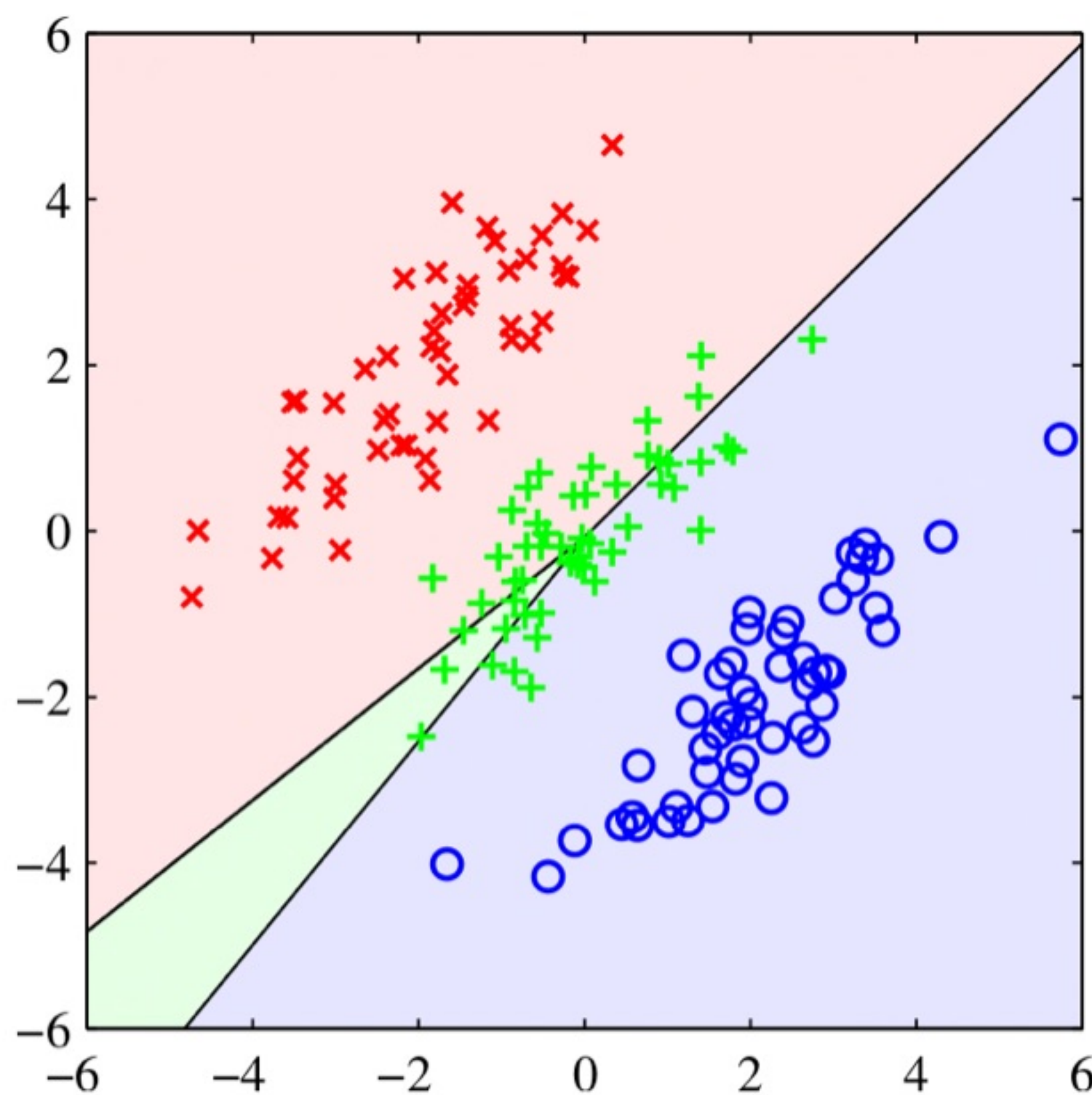


Figure 4.5 Example of a synthetic data set compri

