

Non-parametric Methods

$$D = \{x_i, y_i\}$$

e.g. regression

1) Assume PDF for y_i

$$y(x) \sim N(y \mid \mu(\vec{x}, \vec{w}), \sigma^2)$$

2) Take a model for $\mu(\vec{x}, \vec{w})$

(e.g. neural net,
basis functions)

3) Fit $\vec{w}$

parametric methods

→ distribution is fixed

- its shape

- the mean function

Non-param methods

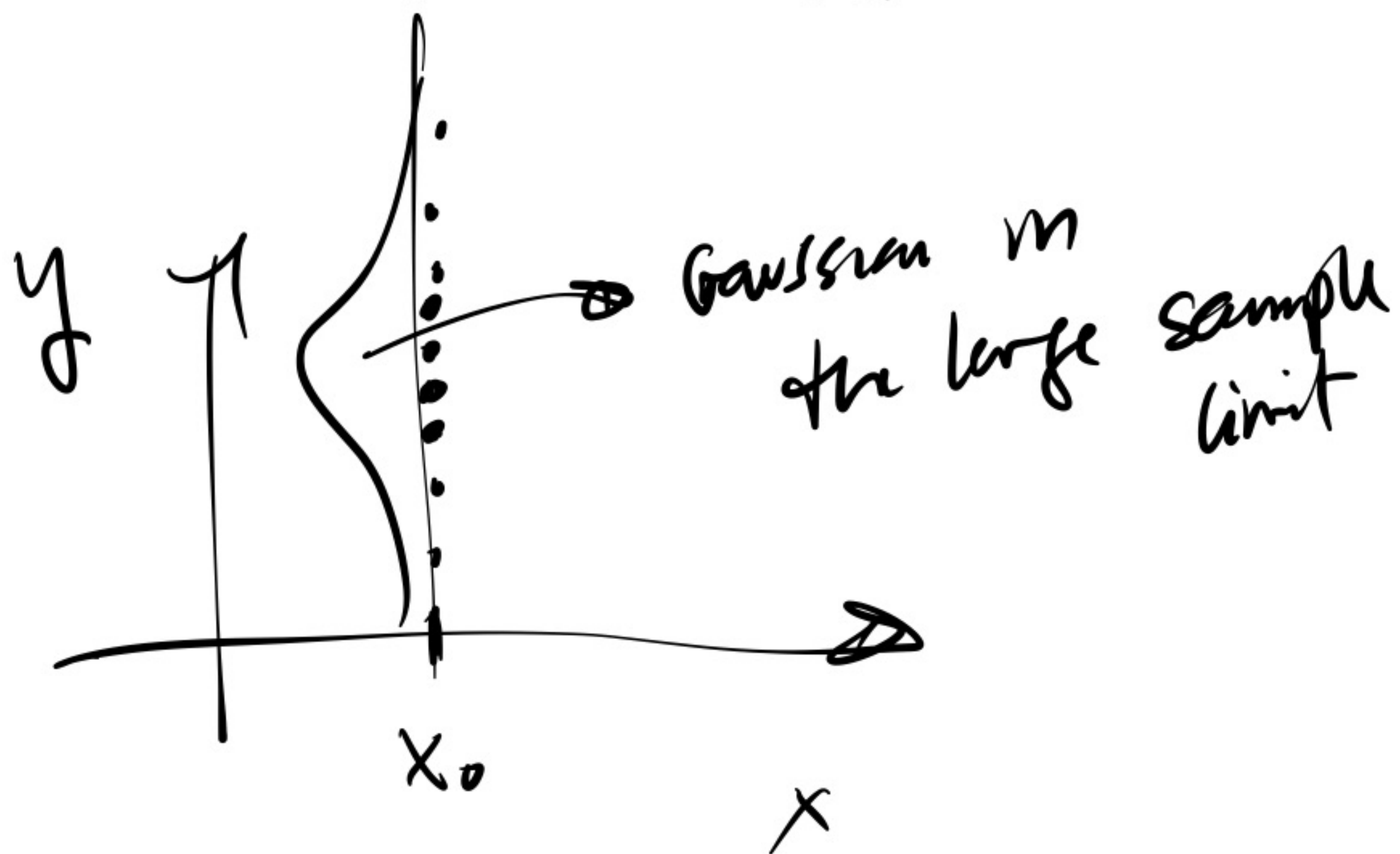
today

→ distribution is not fixed
a priori

Gaussian
Processes

→ mean function is not fixed
a priori

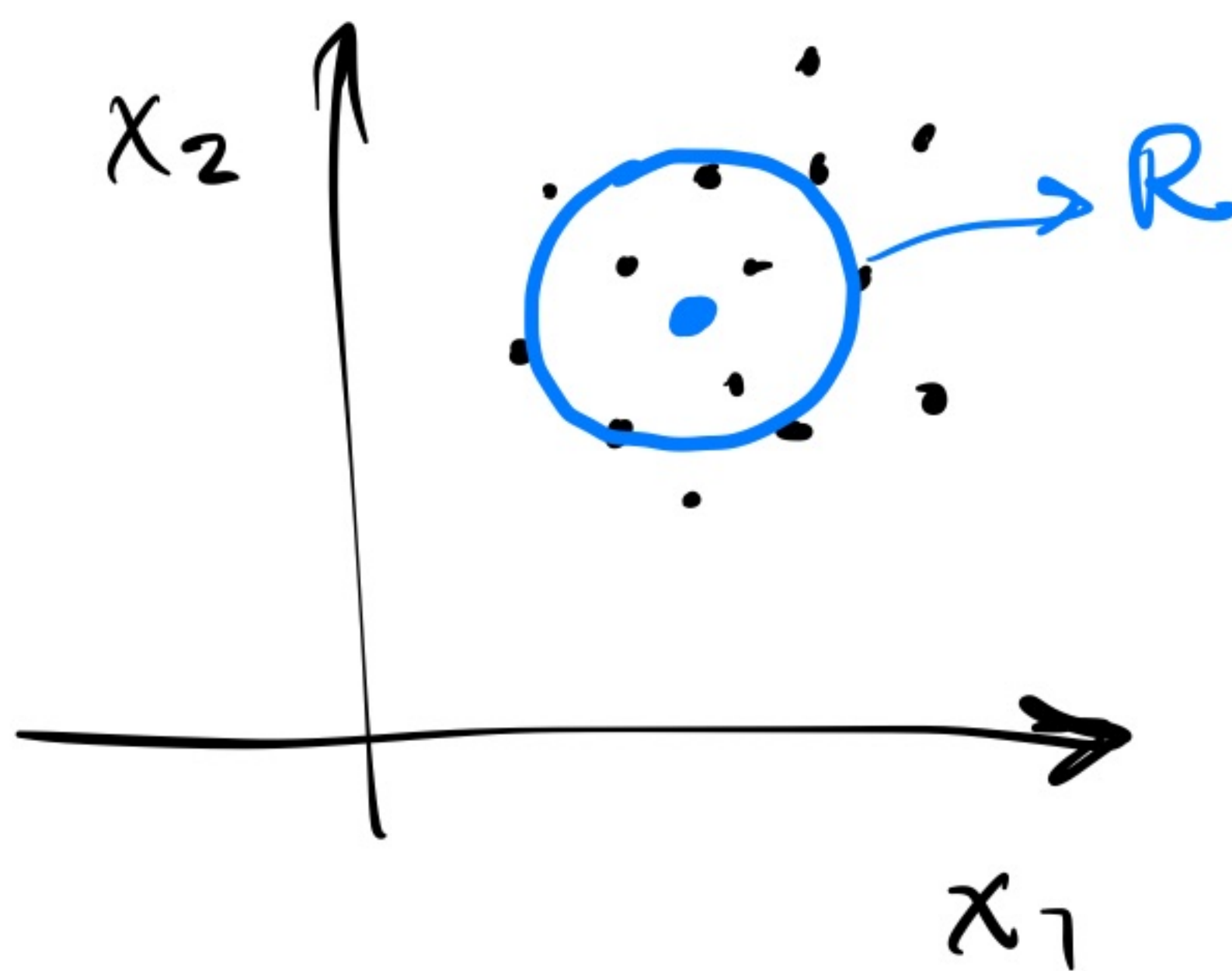
The distn is fixed by data
 (if dataset is large N) the distn.
 will be more flexible



$$D = \{ \vec{x}_i \}_{i=1}^N$$

estimate the
 PDF of $\vec{x}$

$$p(\vec{x})$$



Prob. mass associated to R

$$P = \int_R p(\vec{x}) d\vec{x} \quad \left\{ \begin{array}{l} \text{prob of each} \\ \text{point to fall in } R \end{array} \right.$$

