

Lecture 4

KDE
selecting h
by CV

K-NN density estimation

HW 1 TB posted

OH 4-5 Tue
B-321

CV $\rightarrow$ notes t.b.
posted

HTF :

Ch 6.1-3, 6.6
7.1-3, 7.10

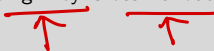
Lecture III – Kernel and k-NN density estimation

Marina Meilă
mmp@stat.washington.edu

Department of Statistics
University of Washington

STAT/BIOST 527
Spring 2023

Outline

- 1 Kernels ✓
- 2 Kernel density estimators ✓
- 3 Choosing h by Cross-Validation and the Bias-Variance trade-off ✓

- 4 The k-Nearest Neighbor density estimator 

Reading AoNPS Ch.: 4.1-4.3, 6.1-6.3, HTF Ch.: 

All of NP Statistics

Kernel density estimators

Data

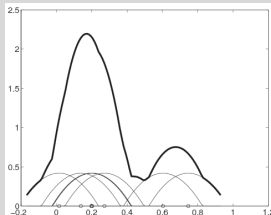
$$\mathcal{D} = \{x^i\}_{i=1}^n$$

sample indices

$$p_X = ?$$

Model family

$$b(\cdot) [\equiv K(\cdot)] \text{ kernel}$$

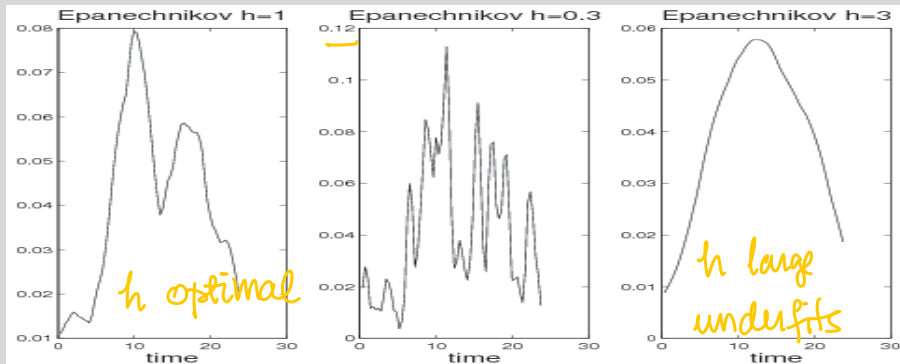
 h kernel width

$$\begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_d \end{bmatrix} = X$$

dim

$$\hat{p}_X(x) = \frac{1}{nh} \sum_{i=1}^n b\left(\frac{x - x^i}{h}\right)$$

Choosing h



* for this 2)
and this n*

h small
overfit

The Bias-Variance trade-off and CV

with separate $\mathcal{D}^V$ validation set

train on $\mathcal{D}$, validate on $\mathcal{D}^V$

$$|\mathcal{D}| = n$$

$$|\mathcal{D}^V| = n'$$

compute $l(\mathcal{D}^V | \mathcal{D}, h)$
 log-likelihood

K-fold CV

$\mathcal{D}$ split into $\mathcal{D}^1, \dots, \mathcal{D}^K$ equal size partition

$$\mathcal{D}^{-k} = \mathcal{D} \setminus \mathcal{D}^k$$

for $k = 1 : K$

• estimate $\hat{p}_x$ on $\mathcal{D}^{-k}$

• $l_{cv}^{(k)}(\mathcal{D}^k | \mathcal{D}^{-k}, h)$

$$\bar{l}_{cv} = \frac{1}{K} \sum_k l_{cv}^{(k)}, \quad \text{std } \bar{l}_{cv}$$

$\bar{l}_{cv}(h)$

The Bias-Variance trade-off

$$\bar{\ell}_{cv}(\mathcal{D}^k | \mathcal{D}^{-k}, h) = \frac{K}{n} \sum_{i \in \mathcal{D}^k} \ell(x^i | \mathcal{D}^{-k}, h)$$

