

24.10.2022. Part 2. Nonparametric Estimators

• Sketch of the lecture:

Introduction to Motivating examples.

Chapter 1. Density Estimators. (univariate) $\{x_i\}_{i=1}^n$

Chapter 2. Regression. (Hồi quy).

Chapter 3. Multivariate and Functional Cases. ($\mathbb{R}^k$)

Introduction.

Motivating examples:

In this lecture we assume that we observe for each individual either:

- A variable X (density)
 - A pair (x, y) where
 - y response variable
 - x explanatory
- Regression

We assume we get i.i.d random samples

1) $(x_1, \dots, x_n) \stackrel{i.i.d}{\sim} X$

$$\sigma_2 \left((x_1, y_1), \dots, (x_n, y_n) \right) \stackrel{\text{ind}}{\sim} (x, y)$$

A Galaxies dataset.

```
library(MASS) | data(galaxies) | hist(galaxies, freq = F, nclass = 100)
```

$\vec{v}$ theo trục $\vec{r}$ thì
 là d'chia

Ex. 1) get the data.

2) How to represent the distribution of these data.

3) Pros and cons.

↙ in d^2 chuyển d^2 .

hàm mật độ
xs phụ thuộc
vào không gian.

Pros	Cons.
- non parametric. - no parametric density assumed.	- stepwise not continuous constant on classes. - Strongly depending on the number and the nature of the classes.

Aim: Define a new estimator being:

- smooth.
- still non parametric.
- depending on less estimators.

↪ Discuss the non parametric estimators (o biết phân phối gì).

- Bắt đầu làm với non-parametric → như là mô hình parametric → đưa về parametric estimators.

	Pros	Cons.
p	- Great speed of convergence. - Easy to interpret.	- No flexibility. - A priori needed.
np	- Flexibility. - No a priori	- Low convergence rate. - More difficult to interpret. gần như.

⑧ Pressure data.

- Ex. 1) Get data. ^{phần 1}
- 2) Draw the scatter plot of pressure w.r. t temperature.
- 3) how to estimate the link function.

. data (pressure)

. attach (pressure)

$\left\{ \begin{array}{l} p = \text{pressure} \quad \$ \text{ pressure} \\ t = \text{ } \quad \$ \text{ temperature} \end{array} \right.$

. plot (t, p) $\rightarrow$ vẽ liên hệ giữa t và p.

$\hookrightarrow$ non linear.

. We are focusing on a regression model.

non parametric: $Y = r(x) + \varepsilon$, with $E[\varepsilon|x] = 0$.

response variable. $\uparrow$ explanatory variable. $\uparrow$ error

aim: estimate function r .

Remark.

$$\begin{aligned} E[Y|x] &= E[r(x) + \varepsilon|x] \\ &= E[r(x)|x] + \underbrace{E[\varepsilon|x]}_0 \\ &= r(x) \end{aligned}$$

where r is an unknown function we want to estimate.

Aim: Get a new estimator

- non parametric (no a priori on the nature)
- smooth.
- very flexible.

③ Doping data.

1) get the data

add file vào R : Session → Set working Directory → Choose Directory

2) If a new cyclist comes and we observe his haematocrit rate x_0 ,
could we give the class y_0 corresponding to that cyclist?

• Discrimination problem. Y : class $\{+, -\}$.

X : explanatory

Aim: Get a classifier, (plan box), constructed from the data that will propose a relevant class for any new value of X .

• Bayes classifier.

$$\hat{y}_0 = \underset{y \in Y}{\operatorname{argmax}} P(Y=y | X=x_0).$$

unknown → estimate

↪ don't optimize $\hat{y}_0(x)$ if Y is finite.

Problem: $y \mapsto P(Y=y | X=x_0)$ are unknown.

How to estimate r_y ?

$$Z_y = \mathbb{1}_{\{Y=y\}}$$

$$E[Z_y | X] = E[\mathbb{1}_{\{Y=y\}} | X] = P(Y=y | X) = r_y(x)$$

Note: $Y = r(x) + \varepsilon$

$$r(x) = E[Y | X]$$

$$\hookrightarrow Z_y = r_y(x) + \varepsilon_y$$

• Estimated Bayes classifier.

$$\hat{y}_0 = \underset{y \in \mathcal{Y}}{\operatorname{argmax}} \hat{r}_y(x_0),$$

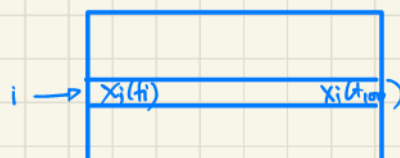
where $\hat{r}_y$ is a non parametric estimator of r_y .

① Tecator data.

```
.tec = read.table('tec', as.is=T)
```

```
X = tec[, 1:100]
```

```
Y = tec[, 101]
```



```
.install.packages('ggplot2')
```

```
library('ggplot2')
```

```
wl = seq(750, 948, 2)
```

```
plot(wl, X[1,], type='l', ylim = range(X), xlab='wavelength',  
      ylab='absorbance')
```

ghi trên trục Oy.

line

```
for (i in 2:213)
{
  lines(wl, x[i, ])
}

```

• Regression on functional variable: $Y = r(X) + \varepsilon$.

Y : real valued

X : functional variable.

⑤ Phoneme dataset

unique (PHONEME)

CURVES

PHONEME.

• par (mfrow = c(2,3))

↖ $\vec{v} \in \text{CURVES theo ctt.}$

```
(plot (CURVES[1, -], type = 'l', main = 'STT', ylim = range (CURVES)))
for (i in 2:250) {lines (CURVES[i, ])}

```

```
(
  _____ [201,] _____ = 'IY', _____
  _____ 202:400 _____
  _____ [401] _____ = 'De', _____
  _____ 402-650 _____
)

```

Discrimination problem with X functional.