The prob for k (out of N) points

$$\text{Bin}(K|N, P) = \frac{N!}{K!(N-K)!} P^K (1-P)^{N-K}$$

$$\mathbb{E}[K] = NP$$

$$\mathbb{E}[K/N] = P \quad (\text{I})$$

$$\text{Var}[K] = NP(1-P)$$

variance of K/N ?

$$\text{Var}[K/N] = \mathbb{E}\left[\left(\frac{K}{N}\right)^2\right] - \left(\mathbb{E}\left[\frac{K}{N}\right]\right)^2$$

$$= \frac{1}{N^2} \mathbb{E}[K^2] - P^2$$

$$= \frac{1}{N^2} \left(\underbrace{\text{Var}[K]}_{NP(1-P)} + \underbrace{(\mathbb{E}[K])^2}_{N^2 P^2} \right) - P^2$$

$$= \frac{P(1-P)}{N} \left\{ \begin{array}{l} \text{decreases with } N \end{array} \right.$$

large N distn peaked at the mean

- Assume that R is small enough

$$(II) \quad P \approx \underbrace{p(\vec{x}')}_{\text{roughly const in } R} V$$

$$\frac{k}{N} \approx p(\vec{x}') V \Rightarrow \boxed{p(\vec{x}') \approx \frac{k}{NV}}$$

