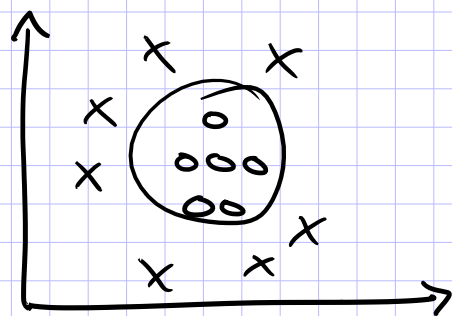
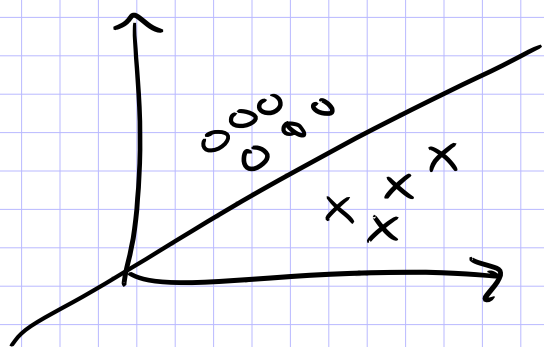


KNN Decision Rule

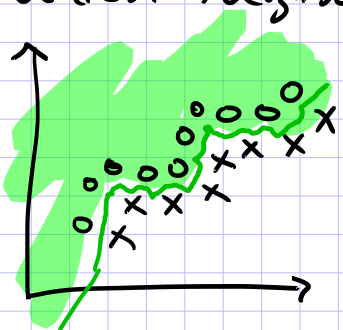


Depends:

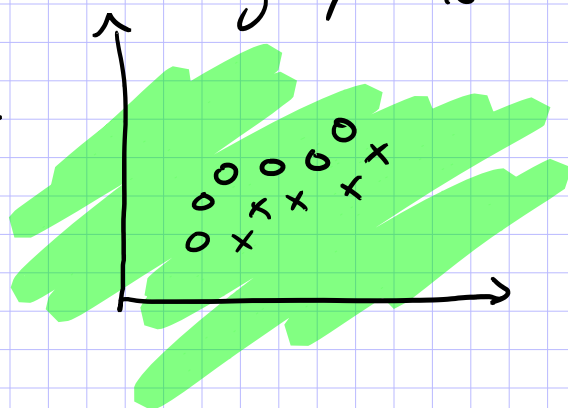
- ① our dataset D
- ② our distance d
- ③ our choice of k

Choosing k

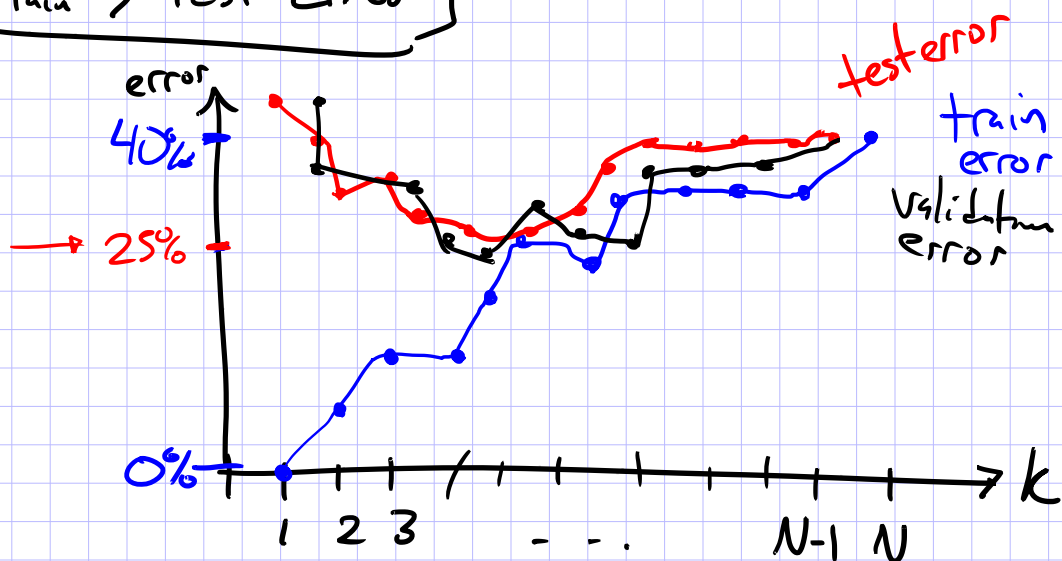
Special Case: $k=1$
"Nearest Neighbor"