Chapter 1.

Density Estimation.

$$X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} X \quad \text{Aim } \hat{f}_X$$

A) History.

• Parametric model

Assume $f_X = f_\theta$, $\theta \in \Theta \subseteq \mathbb{R}^k$.

$$\hat{f}_X = f_\theta$$

• Non parametric model:

1) histogram: thông qua biểu đồ.
properties discussed in the introduction part.

2) moving window estimator.

2a) non parametric estimation of the cumulative distribution function:

$$F_X(x) = P(X \leq x)$$

$$\text{for all } x \in \mathbb{R}, \hat{F}_X(x) = \frac{1}{n} \cdot \sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}$$

• Remark. $E[\mathbb{1}_{\{X \leq x\}}] = P(X \leq x) = F_X(x)$.

↳ unbiased estimator for F_X .

• Properties

1) for all $x \in \mathbb{R}$, $\hat{F}_X(x)$ is an unbiased estimator of $F_X(x)$.

$$2) \text{ ————— }, \text{var}(\hat{F}_X(x)) = \frac{F_X(x)(1-F_X(x))}{n} \leftarrow \text{giving Bar}$$

$$3) \text{ ————— }, \hat{F}_X(x) \xrightarrow[n \rightarrow +\infty]{\|\cdot\|^2} F_X(x)$$

$$4) \text{ ————— }, \hat{F}_X(x) \xrightarrow[n \rightarrow +\infty]{a.s.} F_X(x) \text{ (LLN)}$$

$$5) \text{ ————— }, \sqrt{n}(\hat{F}_X(x) - F_X(x)) \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N}(0, \sqrt{F_X(x)(1-F_X(x))})$$

2b) Moving window estimator.

Remark. $f_X(x) = \lim_{h \rightarrow 0^+} \frac{F_X(x+h) - F_X(x-h)}{2h}$

Proof.

$$\begin{aligned} \frac{F_X(x+h) - F_X(x-h)}{2h} &= \frac{F_X(x+h) - F_X(x) + F_X(x) - F_X(x-h)}{2h} \\ &= \frac{F_X(x+h) - F_X(x)}{2h} + \frac{F_X(x) - F_X(x-h)}{2h} \\ &= \frac{1}{2} \left[\frac{F_X(x+h) - F_X(x)}{h} + \frac{F_X(x) - F_X(x-h)}{h} \right] \\ &\xrightarrow{h \rightarrow 0^+} \frac{1}{2} [f_X(x) + f_X(x)] = f_X(x). \end{aligned}$$

hence, for h small

$$f_X(x) \approx \frac{F_X(x+h) - F_X(x-h)}{2h}$$

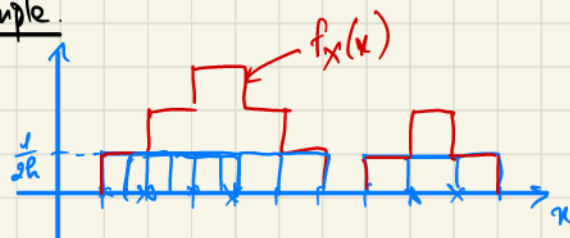
Hence a natural estimator could be

$$\tilde{f}_X(x) = \frac{\hat{F}_X(x+h) - \hat{F}_X(x-h)}{2h}$$

$$\begin{aligned} \hookrightarrow \tilde{f}_X(x) &= \frac{\frac{1}{n} \sum_{i=1}^n [\mathbb{1}_{\{x_i \leq x+h\}} - \mathbb{1}_{\{x_i \leq x-h\}}]}{2h} \\ &= \frac{\frac{1}{n} \sum_{i=1}^n \mathbb{1}_{\{x-h < x_i \leq x+h\}}}{2h} \\ &= \frac{1}{nh} \sum_{i=1}^n \frac{1}{2} \cdot \mathbb{1}_{]-1,1]} \left(\frac{x_i - x}{h} \right) \\ &= \frac{1}{nh} \sum_{i=1}^n w \left(\frac{x_i - x}{h} \right), \end{aligned}$$

where $w(t) = \frac{1}{2} \cdot \mathbb{1}_{]-1,1]}(t)$: kernel = 'r'.

• Example



$h=1$.

Pros	Cons
depends on less parameter	- still stepwise (non smooth and constant on each interval)

• Remark

$$\tilde{f}_X(x) = \frac{1}{nh} \sum_{i=1}^n \underbrace{\frac{1}{2} \mathbb{1}_{[-1,1]} \left(\frac{x_i - x}{h} \right)}_{\text{almost surely } \sim U[-1,1]}$$

Connect back to galaxies

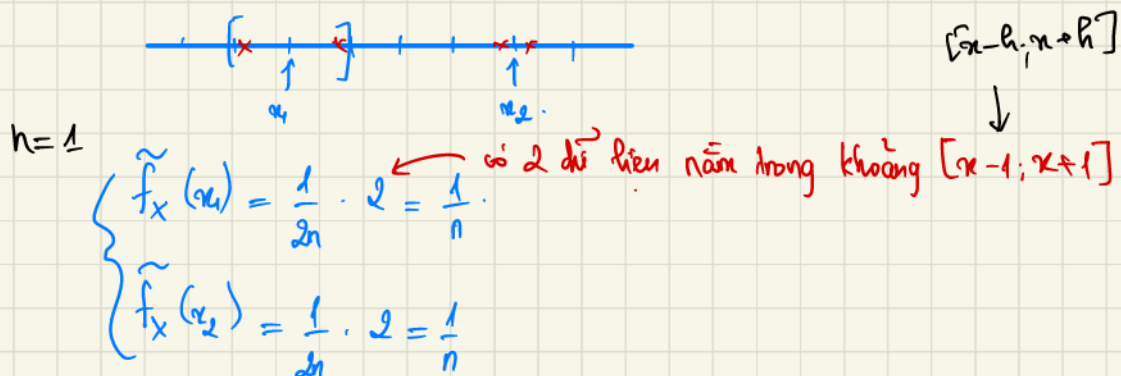
`plot(density(galaxies, kernel='r', bw=100))`

In file Commandes - Cours NP.R.