Two possibilities to proceed:

1) Fix k and determine V from data

↳ k -nearest neighbors

2) Fix V and determine k from data

↳ (kernel methods)

k -nearest neighbors

k fixed $\Rightarrow$ determines the model complexity

$\vec{x}'$

eg. classification M classes $\{C_m\}$

dataset N points $N = \sum_m N_m$

Task is to classify $\vec{x}$

↳ consider k nearest neighbors to $\vec{x}$

K_m = # of points (out of k)
belonging to C_m

$$p(\vec{x} | C_m) = \frac{K_m}{N_m \cdot V}, \quad p(\vec{x}) = \frac{k}{N \cdot V}$$

Take class priors $p(C_m) = N_m / N$

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

$$= \frac{K_m}{\cancel{N_m \cdot V}} \cdot \frac{\cancel{N_m}}{\cancel{N}} \cdot \frac{\cancel{N \cdot V}}{k} = \frac{K_m}{k}$$

The best class is the one where K_m/k
is maximal

↳ majority vote

For regression :

Consider both $\vec{x}$ and y to be random

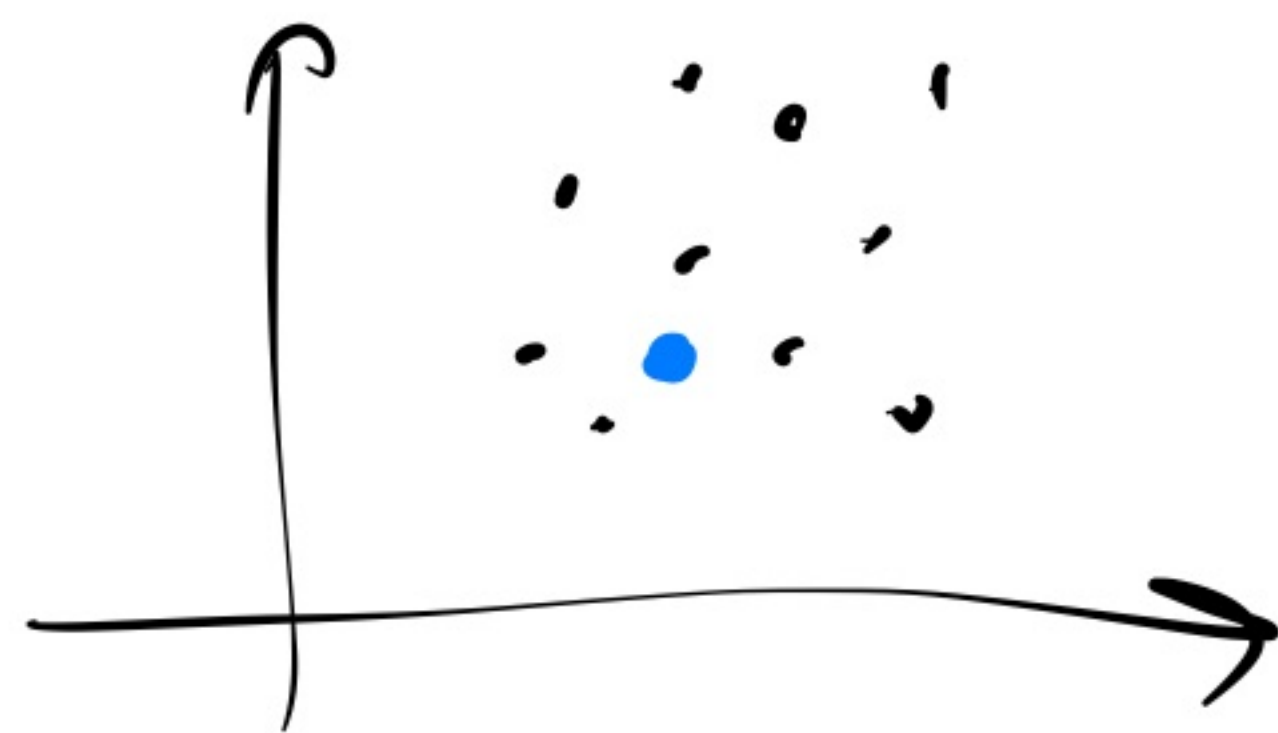
$$\begin{aligned} \langle y(\vec{x}) \rangle &= \frac{\int dy \cdot y \cdot p(\vec{x}, y)}{\int dy \cdot p(\vec{x}, y)} = \\ &\downarrow \\ \text{"regression function"} &= \frac{\int dy \cdot y \cdot p(y|\vec{x}) \cdot \cancel{p(\vec{x})}}{\cancel{p(\vec{x})}} \\ &= \int dy \cdot y \cdot p(y|\vec{x}) = \mathbb{E}[y|\vec{x}] \end{aligned}$$