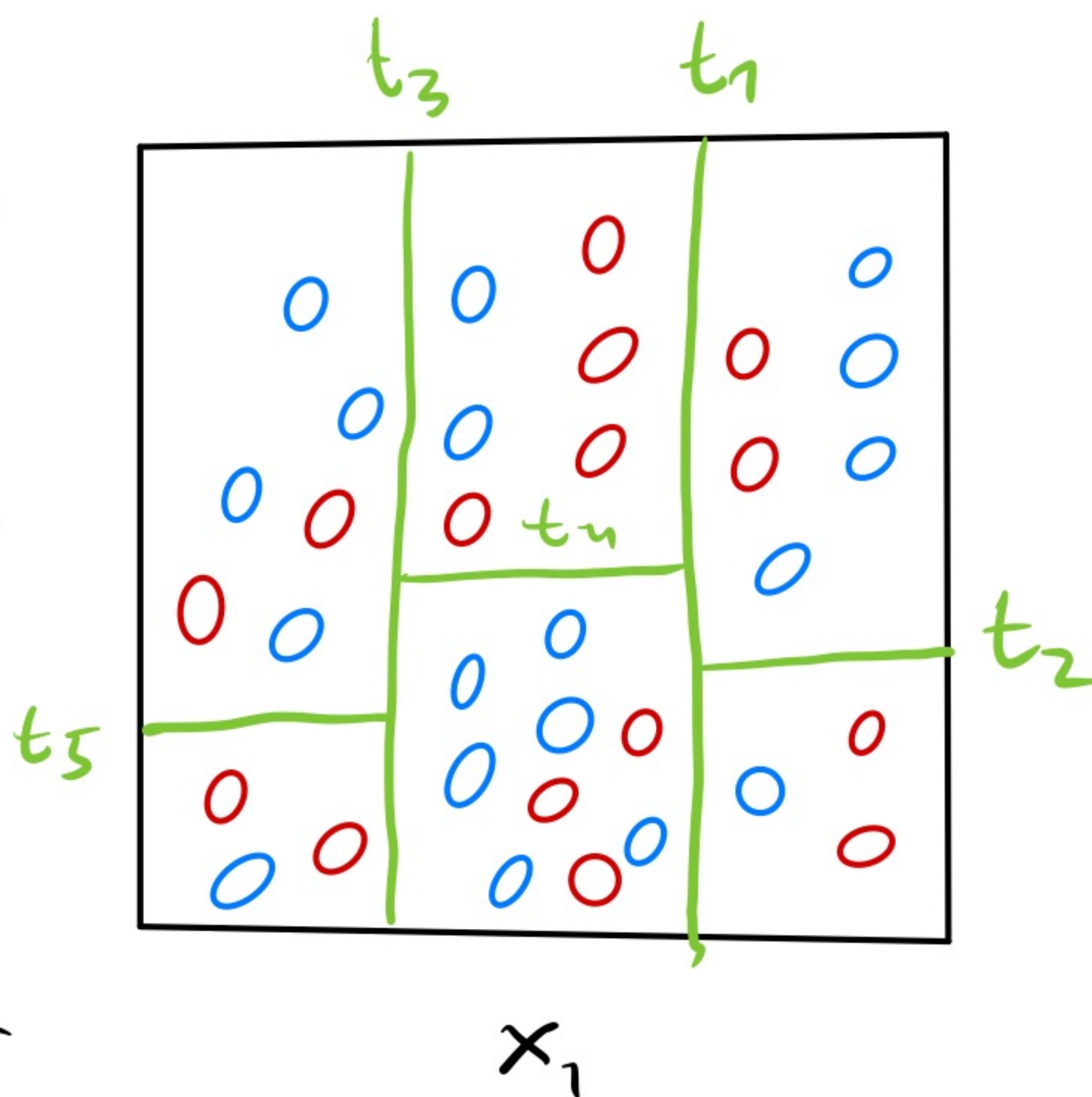
categorical variables don't follow a Gaussian