Special Case: $k=N$
"Majority Vote"



Train $\frac{1}{2}$ Test Error



D is 40% $y^{(i)}=0$
60% $y^{(i)}=1$

$$\begin{aligned}
 \hat{y}^{(test)} &= \begin{bmatrix} h(x^{(t1)}) \\ \vdots \\ h(x^{(tT)}) \end{bmatrix} \\
 D^{train} &= \begin{bmatrix} y^{(1)} \\ \vdots \\ y^{(N)} \end{bmatrix} \begin{bmatrix} x^{(1)} \\ \vdots \\ x^{(N)} \end{bmatrix} \quad \begin{matrix} 80\% \text{ train} \\ 10\% \text{ validation} \end{matrix} \\
 D^{test} &= \begin{bmatrix} y^{(t1)} \\ \vdots \\ y^{(tT)} \end{bmatrix} \begin{bmatrix} x^{(t1)} \\ \vdots \\ x^{(tT)} \end{bmatrix} \quad \begin{matrix} 10\% \text{ test} \end{matrix} \\
 &\quad \left[\begin{matrix} 100\% \\ N+T \end{matrix} \right]
 \end{aligned}$$

★ Choose k based on validation error

Compare for test error

Lecture 3

Function Approx. View

- ① There exists some unknown distribution p^* that generates our unlabeled data instances, $x^{(i)}$

$$x^{(i)} \sim p^*(x) \quad \forall i$$

↑ denotes "is sampled from"

- ② Human expert annotated each training instance in D using a fixed unknown function h^*

$$y^{(i)} = h^*(x^{(i)}) \quad \forall i$$

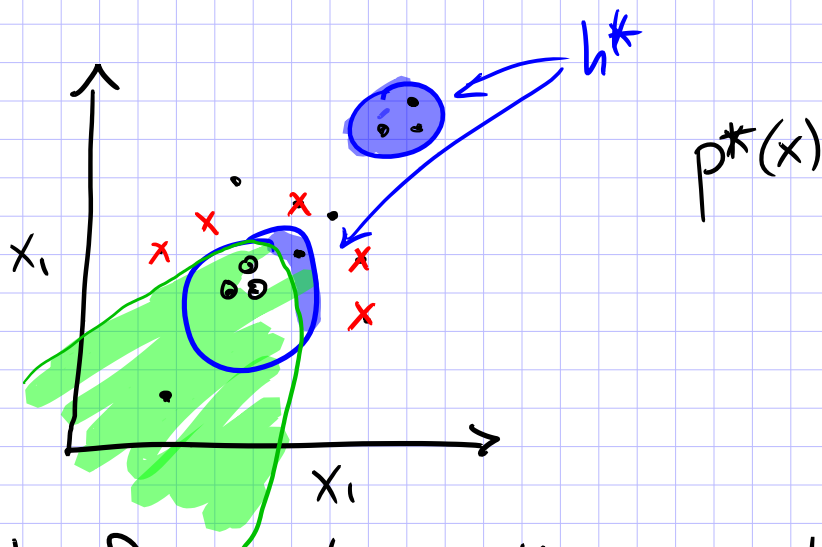
- ③ Learning algo. takes data D and output a (good) hypothesis $h \in H$

$$y^{(ti)} = h(\vec{x}^{(ti)})$$

- ④ Our goal: choose h with low error on data from p^*

$$\text{error}(h, p^*) = \mathbb{P}_{x \sim p^*}(h(x) \neq h^*(x))$$

↑ we can't compute this but validation error gives a good approximation



Question: What if we know that our data was not created as described above?