→ approximated in two ways

- $\mathbb{E} \approx$ average over training sample
- conditioning at a point
↳ conditioning at a region close to the point

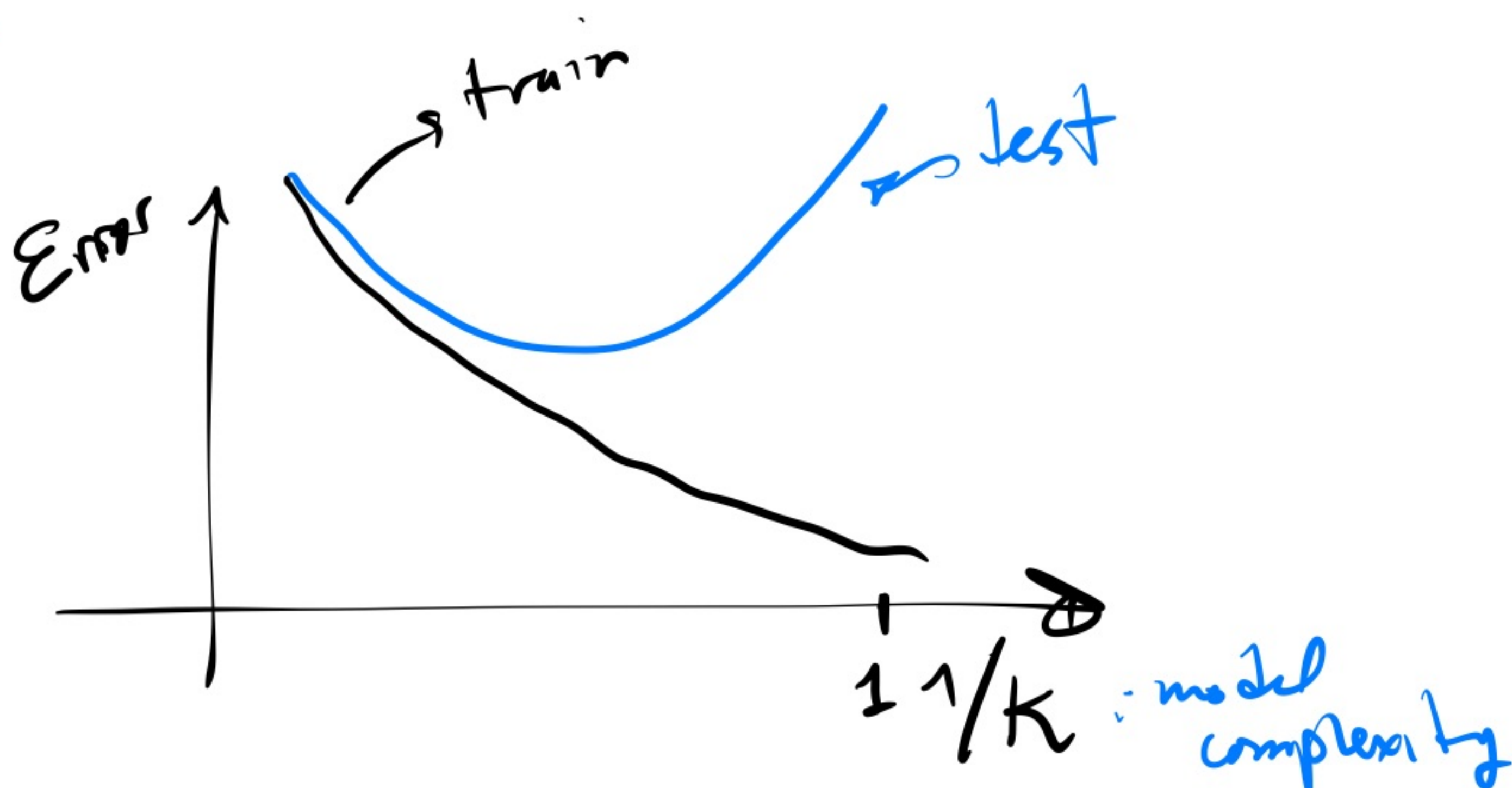
$$\langle y(\vec{x}) \rangle \approx \frac{1}{k} \sum_{i \in V(\vec{x})} y_i \quad \left\{ \begin{array}{l} \text{average output} \\ \text{of all the} \\ k \text{ neighbors} \end{array} \right.$$

$K=N \Rightarrow$ simplest model
 $\downarrow$ high bias
 $\downarrow$ low variance
 $\downarrow$ average of all the points



$K=1 \Rightarrow$ more complex model

$\downarrow$ overfit
 $\downarrow$ large variance (and low bias)