$x = smelange$
 $y = dmelange$

$$f(x) = \alpha f_1(x, p) + (1-\alpha) f_0(x, p).$$

• Comment:



We would like to improve our estimator in order to give more weight to observation that are closer to x than the other.

• Recall Properties of estimators of a function g

1) Pointwise properties. (tính d^2)

a) Bias, for all $x \in \mathbb{R}$, $\text{Bias}(\hat{g}(x)) = E[\hat{g}(x)] - g(x)$.

b) Variance ——— $\text{Var}(\hat{g}(x)) = E[(\hat{g}(x) - E[\hat{g}(x)])^2]$

c) Quadratic error ——— $E[(\hat{g}(x) - g(x))^2] = \text{Var}(\hat{g}(x)) + (\text{Bias}(\hat{g}(x)))^2$

2) Global (tính quát cho hàm g)

a) Sum of squared error: $SSE(\hat{g}) = \sum_{i=1}^n (\hat{g}(x_i) - g(x_i))^2$

b) Mean _____: $MSE(\hat{g}) = E\left[\sum_{i=1}^n (\hat{g}(x_i) - g(x_i))^2\right]$

c) Integrated Squared Error: $ISE(\hat{g}) = \int [\hat{g}(x) - g(x)]^2 dx$

d) Mean Integrated Squared Error:

$$MISE(\hat{g}) = \int E[\hat{g}(x) - g(x)]^2 dx$$

• Property

1) If $h = h_n \xrightarrow{n \rightarrow +\infty} 0$ then for all $x \in D_{f_x}$, $\tilde{f}_x(x)$ is an asymptotically unbiased estimator of $f_x(x)$.

2) If $h = h_n \xrightarrow{n \rightarrow +\infty} 0$ and $nh_n \xrightarrow{n \rightarrow +\infty} +\infty$, then for all $x \in D_{f_x}$, with $f_x(x) > 0$: $\tilde{f}_x(x) \xrightarrow[n \rightarrow +\infty]{L^2} f_x(x)$

• Remarks

$$E[\tilde{f}_x(x)] = \frac{f_x(x+h) - f_x(x-h)}{2h} \xrightarrow{h \rightarrow 0^+} f_x(x).$$

$$\text{Var}(\tilde{f}_x(x)) \sim \frac{1}{2nh} f_x(x).$$

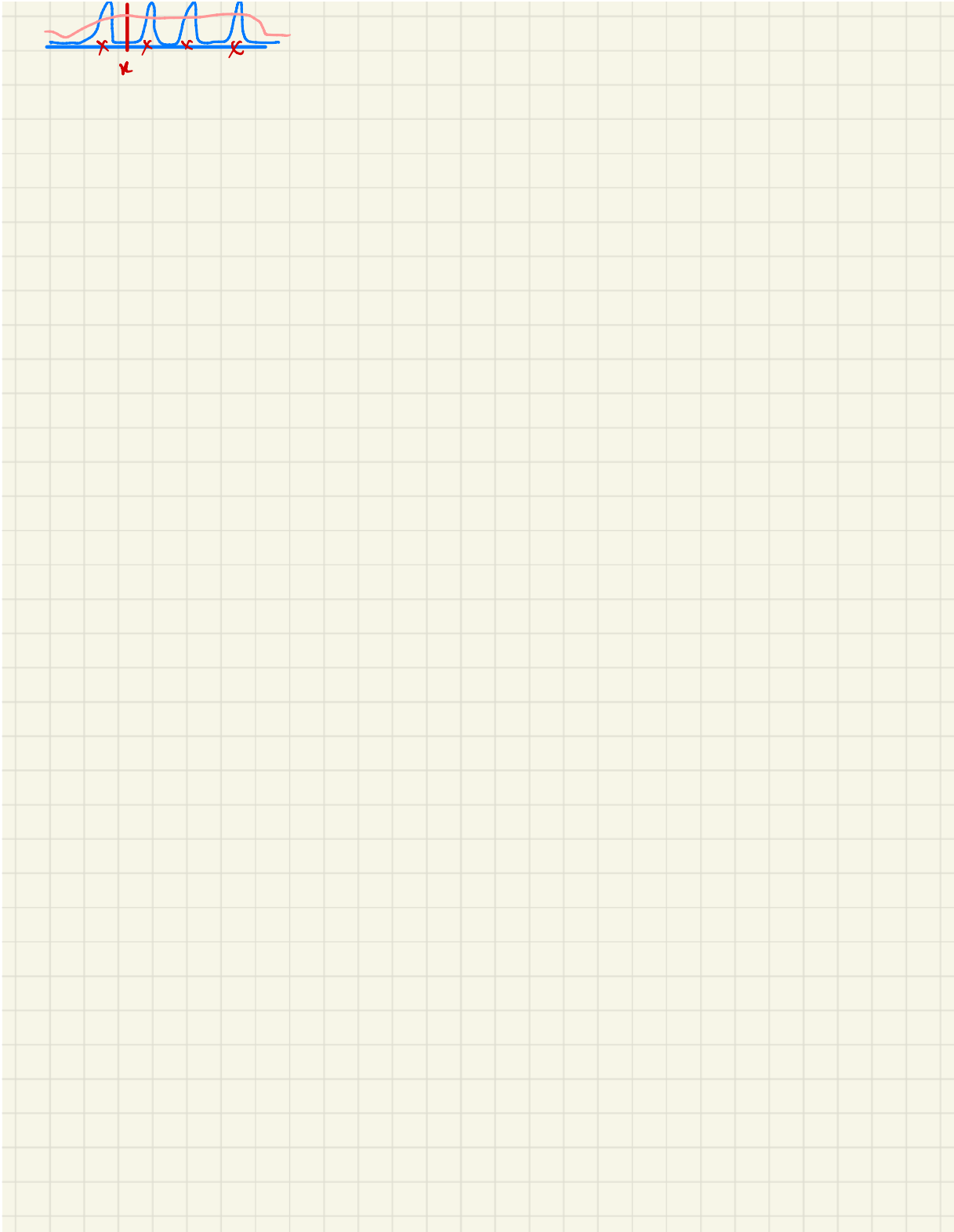
For n fixed $(-)$ h small $|$ h great $(-)$