$$\ell(x | \mathcal{D}, h) = -\ln \hat{p}_x(x | \mathcal{D}, h)$$

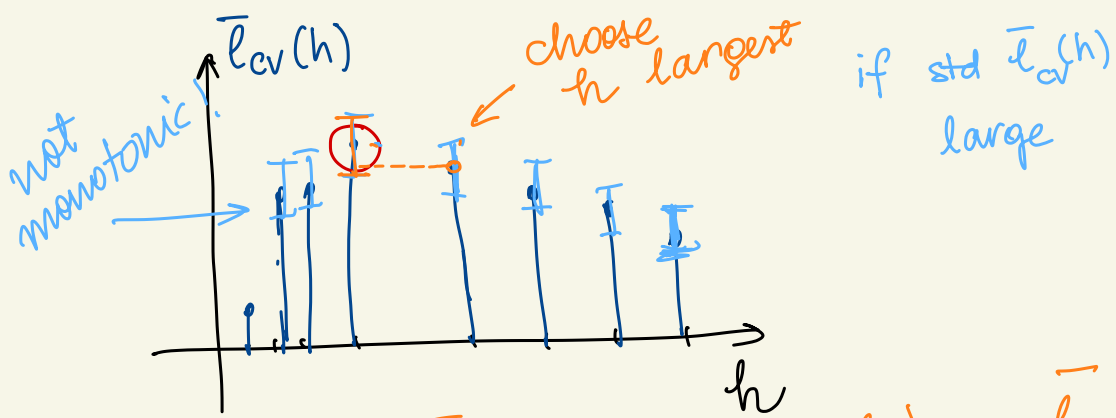
$$|\mathcal{D}^k| = \frac{n}{K}$$

$$\bar{\ell}_{cv}(h) = \frac{1}{K} \sum_{k=1}^K \sum_{i \in \mathcal{D}^k} \ell(x^i | \mathcal{D}^{-k}, h)$$

fold k

for every h
 Choose $h^{opt} = \underset{h}{\operatorname{argmax}} \bar{\ell}_{cv}(h)$

Ex: how many operations
 $\uparrow$
 evaluations of $b(\cdot)$



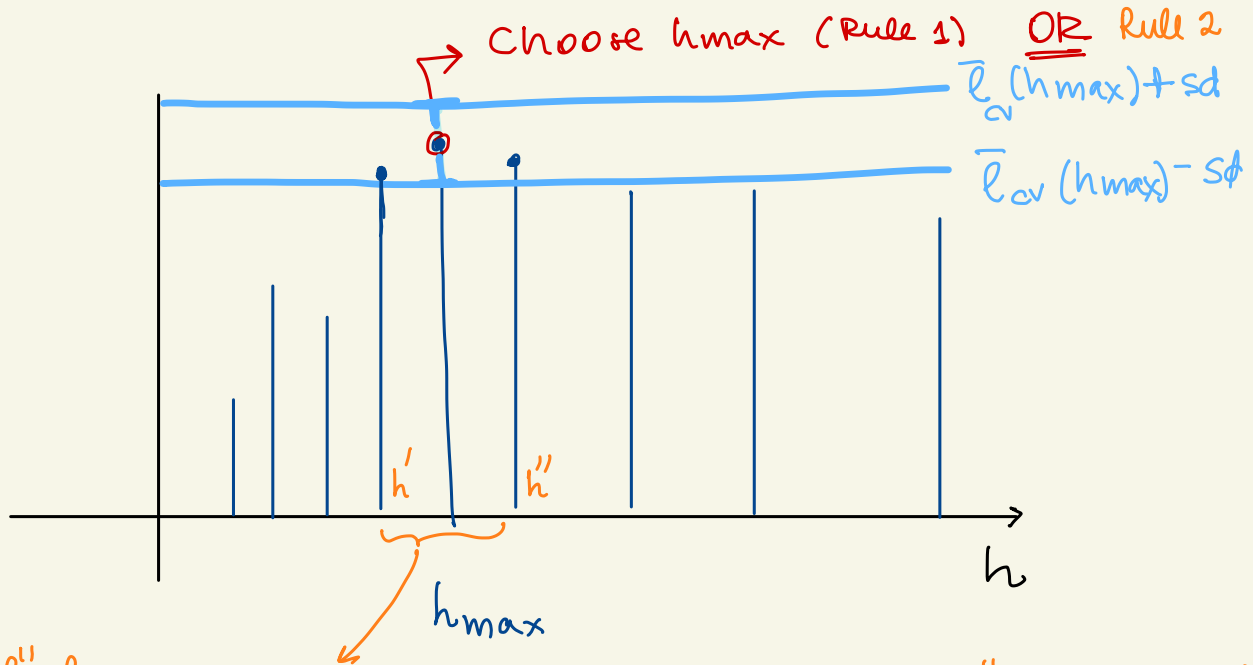
so that $\bar{l}_{cv}(h) + \text{std}(h) > \bar{l}_{cv}(h_{\max}) - \text{std}(h_{\max})$