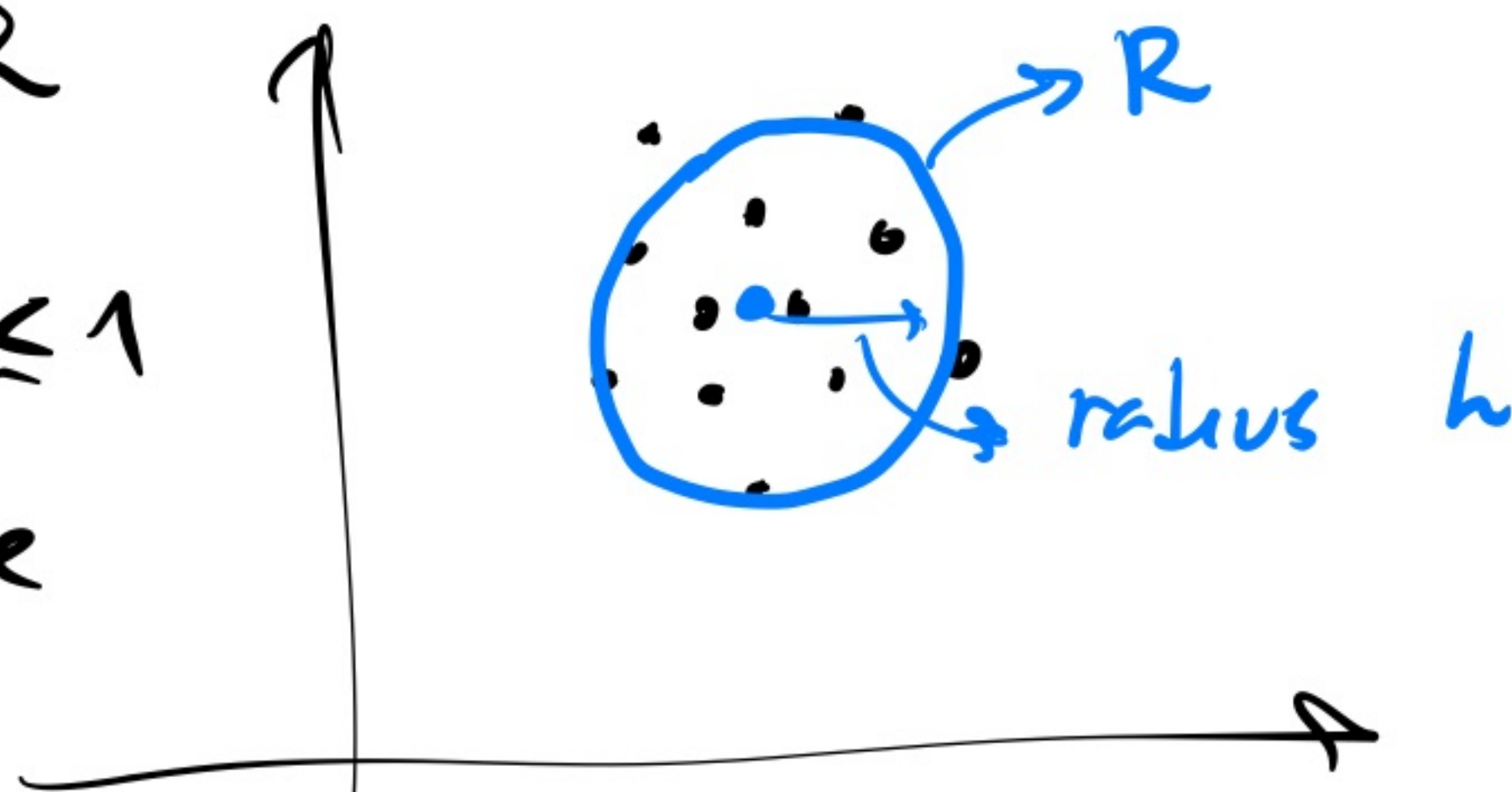


Kernel methods

Fix V and estimate k from the data

Count k enter in R

$$k(\vec{u}) = \begin{cases} 1 & \text{if } |\vec{u}| \leq 1 \\ 0 & \text{otherwise} \end{cases}$$



$$k\left(\frac{\vec{x} - \vec{x}_i}{h}\right) = \begin{cases} 1 & \forall i \text{ inside } R \\ 0 & \text{otherwise} \end{cases} \quad \left. \vphantom{k\left(\frac{\vec{x} - \vec{x}_i}{h}\right)} \right\} \Rightarrow \text{"kernel function"}$$

$$K = \sum_{i=1}^N k\left(\frac{\vec{x} - \vec{x}_i}{h}\right) \quad (II)$$

↙ reinterpreted as sum of N balls centered on $\vec{x}_i$

$$p(\vec{x}) = \frac{K}{N \bar{V}} = \frac{1}{N} \sum_{i=1}^N \frac{1}{V(h)} k\left(\frac{\vec{x} - \vec{x}_i}{h}\right)$$