Bias small	Bias great
Var great	Var small



depend on vị trí của dữ liệu.

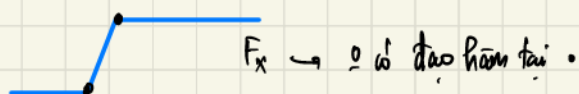
phụ thuộc vào dữ liệu.



26.10.2022.

B. Kernel estimator of the density function.

Either f_x is continuous at x .
or f_x is defined as F'_x



B. Definition.

Idea: Replace the indicator function $w(t) = \frac{1}{2} \cdot \mathbb{1}_{[-1,1]}(t)$ by a smoother function called kernel and usually denoted K .

We assume that the kernel fulfills the following assumptions:

1) $\forall t \in \mathbb{R}, K(t) \geq 0$

2) $\int K(t) dt = 1$.

3) $\int t \cdot K(t) dt = 0$. $\swarrow$ là giá trị nhất tại $t=0$

4) K is maximum at 0 and decrease when $|t| \uparrow$.

Ex: $K(t) = \frac{1}{\sqrt{2\pi}} e^{-t^2/2} \rightarrow f_x$ via $\mathcal{N}(0,1)$.

Definition. Kernel density function estimator Parzen 1964:

$$\hat{f}_x(x) = \frac{1}{n h} \sum_{i=1}^n K\left(\frac{x_i - x}{h}\right)$$

where K is a kernel, $h > 0$ is a smoothing parameter.

Main kernel examples

name	expression $K(t)$	graph
uniform	$\frac{1}{2} \mathbb{1}_{[-1,1]}(t)$	
gaussian	$\frac{1}{\sqrt{2\pi}} e^{-t^2/2}$	
triangular	$(1- t) \mathbb{1}_{[-1,1]}(t)$	
Epanechnikov	$\frac{3}{4} (1-t^2) \mathbb{1}_{[-1,1]}(t)$	
Quartic	$\frac{15}{16} (1-t^2)^2$	
Triweight	$\frac{35}{32} (1-t^2)^3$	
Cotinus	$\frac{\pi}{4} \cos\left(\frac{\pi t}{2}\right)$	

Remarks

- 1) K is chosen to be maximal at 0 to give more weight to X_i that are close to x than to the others.
- 2) Using a smooth kernel will lead to a smooth estimator.
- 3) The role of h is exactly the same as for the moving window estimator.

Comeback to R software.

- 1) Replace in the previous ex kernel = 'rectangular' by kernel = 'g' or 'e'

gaussian
 ↙
 epanechnikov

 par(mfrow = c(2,3))

```

par(mfrow = c(1,1))
res = density(x, kernel = 'r')
plot(res)
lines(density(x, kernel = 'g', bw = res$bw), col = 'blue')
lines(s, dmlarge, green)

```

rectangular
 Gaussian
 kernel
 : true density

• Come back to the galaxies.

```

library(MASS)
data(galaxies)
x = galaxies

```

• Remarks.

- 1) The effect of K is not so important (only for the smoothness it could be)
- 2) observe errors : at the points where the density is high the curvature of the density is important.
 do cong.

Pros	Cons.
Smooth and non constant by part.	depend on K (no important effect) R
weights of data depends on their proximity w.r.t x .	

8.2. Properties of the convergence rate.

Proposition:

- 1) If $K \geq 0$ then $\hat{f} \geq 0$
- 2) If K is a density $\hat{f}$ is a density.
- 3) If K is continuous then $\hat{f}$ is continuous
- 4) $C^p \xrightarrow{\quad} C^p$

Proof. Exercise.

Theorem — If f_x is C^2 on a neighborhood of x , a bounded (this could be relaxed if K has a compact support).

If $h_n \xrightarrow{n \rightarrow \infty} 0$ and $nh_n \xrightarrow{n \rightarrow \infty} +\infty$ then:

$$1) \text{ Bias } (\hat{f}(x)) = \frac{h_n^2}{2} f_x''(x) \mu_2(K) + o(h_n^2)$$

h_n small $\rightarrow$ Bias small
plus $\in$ vào $f_x''(x)$.

where $\mu_2(K) = \int t^2 K(t) dt$.

2) If $f_x(x) > 0$ then:

$$\text{MSE}(\hat{f}(x)) = \frac{h_n^4}{4} (f_x''(x))^2 \mu_2(K) + \frac{1}{nh_n} f_x(x) R(K) + o(h_n^4) + o\left(\frac{1}{nh_n}\right)$$

where $R(K) = \int K^2(t) dt$.

h_n small, $f_x(x)$ great $\rightarrow$ MSE great.