(Fisher Discriminant, Perceptron)  
better models

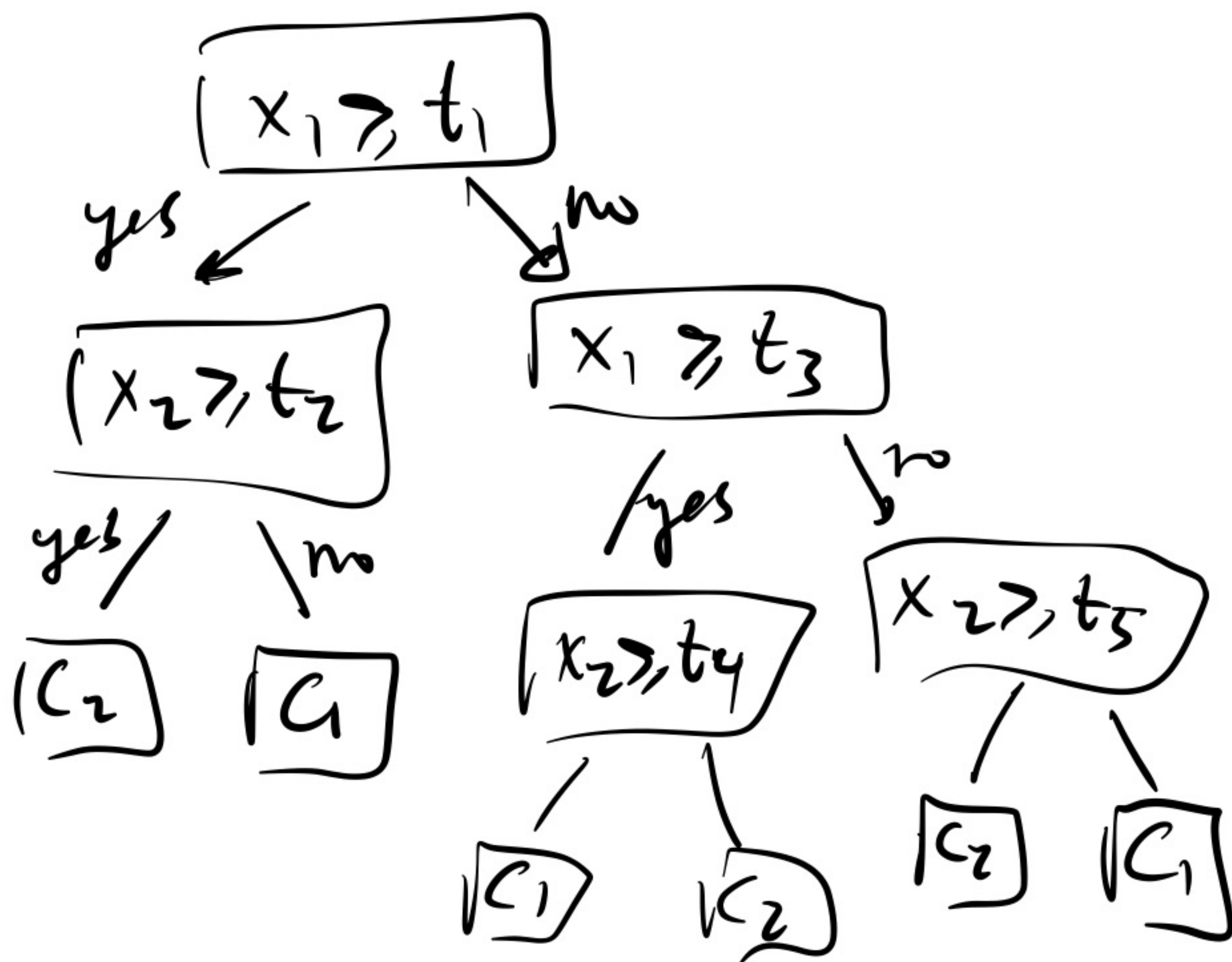