"as long as $\pm \text{std}$ intervals overlap"

$K = n$
or K large $\} \Rightarrow \text{std } \bar{l}_{cv}(h)$ large

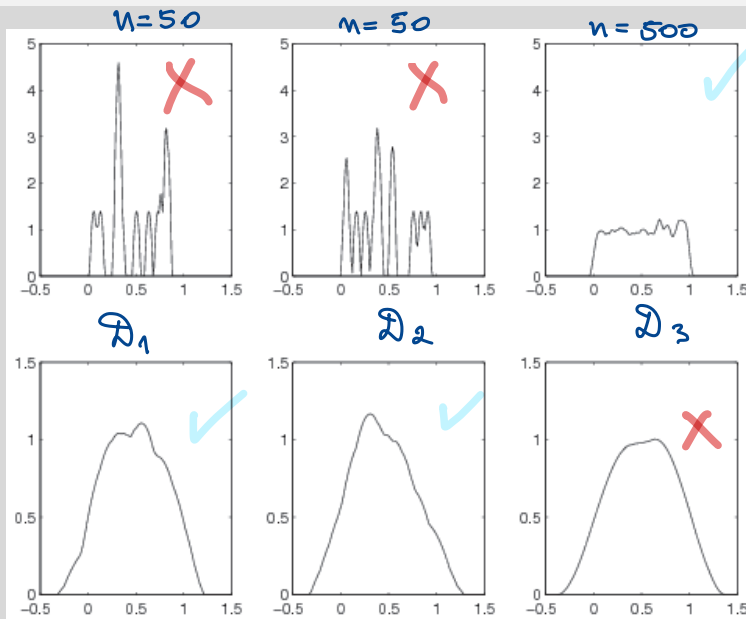
K small $\rightarrow$ too little data for training

Rule of thumb $\left\{ \begin{array}{l} n > 100 \\ n < \end{array} \right. \begin{array}{l} K = 5 \cdots 10 \\ K = n \text{ Leave One Out (LOO)} \end{array}$



$h', h'', h_{\max}$
 have $\bar{e}_{cv}$ in interval $\rightarrow$ choose largest $= h''$ (Rule 2)
 $\Rightarrow$ smoothest model
 $h' < h_{\max} < h''$

The Bias-Variance trade-off



h small

h large

The k-Nearest Neighbor density estimator

$$\hat{p}_{kNN}(x) = \frac{k}{n} \frac{1}{\omega_d r_k(x)^d} \quad (1)$$

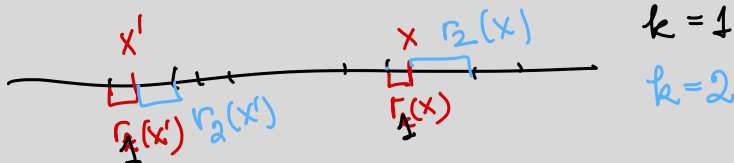
where

- n is sample size, $\mathcal{D} = \{x^1, \dots, x^n\}$
- $r_k(x)$ is distance to k -th nearest neighbor of x in $\mathcal{D}$
- d is the dimension of the data
- $\omega_d = \frac{\pi^{d/2}}{\Gamma(1+d/2)}$ is the volume of the unit ball in d dimensions

distance to k -th NN
of x

- $d = 1, \omega_1 = 2$
- $d = 2, \omega_2 = \pi$
- $d = 3, \omega_3 = \frac{4}{3}\pi$

$\omega_d r^d = \text{Volume Ball}(x, r)$



The k-Nearest Neighbor density estimator – Motivation

$$\hat{p}_{kNN}(x) = \frac{k}{n \underbrace{\omega_d r_k(x)^d}_{\text{Vol[Ball]}}}$$



Idea

- let $B(x, r)$ be the ball of radius r centered at x
- probability of any set approximated by the empirical distribution

$$\underbrace{Pr[B(x, r)]}_A \approx \frac{\# \text{ data points in } B(x, r)}{n} = \frac{|D \cap B(x, r)|}{n} \quad (2)$$