$\hookrightarrow$ Cần dùng định lý Borel bị chặn.

Almost sure convergence result:

Under some technical assumptions:

$$\hat{f}_x(x) = f_x(x) + O_{a.s.}(h^2) + O_{a.s.}\left(\sqrt{\frac{\ln(n)}{nh}}\right)$$

• Theorem If f_x is bound and C^2 , $h_n \rightarrow 0$

$$n \cdot h_n \rightarrow +\infty$$

$$f_x, f'_x, f''_x \in L^2(\mathbb{R})$$

then:

$$MISE(\hat{f}) = \frac{h^4}{4} \int (f''_x(t))^2 dt \mu_2^2(K) + \frac{1}{nh} R(K) + O\left(h^4 + \frac{1}{nh}\right)$$

↳ global. → aim: minimize MISE

• Remarks. from the asymptotic expression of MISE and AMISE, we can try to get optimal values of h .

$$h \mapsto AMISE(h) = \frac{h^4}{4} \int (f''_x(t))^2 dt \mu_2^2(K) + \frac{1}{nh} R(K)$$

. Then we try to optimize AMISE w.r.t h :

$$AMISE'(h) = h^3 \int (f''_x(t))^2 dt \mu_2^2(K) - \frac{1}{nh^2} R(K).$$

$$AMISE'(h) = 0 \Rightarrow h^3 \int (f''_x(t))^2 dt \mu_2^2(K) = \frac{1}{nh^2} R(K)$$

$$\Leftrightarrow h^5 = \frac{n^{-1} \cdot R(K)}{\int (f''_x(t))^2 dt \mu_2^2(K)}$$

$$\Leftrightarrow h = \left(\frac{R(K)}{R(f''_x) \mu_2^2(K)} \right)^{1/5} \cdot n^{-1/5}$$

↳ check that $nh \rightarrow +\infty$ and $h \rightarrow 0$.

$$\bullet \hat{p}_{AMISE} = \left(\frac{R(K)}{R(f_x'') \mu_2^2(K)} \right)^{1/5} \cdot n^{-1/5} : \text{global (optimal } \forall x)$$

• On the same way:

$$\hat{p}_{AMISE} = \left[\frac{f_x(x) \cdot R(K)}{(f_x''(x))^2 \cdot \mu_2^2(K)} \right]^{1/5} \cdot n^{-1/5} : \text{local (optimal } \forall x)$$

$$\begin{aligned} AMISE(\hat{p}_{AMISE}) &= \frac{1}{4} \left[\frac{R(K)}{R(f_x'') \mu_2^2(K)} \right]^{4/5} \cdot R(f_x'') \mu_2^2(K) \\ &\quad + n! \cdot n^{1/5} \left[\frac{R(K)}{R(f_x'') \mu_2^2(K)} \right]^{-1/5} \cdot R(K) \\ &= \frac{5}{4} \cdot [R^4(K) \cdot R(f_x'') \cdot \mu_2^2(K)]^{1/5} n^{-4/5} \text{ (min)} \end{aligned}$$

• Remarks

• in the parametric case usually: $E_{\theta}[(\hat{\theta} - \theta)^2] \sim \frac{c}{n}$

if in addition:

1) for all x : $\theta \rightarrow f_{\theta}(x)$ is C^2 .

2) $\exists h \in L^2(\mathbb{R})$, $\frac{\partial^2 f_{\theta}(x)}{\partial \theta^2} \leq h(x)$. then

$$MISE(\hat{f}_{\hat{\theta}}) \sim \frac{C}{n}$$

• for a histogram, with all the classes of the same amplitude h
we may get $MISE(\text{hist}) \sim C_2 \cdot n^{-2/3}$ biến số

↑
histogram.

→ parametric estimator $\hat{\theta}$ tốt hơn.

↪ 2 parameters $\leftarrow$ kernel K
smoothing h .

• $\hat{p}_{AMISE}$ đã và có f_x bên trong $\rightarrow$ tìm sieve estimate.

B₃. Optimal choice of the kernel K .

Idea: the optimal expression of $AMISE(\hat{p}_{AMISE})$ depends on:

$$C(K) = [R^4(K) \cdot \mu_2^2(K)]^{1/5}.$$

We naturally want to minimize $C(K)$ w.r.t K under the constraint that K is a kernel.

Result: The winner is the **Epanechnikov kernel (1969)**

$$K_{opt}(t) = \frac{3}{4} \cdot (1-t^2) \cdot \mathbb{1}_{[-1,1]}(t)$$

The **efficiency** of K is defined:

$$eff(K) = \left[\frac{C(K_{opt})}{C(K)} \right]^{5/4}$$

What does it mean?

$AMISE(n, K) = AMISE(n_0, K_{opt}) \rightarrow$ nếu $K \neq K_{opt}$:
phải chọn n lớn hơn n_0

$$\Leftrightarrow n^{-4/5} \cdot C(K) = n_0^{-4/5} C(K_{opt})$$

$$\Leftrightarrow \left(\frac{n_0}{n} \right)^{4/5} = \frac{C(K_{opt})}{C(K)}$$

$$\Leftrightarrow \frac{n_0}{n} = eff(K) < 1 \quad \left(\text{since } \frac{C(K_{opt})}{C(K)} < 1 \right)$$

$$\Leftrightarrow n_0 = eff(K) \cdot n.$$

$$\Rightarrow n = n_0 \cdot \frac{1}{\text{eff}(k)}$$

↳ nếu $k \neq k_{\text{opt}} \rightarrow$ cần nhiều dữ liệu hơn n_0 .

values $\text{eff}(k)$		$K(t)$
	1	Exponential $\frac{3}{4}(1-t^2) \cdot \mathbb{1}_{[-1,1]}(t)$
✓	0,994	quartic $\frac{15}{16}(1-t^2)^2 \cdot \mathbb{1}_{[-1,1]}(t)$
	0,987	Triweight $\frac{35}{32}(1-t^2)^3$
↖	0,951	Gaussian
?	0,986	Triangular $(1- t)$
	0,93	Uniform

b4) Automatic choice of the smoothing parameter.