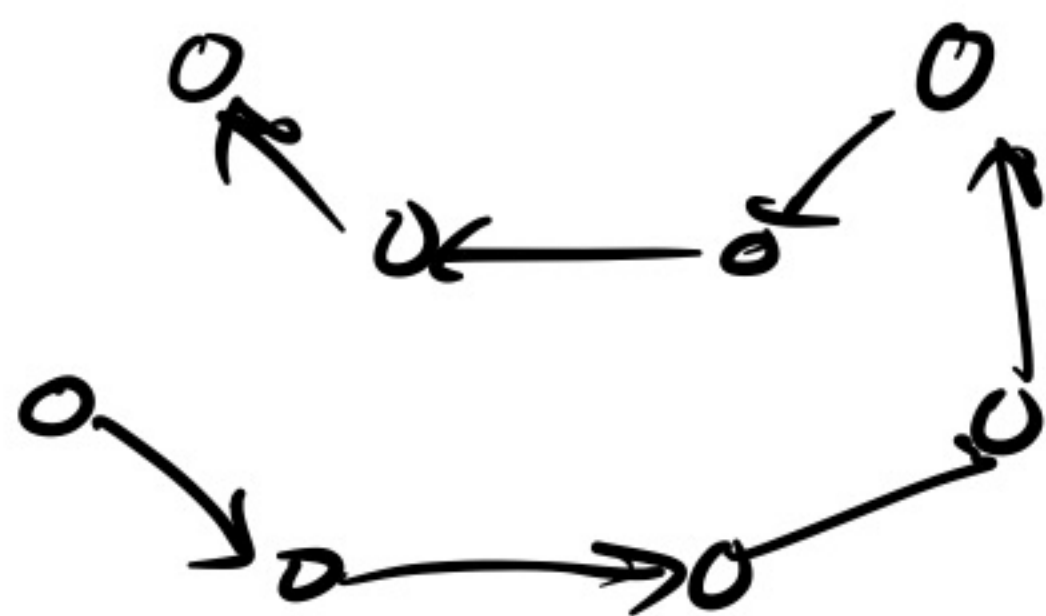
# Decision Trees