1962 Parzen proved $\uparrow$ is an unbiased estimator of the true PDF

$$\text{i.e. } \lim_{N \rightarrow \infty} p(\vec{x}) = p_{\text{true}}(\vec{x})$$

under mild conditions on $k(\vec{u})$

- $\sup_{-\infty < u < \infty} |k(u)| < \infty$
- $\int_{-\infty}^{\infty} du |k(u)| < \infty$
- $\lim_{u \rightarrow \infty} |u k(u)| < \infty$
- $\int_{-\infty}^{\infty} k(u) du = 1$ $\left\{ \begin{array}{l} \text{the volume } V(h) \\ \text{is absorbed inside } k(u) \end{array} \right.$

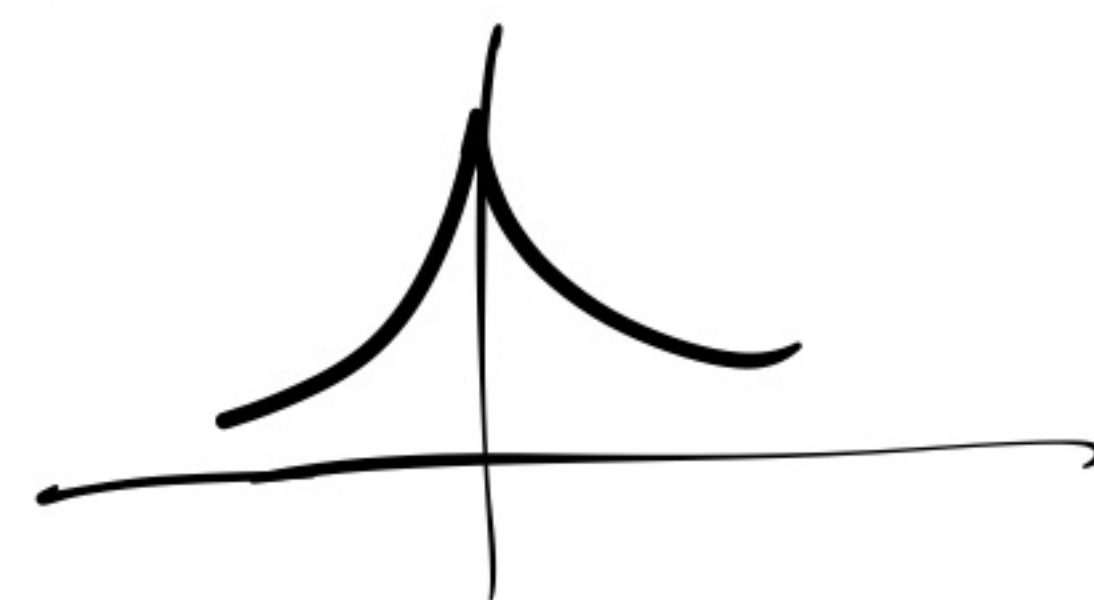
Examples of kernels

- Gaussian kernel

$$k(u; h) = \frac{1}{\sqrt{\pi}h} e^{-u^2/2h^2}$$

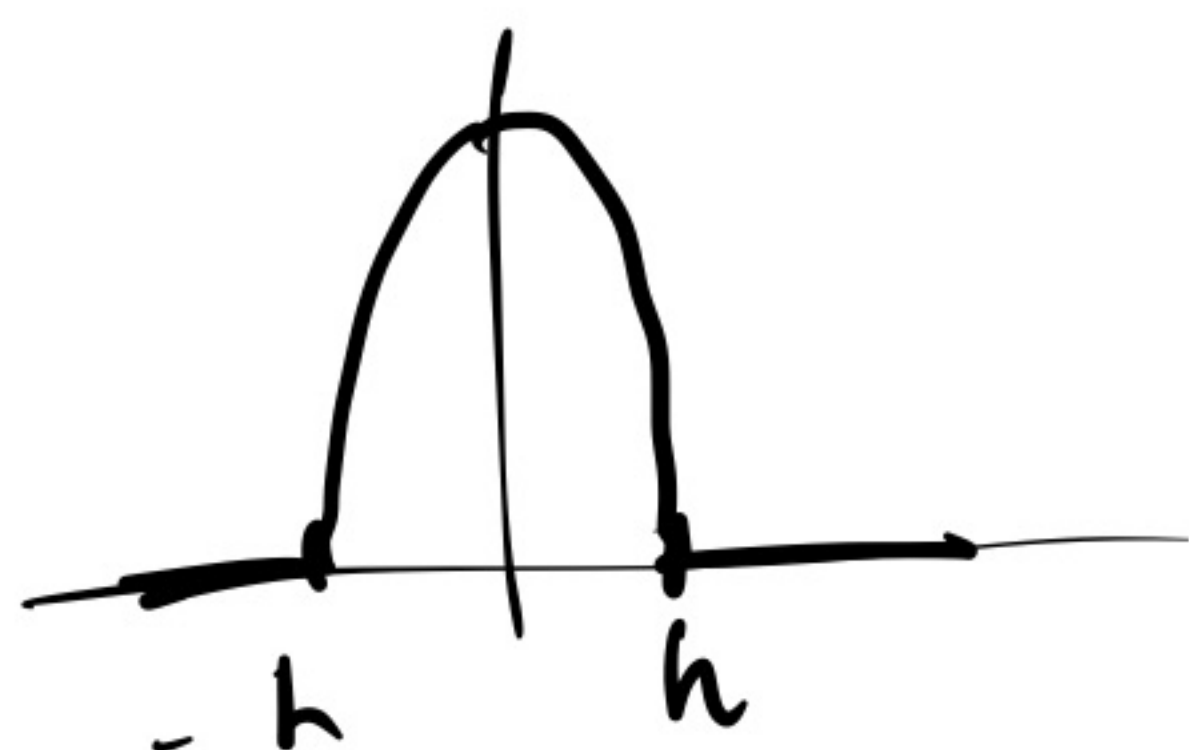
- Exponential

$$k(u; h) = \frac{1}{2h} e^{-|u|/h}$$



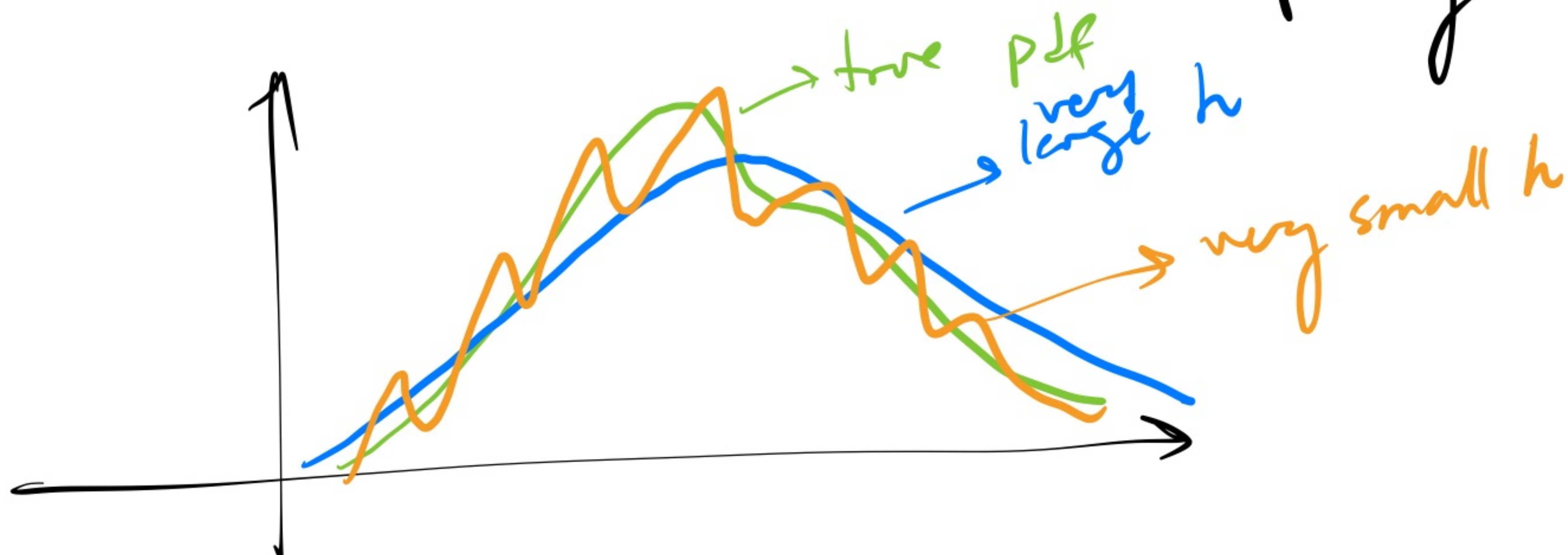
- Cosine kernel

$$k(u; h) = \begin{cases} \frac{\pi}{4h} \cos\left(\frac{\pi u}{2h}\right) & u \leq h \\ 0 & u > h \end{cases}$$



$\therefore$ many other kernels

Parameter $h \Rightarrow$ controls the model complexity



Kernel regression

$$D = \{ \vec{x}_i, y_i \}$$

$$\langle y(\vec{x}) \rangle = \frac{\int dy \cdot y \cdot p(\vec{x}, y)}{\int dy p(\vec{x}, y)}$$

$$p(\vec{x}, y) = \frac{1}{N} \sum_{i=1}^N k_x(\vec{x} - \vec{x}_i; h_x) k_y(y - y_i; h_y)$$

Generalisation of Specht '91

$$\langle y(\vec{x}) \rangle = \frac{\sum_i k_x(\vec{x} - \vec{x}_i; h_x) \int_{-\infty}^{\infty} dy \cdot y \cdot k(y - y_i; h_y)}{\sum_i k_x(\vec{x} - \vec{x}_i; h_x) \int_{-\infty}^{\infty} dy \cdot k(y - y_i; h_y)}$$

if $k_y(u; h_y) \Rightarrow$ follows a dist'n.