- 1) The normal reference.
- 2) The plug-in method.
- 3) The cross-validation method.

1) The normal reference.

$$h_{\text{NWSE}} = \left[\frac{R(k)}{R(f_x'') \cdot \mu_2^2(k)} \right]^{1/5} \cdot n^{-1/5}$$

• idea: compute $R(f_x'')$ using the Gaussian hypothesis on the density:

$$f_x(t) = \frac{1}{\sigma \sqrt{2\pi}} e^{-(t-\mu)^2 / 2\sigma^2}$$

$$f'_x(t) = \frac{1}{\sigma\sqrt{2\pi}} \cdot \frac{-2(t-\mu)}{2\sigma^2} \cdot e^{-(t-\mu)^2/2\sigma^2}$$

$$= \frac{1}{\sigma\sqrt{2\pi}} \cdot \frac{-(t-\mu)}{\sigma^2} \cdot e^{-(t-\mu)^2/2\sigma^2}$$

$$f''_x(t) = \frac{1}{\sigma\sqrt{2\pi}} \left[\frac{-1}{\sigma^2} e^{-(t-\mu)^2/2\sigma^2} + \frac{(t-\mu)^2}{\sigma^4} \cdot e^{-(t-\mu)^2/2\sigma^2} \right]$$

$$= \frac{1}{\sigma\sqrt{2\pi}} \left[\frac{(t-\mu)^2}{\sigma^4} - \frac{1}{\sigma^2} \right] e^{-(t-\mu)^2/2\sigma^2}$$

$$\Rightarrow R(f''_x) = \int (f''_x(t))^2 dt$$

$$= \int \frac{1}{2\pi\sigma^2} \cdot \frac{1}{\sigma^4} \cdot \left[\frac{(t-\mu)^2}{\sigma^2} - 1 \right]^2 e^{-(t-\mu)^2/\sigma^2} dt$$

Let $z = \frac{\sqrt{2}(t-\mu)}{\sigma} \Rightarrow dz = \frac{\sqrt{2}}{\sigma} dt$ (để tiện p² hơn)

$$\Rightarrow R(f''_x) = \int \frac{1}{2\pi\sigma^2} \cdot \frac{1}{\sigma^4} \cdot \left(\frac{z^2}{2} - 1 \right)^2 \cdot e^{-z^2/2} \cdot \frac{\sigma}{\sqrt{2}} dz$$

$$= \frac{1}{2\sqrt{\pi}\sigma^5} \int \frac{1}{\sqrt{2\pi}} \left(\frac{z^2}{2} - 1 \right)^2 e^{-z^2/2} dz$$

$$= \frac{1}{2\sqrt{\pi}\sigma^5} \int \frac{1}{\sqrt{2\pi}} \left(\frac{z^4}{4} - z^2 + 1 \right) e^{-z^2/2} dz$$

$$= \frac{1}{2\sqrt{\pi}\sigma^5} \left(\frac{E[z^4]}{4} + 1 - \underbrace{E[z^2]}_1 \right) \text{ where } z \sim N(0,1)$$

$$= \frac{1}{2\sqrt{\pi}\sigma^5} \cdot \frac{3}{4} = \frac{3}{8\sqrt{\pi}\sigma^5}$$

$$\bullet R(f'_x) = \frac{3}{8\sqrt{\pi} \cdot \hat{\sigma}^5}, \text{ where } \hat{\sigma} = \min\left(S_{n-1}, \frac{R}{1.349}\right)$$

$$R = q_{0.75} - q_{0.25} \text{ : tính dựa vào data}$$

↑
phần vị thứ 0.75

• Remark.

If $Z \sim N(0,1)$ then $\overset{\text{tính dựa vào } p^2}{q_2(0.75)} - q_2(0.25) = 1.349$. ← theory quantile.

If $X \sim N(\mu, 1)$ then $q_x(0.75) - q_x(0.25) = 1.349$.

If $Y \sim N(\mu, \sigma)$ then $q_Y(0.75) - q_Y(0.25) = \sigma \times 1.349$.

Indeed, $P(Y \leq t) = 0.75$

$$\Rightarrow P\left(\underbrace{\frac{Y-\mu}{\sigma}}_Z \leq \underbrace{\frac{t-\mu}{\sigma}}_{q_2(0.75)}\right) = 0.75.$$

$$\Rightarrow q_Y(0.75) = \sigma \cdot q_2(0.75) + \mu.$$

$$\underline{q_Y(0.25) = \sigma \cdot q_2(0.25) + \mu.}$$

$$\Rightarrow q_Y(0.75) - q_Y(0.25) = \sigma \cdot (q_2(0.75) - q_2(0.25)) = \sigma \times 1.349.$$

$\overset{\text{normal reference}}{h_{NR}} = \left(\frac{8\sqrt{\pi} \cdot \hat{\sigma}^5}{3} \cdot \frac{R(K)}{\mu_2^2(K)} \right)^{1/5} \cdot n^{-1/5}.$

$$= \hat{\sigma} \cdot \left(\frac{8\sqrt{\pi} \cdot R(K)}{3\mu_2^2(K)} \right)^{1/5} \cdot n^{-1/5}$$

Kernel	$\left(\frac{8\sqrt{R(K)}}{3\mu_2(K)} \right)^{1/5}$
Epanechnikov	2.34
Gaussian	1.06
Triweight	2.78

- Pros: Easy to compute and implement.
- Cons: If our normal assumption is incorrect, we may be wrong.