Idea: Split input space  
in a binary way  
"in a greedy  
manner"

$x_2$



e.g. Travelling salesman  
problem



The splitting criterion has to do with some measure of 'impurity'

$$P_{mk} = \frac{1}{N_m} \sum_{x_i \in R_m} I(y_i = k)$$

↓  
 proportion  
 of class  
 "k" in  
 $R_m$

Different ways to use it:

\* Misclassification error

$$1 - P_{m\hat{k}}$$

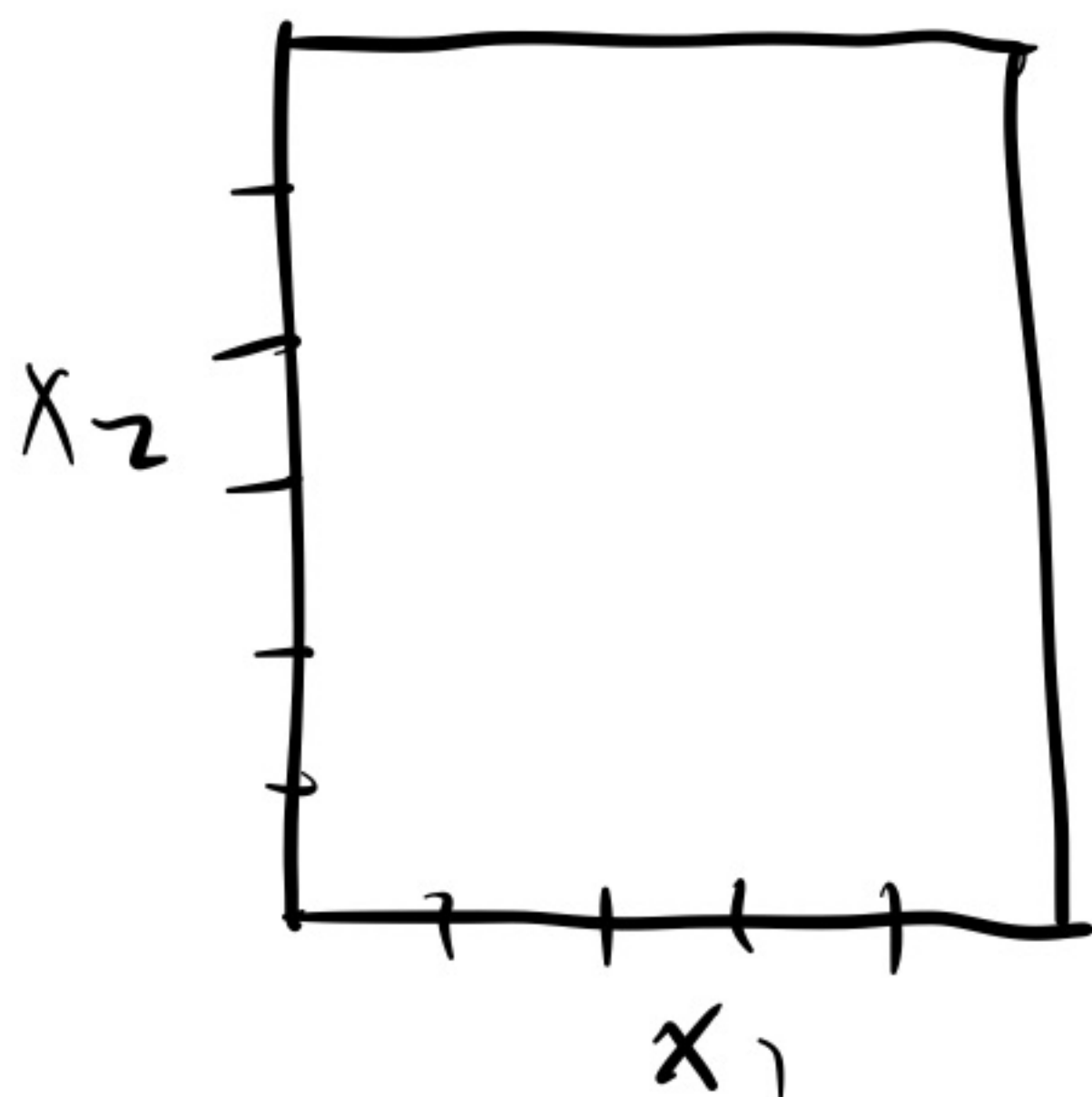
$$\hat{k} = \arg \max_k P_{mk}$$

\* Gini-index:

$$\sum_{k=1}^K P_{mk} (1 - P_{mk})$$

\* Cross-entropy:

$$-\sum_{k=1}^K P_{mk} \log P_{mk}$$



for 1st split.

evaluate Gini for all the possible points of the grid

Advantages  $\Rightarrow$  - Interpretability

- data doesn't need to be linearly separable

Disadvantages  $\Rightarrow$  makes splits are orthogonal to the axes

(hard assignments of class  
& (today this is improved by  
making probabilistic  
assignments))

- A single tree is likely to do  
 $\Downarrow$  overfitting

solution is to rely  
= ensemble methods

(Decision Forests e.g.)

Situations where Discriminative  
Functions may not be convenient

1) inherently interested in prob. of  
given class

2) want to purchase <sup>more</sup> misclassifications in class A than in class B

⇓  
design a particular loss function

$$L: \text{true} \begin{matrix} A & B \end{matrix} \begin{matrix} A \\ B \end{matrix} \begin{pmatrix} 0 & 1000 \\ 1 & 0 \end{pmatrix}$$