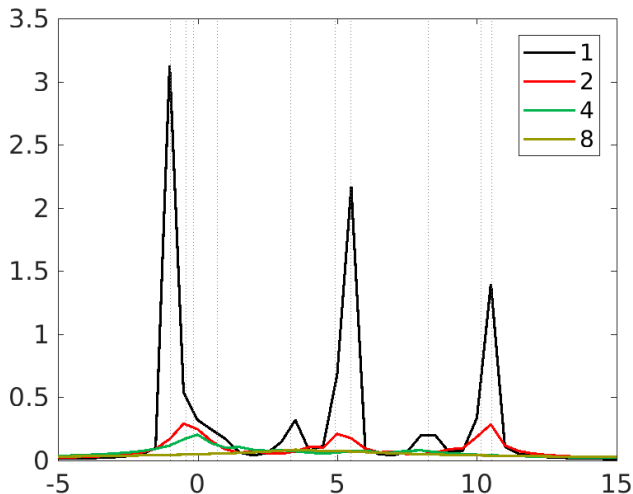
- density at x approximated by

$$\hat{p}_{kNN}(x) \approx \frac{Pr[B(x, r)]}{Vol[B(x, r)]} \quad (3)$$

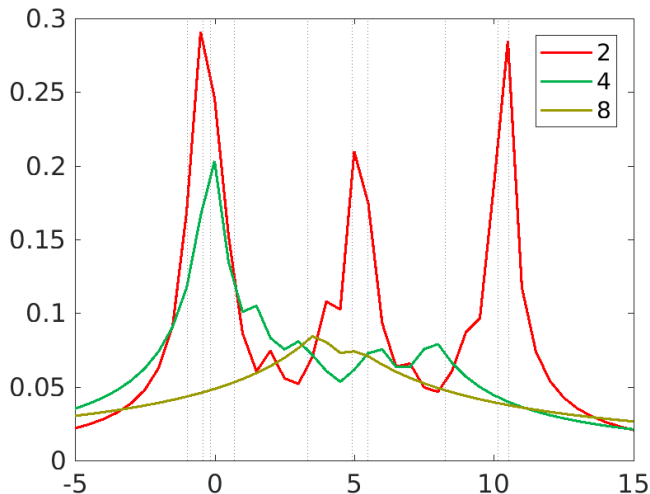
$\frac{P[\text{set}]}{Vol[\text{set}]}$

- $Vol B(x, r) = r^d \omega_d$
 - if $r = r_k(x)$, then $\# \text{ data points in } B(x, r) = k$
- ... putting it all together we get (1)

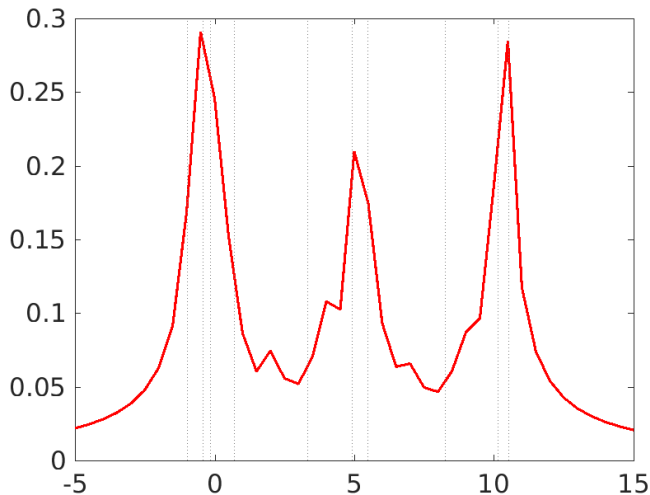
smoothing parameter k } $k=1$ Var $\uparrow$
 bias $\downarrow$
 k large bias $\uparrow$ var $\downarrow$

k-NN density estimator for $n = 10$ points

Note that when $d = 1$, $\int_{\mathbb{R}} \hat{p}_{kNN}(x) dx = \infty$!!!!!

k-NN density estimator for $n = 10$ points

Note that when $d = 1$, $\int_{\mathbb{R}} \hat{p}_{kNN}(x) dx = \infty$!!!!!

k-NN density estimator for $n = 10$ pointsonly $k = 2$ Note that when $d = 1$, $\int_{\mathbb{R}} \hat{p}_{kNN}(x) dx = \infty$!!!!!

Convergence rates

- for $d = 1$: $k^* \sim n^{4/5}$ $\text{MSE} \sim n^{-4/5}$ $\rightarrow$ 1) $\frac{k^*}{n} \rightarrow 0$ for $n \rightarrow \infty$
increases slowly

$$\text{KR } d=1 \quad \text{MSE} \sim n^{-\frac{4}{5}}$$

$$h \sim n^{-\frac{1}{5}}$$

KDE

— " —

$$n^{\frac{1}{d+4}}$$

$$E_{P_X}[(\hat{P}_X - P_X)^2 | n] = \text{MISE}$$