2) The plug-in method

Idea: Estimate $R(f_x'')$ by $R(\hat{f}_0'')$ where $\hat{f}_0$ is a non optimal kernel estimator of f_x . A C^2 kernel is used. (Sheater-Jones, 1991)

b) Cross-validation

$$\begin{aligned}
 \text{MISE}(\hat{f}) &= E \left[\int [\hat{f}_x(t) - f_x(t)]^2 dt \right] \\
 &= E \left[\int [\hat{f}_x^2(t) - 2\hat{f}_x(t)f_x(t) + f_x^2(t)] dt \right] \\
 &= E \left[\int \hat{f}_x^2(t) dt \right] - 2E \left[\int \hat{f}_x(t)f_x(t) dt \right] + \int f_x^2(t) dt.
 \end{aligned}$$

• Remark: minimizing $\text{MISE}(\hat{f})$, is equivalent to minimizing $\text{MISE}(\hat{f}) - \int f_x^2(t) dt$

• Remark: $\int \hat{f}_x^2(t) dt$ is an unbiased estimator of $E \left[\int \hat{f}_x^2(t) dt \right] = P(h)$

• $\frac{1}{n} \sum_{i=1}^n \hat{f}^{-i}(X_i)$, where $\hat{f}^{-i}$ is the kernel estimator computed from the subsample where X_i has been removed.

$$\hat{f}^{-i}(t) = \frac{1}{(n-1) \cdot h} \sum_{\substack{j=1 \\ j \neq i}}^n K\left(\frac{X_j - t}{h}\right)$$

↳ X_i ra $\frac{1}{n-1}$ der $f_{X,h}$.

• Fact: $E\left[\frac{1}{n} \sum_{i=1}^n \hat{f}^{-i}(X_i)\right] = E\left[\int f_X(t) \cdot \hat{f}_X(t) dt\right]$

Proof $E\left[\frac{1}{n} \sum_{i=1}^n \hat{f}^{-i}(X_i)\right] = E[\hat{f}^{-1}(X_1)]$ (since $\hat{f}^{-i}(X_i) \sim \text{i.i.d.}$)

$$= E[E[\hat{f}^{-1}(X_1) | X_2, \dots, X_n]] = E\left[\int \hat{f}^{-1}(t) f(t) dt\right]$$

$$= E\left[\int \frac{1}{(n-1) \cdot h} \sum_{j=2}^n K\left(\frac{X_j - t}{h}\right) f_X(t) dt\right]$$

$$= \int \frac{1}{(n-1) \cdot h} \sum_{j=2}^n E\left[K\left(\frac{X_j - t}{h}\right)\right] f_X(t) dt$$

$$= \int \frac{1}{h} \int K\left(\frac{y - t}{h}\right) f_X(y) dy \cdot f_X(t) dt.$$

we need to prove that $= E[\hat{f}_X(t)]$

$$= \int \frac{1}{h} \cdot E\left[K\left(\frac{X - t}{h}\right)\right] f_X(t) dt.$$

• $E\left[\int \hat{f}_X(t) f_X(t) dt\right] = \int f_X(t) \cdot E\left[\frac{1}{n \cdot h} \sum_{i=1}^n K\left(\frac{X_i - t}{h}\right)\right] dt$

$$= \int f_X(t) \cdot \frac{1}{n \cdot h} E\left[\sum_{i=1}^n K\left(\frac{X_i - t}{h}\right)\right] dt$$

$$= \int f_X(t) \cdot \frac{1}{h} \cdot \mathbb{E} \left[K \left(\frac{X-t}{h} \right) \right] dt.$$

$$\Rightarrow \mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n \hat{f}^{-i}(X_i) \right] = \mathbb{E} \left[\int f_X(t) \cdot \hat{f}(t) dt \right] \quad \square$$

Now,

$$\hat{\varphi}(h) = \int \hat{f}^2(t) dt - 2 \cdot \frac{1}{n} \sum_{i=1}^n \hat{f}^{-i}(X_i)$$

is hence an unbiased estimator of $\varphi(h)$

And we have:

$$h_{cv} = \underset{h}{\operatorname{argmin}} \hat{\varphi}(h)$$

cross-validation.

Back to R:

• `x = dmelange`

màu định

`plot(density(x, bw = 'nrd'))`

`bw = 'nrd'`: normal

reference

`lines(—(x, bw = 'nrd'), col = 'blue')`

`= 'sj'`: plug-in method

unbiased

`= 'cv'`: cross-validation

`— 'sj' — 'green'`

`:`

`'cv'`

`lines(s, dmelange ...)`: hàm chiếu xdc

↳ `cv` is okay nhất.

• `library(sm)`

`sm.density(x, method = 'normal')` (hoặc cho `h = —`)

`— 'cv', add = T, col = 'blue'`

`— 'sj', — green`

B5. Mode estimation. $\leftarrow \arg \max_i x_i^2$ mode.

$$\theta_0 \in \arg \max_{\theta \in \mathbb{R}} f_x(\theta)$$

$$\hat{\theta} \in \arg \max_{\theta \in \mathbb{R}} \hat{f}_x(\theta)$$

• Idea: compute on a grid $t_1, \dots, t_p$ $\hat{f}_x$
 $\hat{f}(t_1) \dots \hat{f}(t_p)$

and get t_i giving the maximum value.

a, b, M : $t_1, \dots, t_p$ from a to b , M points.

↳ ? : change max to min?

28.10.2022.

Chapter 2. Kernel estimator for nonparametric regression.