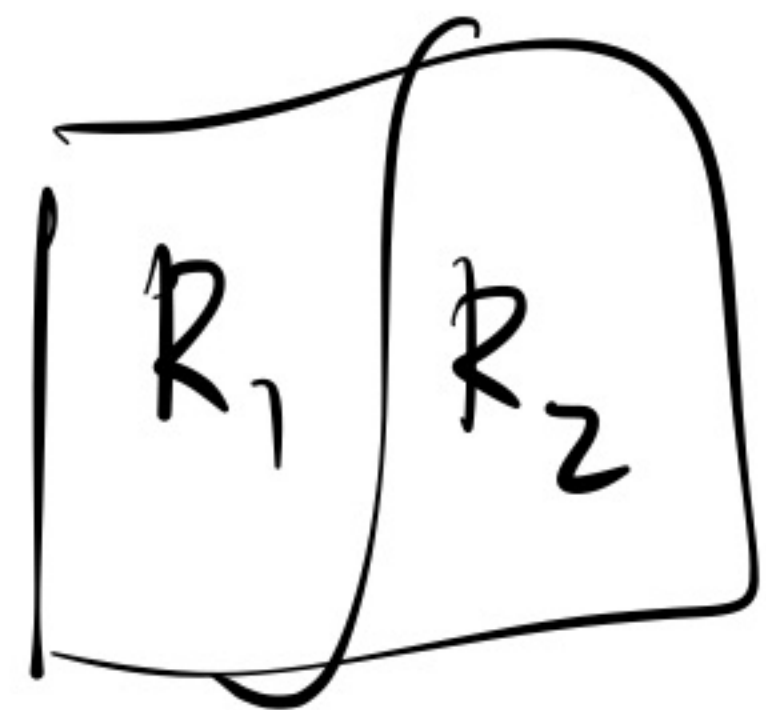
predicted class

Minimising

$$\mathbb{E}[L] = \sum_k \sum_j \int_{R_j} L_{kj} P(\vec{x}, c_k) d\vec{x}$$

look for

$$R_j = \operatorname{argmin} \mathbb{E}[L]$$



⇓

$$\hat{j} = \operatorname{argmin} \sum_k L_{kj} \underbrace{P(c_k | \vec{x})}_{\text{object you are after}}$$

you don't need to train again

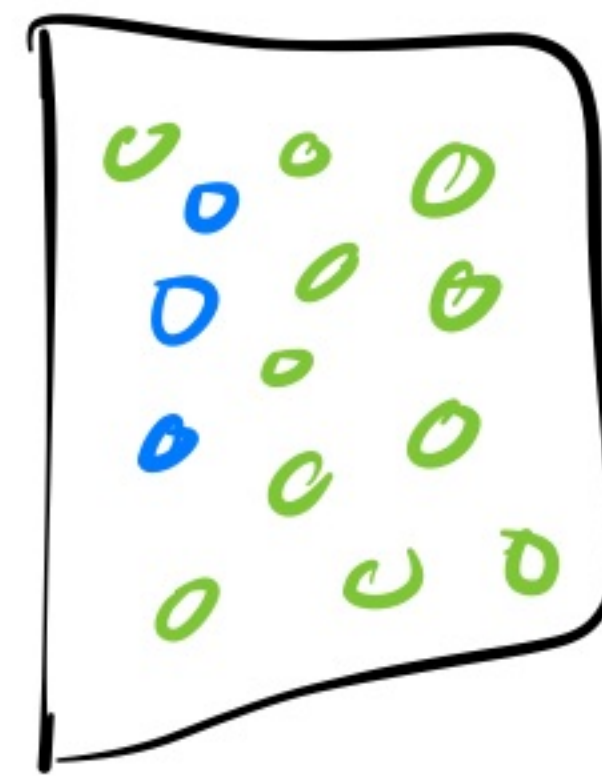
if  $L$  changes

training is

- Compensate for unbalanced data

Assign priors

$p(C_k) \propto$  Fraction of points of each  $k$



99%

1%

- Balance the data

- estimate  $p(C_k | \vec{x}) \rightarrow \alpha$

- correct with  $p(C_k)$

## Probabilistic Approach

Random binary variable  $t$

$$\text{Bern}(t | \mu) = \mu^t (1 - \mu)^{1-t}$$

↳ prob for  $t=1$

Categorical variable  $\vec{t} = \{t_k\}$

$$g_{\text{Bern}}(\vec{t} | \vec{\mu}) = \prod_{k=1}^K \mu_k^{t_k}$$

Idea is to infer  $\mu_k$

estimate  $p(C_k | \vec{x}) \Leftarrow \mu_k$

Suppose dataset  $D = \{ \vec{x}_i, \vec{t}_i \}_{i=1}^N$   
likelihood

$$P(D | \vec{w}) = \prod_{i=1}^N \prod_{k=1}^K p(C_k | \vec{x}_i; \vec{w})^{t_{ik}}$$

Classif. model is a proposition for

Prob. Generative approach

$$p(C_k | \vec{x}) \Leftarrow p(\vec{x} | C_k)$$

$\Downarrow$  via the Bayes theorem, given  $C_k$   
likelihood of point  $\vec{x}$

$$p(C_k | \vec{x}) = \frac{p(\vec{x} | C_k) p(C_k)}{p(\vec{x})}$$

Consider 2 classes:

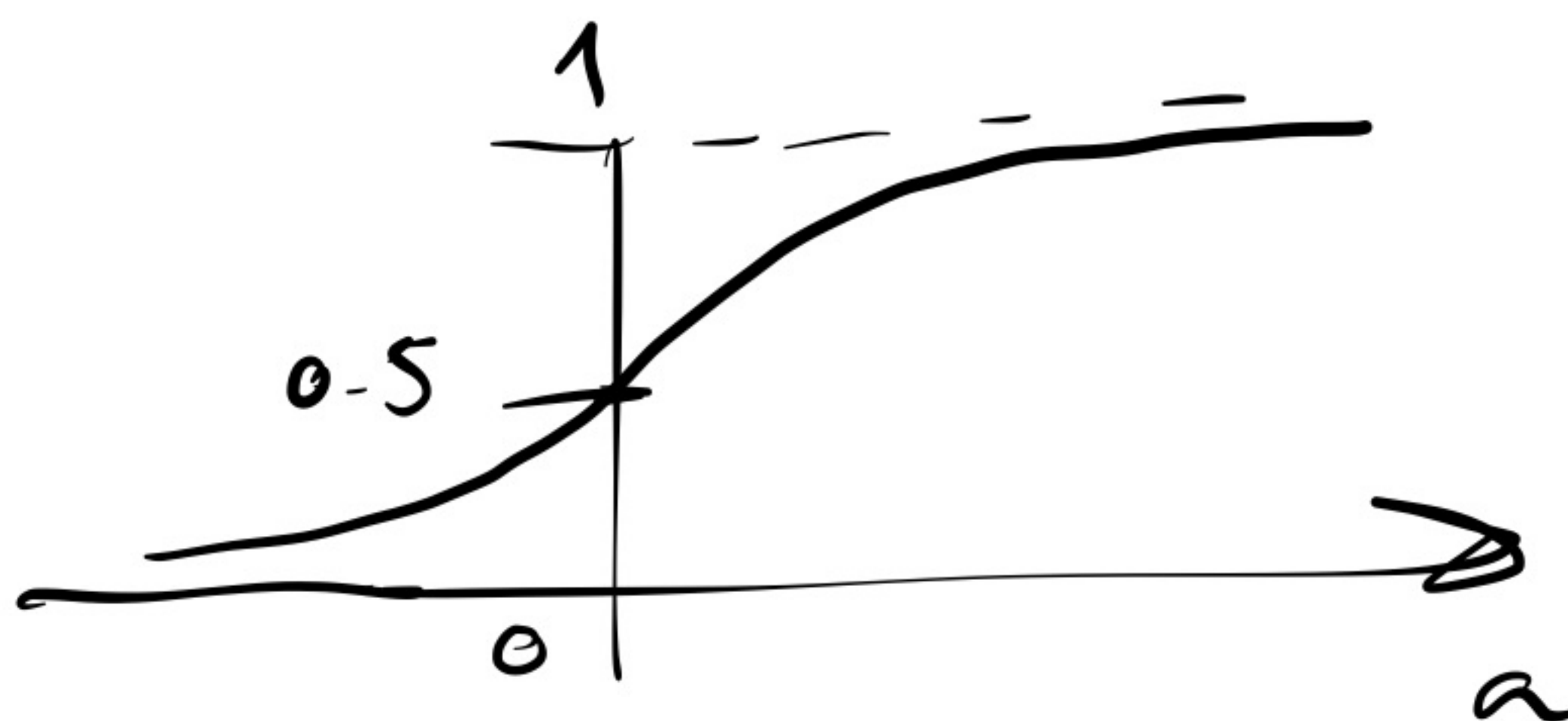
$$p(C_1 | x) = \frac{p(x | C_1) p(C_1)}{p(x | C_1) p(C_1) + p(x | C_2) p(C_2)}$$