$$\int dy \cdot y \cdot k(y - y_i; h_y) = y_i$$

$$\langle y(\vec{x}) \rangle = \frac{\sum_i y_i \cdot k_x(\vec{x} - \vec{x}_i; h_x)}{\sum_i k_x(\vec{x} - \vec{x}_i; h_x)}$$

Nadaraya-Watson model

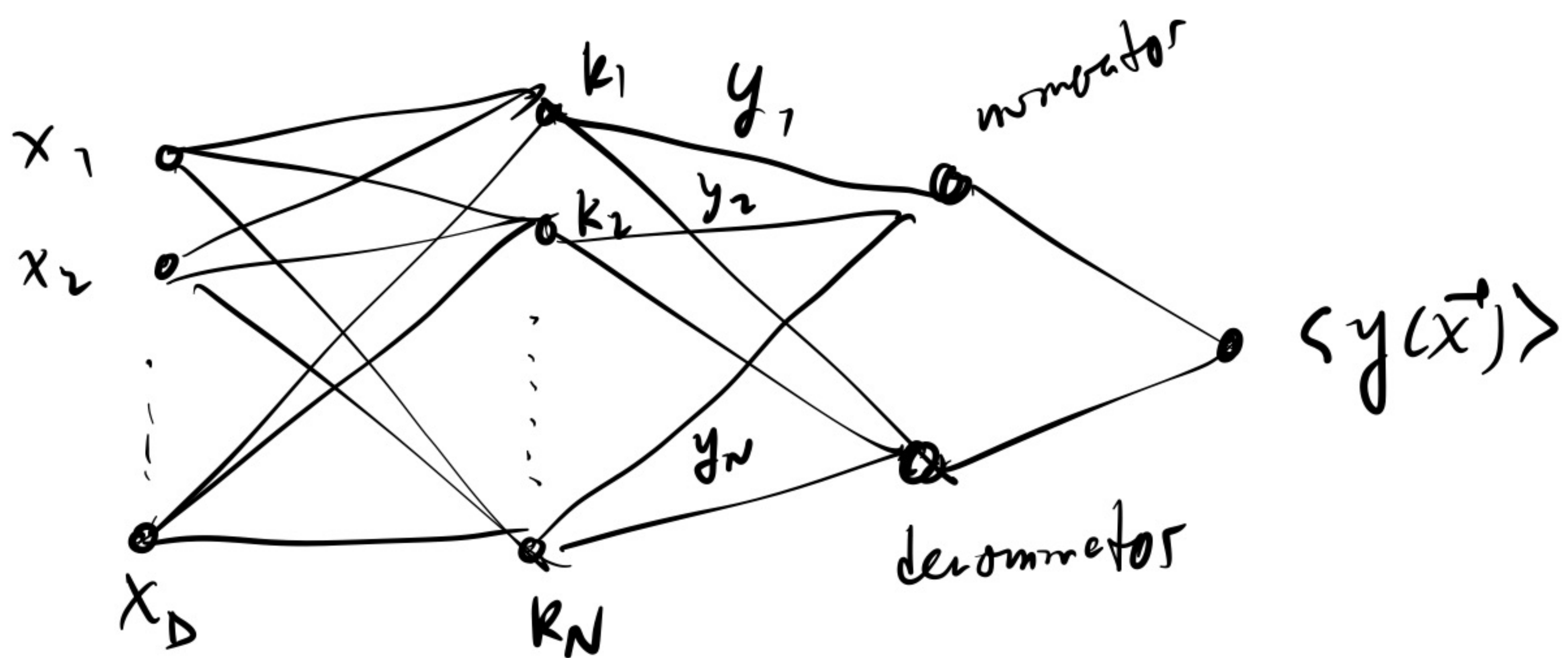
weighted average of all the data points

in the original paper (Specht)

↓
h optimised by minimising MSE

MLE

$$L(h_x, h_y) = \prod_{i=1}^N \frac{1}{N-1} \sum_{j \neq i}^{N-1} k_x(\vec{x}_i - \vec{x}_j; h_x) \times k_y(y_i - y_j; h_y)$$



⇒ "General Regression Neural Network"
(Specht)

Training is a one-step procedure

⇒ "the whole dataset stored in memory"

Prediction is costly

Classification with kernel methods

⇒ Generative approach

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

$$p(\vec{x} | C_k) = \frac{1}{N_k} \sum_{i=1}^{N_k} k_x(\vec{x} - \vec{x}_i; h)$$

$$p(C_k) = N_k / N$$

$$p(C_k | \vec{x}) \propto p(C_k) \frac{1}{N_k} \sum_{i=1}^{N_k} k_x(\vec{x} - \vec{x}_i; h)$$

$$\hat{C}_k = \operatorname{argmax} p(C_k | \vec{x})$$

⇒ "Probabilistic Neural network" ⇒