$$z = \frac{1}{1 + \frac{p(x|C_2)p(C_2)}{p(x|C_1)p(C_1)}}$$

$$a \equiv \ln \left[ \frac{p(x|C_1)p(C_1)}{p(x|C_2)p(C_2)} \right]$$

$$z = \frac{1}{1 + e^{-a}} = \sigma(a)$$

$\Downarrow$   
 sigmoid function



For  $k > 2$  classes

$$p(C_k|x) = \frac{p(x|C_k)p(C_k)}{\sum_j p(x|C_j)p(C_j)}$$

$$\sum_j p(x|C_j)p(C_j)$$

$$a_k \equiv \ln p(x|C_k)p(C_k)$$

$$\rightarrow = \frac{e^{a_k}}{\sum_j e^{a_j}} \equiv \sigma_k(\vec{a})$$

$\Uparrow$   
 softmax function

e.g. suppose  $\vec{x}$  is Gaussian

$$p(\vec{x} | C_k) \propto \exp \left\{ -\frac{1}{2} (\vec{x} - \vec{\mu}_k)^T \Sigma^{-1} (\vec{x} - \vec{\mu}_k) \right\}$$

$\swarrow$  characterises the class  $k$   
 $\searrow$  assume the same  $\forall k$

e.g. 2 classes

$$a = \ln p(\vec{x} | C_1) - \ln p(\vec{x} | C_2) + \ln [p(C_1) / p(C_2)]$$

$$= (\vec{\mu}_1 - \vec{\mu}_2)^T \Sigma^{-1} \vec{x} - \frac{1}{2} \vec{\mu}_1^T \Sigma^{-1} \vec{\mu}_1 + \frac{1}{2} \vec{\mu}_2^T \Sigma^{-1} \vec{\mu}_2 + \ln p(C_1) / p(C_2)$$

$$= \vec{w}^T \vec{x} + w_0$$

$$P(C_k | \vec{x}) = \sigma(a_k) \\ = \sigma(\vec{w}_k^T \vec{x} + w_{k0})$$

↑  
implies a linear  
decision boundary

Now you need to  
find the optimum  $\vec{w}$

Example of a generative model:

## Naive Bayes classifier

Assume that, given a class  $C_k$ ,  
each component of  $\vec{x} = \{x_d\}_{d=1}^D$   
are indep. from each other

$$P(\vec{x} | C_k) = \prod_{d=1}^D p(x_d | C_k)$$

↳ 'conditional independence  
assumption'

$$P(C_k | \vec{x}) = \left[ \prod_{d=1}^D p(x_d | C_k) \right] p(C_k)$$

↑↑

compute this for each  $k$

for a new point  $\vec{x}_{\text{new}}$  we choose  
the class that maximises  $p(C_k | \vec{x}_{\text{new}})$

$$C(\vec{x}_{\text{new}}) = \underset{k}{\operatorname{argmax}} p(C_k | \vec{x}_{\text{new}})$$

## Probabilistic Discriminative Models

directly inferring  $p(C_k | x)$

without defining  $p(x | C_k)$  first

## logistic Regression

$$p(C_k | \vec{x}) = \sigma(\underbrace{\vec{w}_k \cdot \vec{x} + w_0^k}_{\text{linear combination}})$$

it is a linear classifier

$$p(C_k | \vec{x}) = \sigma(\vec{w}_k^T \cdot \vec{\phi}(\vec{x}))$$

To train it: minimize cross-entropy:

$$C = -\log p(T|W)$$

↓  
matrix of labels

$$= -\sum_{i=1}^N \sum_{k=1}^K t_{ik} \log \sigma(\vec{w}_k^T \vec{\phi}_i(\vec{x}))$$

$$\frac{dC}{d\vec{w}_j} = -\sum_i \sum_k t_{ik} \cdot \frac{1}{y_{ik}} \frac{dy_{ik}}{da_j} \frac{da_j}{d\vec{w}_j}$$

$$= \sum_i (y_{ij} - t_{ij}) \vec{\phi}_i$$

\* please do it  
HW

$$= \sum_i [\sigma(\vec{w}_j^T \cdot \vec{\phi}_i) - t_{ij}] \vec{\phi}_i = 0$$

optimum  $\vec{w}_j^*$  will not be analytical

$$\vec{w}_j^* = \underset{\vec{w}_j}{\operatorname{argmin}} C(\vec{w}_j) \quad \text{numerically}$$

Log. Regression  $\Leftarrow$  Neural network  
 $\nwarrow$   
 without hidden layer