Reminder: Regression model:

$$Y(X) = r(X) + \varepsilon, \text{ with } \begin{cases} Y, X \text{ are real valued.} \\ E[\varepsilon|X] = 0 \\ \Rightarrow r(x) = E[Y|X]. \end{cases}$$

The function r is unknown and we want to estimate it.

Linear regression: $r(x) = a \cdot x + b$.

$(x_1, y_1), \dots, (x_n, y_n) \stackrel{i.i.d.}{\sim} (X, Y)$

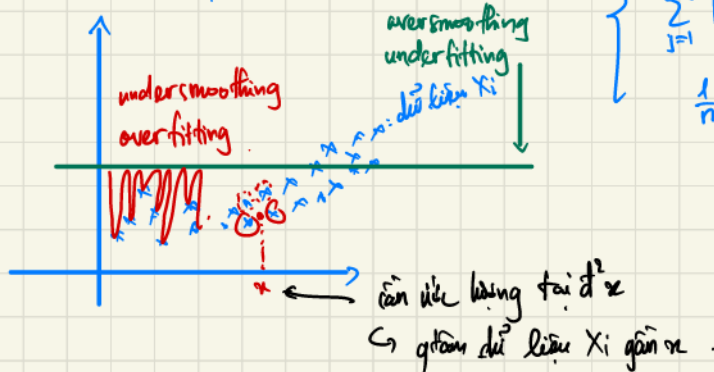
1) Definition

Nadaraya - Watson kernel estimator (1964).

$$\hat{r}(x) = \begin{cases} \frac{\sum_{i=1}^n y_i \cdot K\left(\frac{x_i - x}{h}\right)}{\sum_{i=1}^n K\left(\frac{x_i - x}{h}\right)} & \text{if } \sum_{i=1}^n K\left(\frac{x_i - x}{h}\right) > 0 \\ \bar{y} & \text{else.} \end{cases}$$

Remarks

$$1) \hat{r}(x) = \sum_{i=1}^n y_i w_i(x) \text{ with } w_i(x) = \begin{cases} \frac{K\left(\frac{x_i - x}{h}\right)}{\sum_{j=1}^n K\left(\frac{x_j - x}{h}\right)} & \text{if } \sum_{j=1}^n K\left(\frac{x_j - x}{h}\right) > 0 \\ \frac{1}{n} & \text{else.} \end{cases}$$



If h is great: the weight $w_i(x)$ are all great and similar and you get an estimator of $\bar{y}$.

h is small: the estimator is $\bar{Y}$ unless x is very closed to x_i , $i \leq n$
 and in that case, $w_i(x) \approx 1$ and $w_j(x) \approx 0, j \neq i$
 and hence $\hat{r}(x) \approx Y_i$

$$2) \hat{r}(x) = \frac{\hat{g}(x)}{\hat{f}(x)} \quad \text{where} \quad \hat{g}(x) = \frac{1}{nh} \sum_{i=1}^n Y_i K\left(\frac{x_i - x}{h}\right)$$

$$\hat{f}(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x_i - x}{h}\right)$$

whenever $\hat{f}(x) \neq 0$.

2) Asymptotic properties.

Hypotheses

a) r and f_x are C^2 on a neighborhood of x .

b) $f_x(x) > 0$

c) $\lim_{n \rightarrow +\infty} nh_n = +\infty$ and $\lim_{n \rightarrow +\infty} h_n = 0$

d) Let consider ϕ : $\phi(x) = E[Y^2|x]$

Assume that ϕ is continuous at point x .

e) K is bounded, integrable, positive, symmetric and with compact support.

Theorem 1 — If hypotheses (a) $\rightarrow$ (e) hold then:

$$MSE(\hat{r}(x)) E[(\hat{r}(x) - r(x))^2] = B^2(x) \cdot h^4 + V(x) \cdot \frac{1}{nh} + O\left(h^2 + \frac{1}{nh}\right)$$

$$\text{where } B(x) = \frac{1}{2f_x(x)} \mu_2(K) \cdot (g^{(2)}(x) - r(x) f^{(2)}(x))$$

$$V(x) = \frac{\phi(x) - r^2(x)}{f_x(x)} \cdot R(K)$$

• Complementary hypotheses:

$$\text{MISE}(\hat{r}) = \mathbb{E} \left[\int (\hat{r}(x) - r(x))^2 w(x) dx \right]$$

where w is a weight function, positive, bounded and with compact support C

a') r and f_x are C^2 on a neighborhood of C .

d') ϕ is continuous on a _____.

• Theorem 2 — If hypotheses (a'), (b), (c), (d'), (e), (f) hold then:

$$\text{MISE}(\hat{r}) = B^2 h^4 + V \cdot \frac{1}{nh} + O\left(\frac{1}{nh} + h^4\right)$$

where

$$B^2 = \int B^2(x) w(x) dx \quad \text{and} \quad V = \int V(x) w(x) dx.$$

• Other results:

• We could also get results (under technical assumptions) on almost sure convergence:

$$\hat{r}(x) - r(x) = O_{as}(h^2) + O_{as}\left(\sqrt{\frac{\ln n}{nh}}\right) \quad (3)$$

• If we assume that r and f_x are C^3 and K is a kernel of order p :

$$\hat{r}(x) - r(x) = \mathcal{O}_{as}(h^p) + \mathcal{O}_{as}\left(\sqrt{\frac{pnn}{nh}}\right) \quad (4)$$

Convergence rates

• If (4) holds : $h \sim c \left(\frac{n}{\ln(n)} \right)^{\frac{-1}{2p+1}}$

$$\hat{r}(x) - r(x) = \mathcal{O}_{as} \left[\left(\frac{n}{\ln(n)} \right)^{\frac{p}{2p+1}} \right]$$

• If condition of theorem 1 (resp. Theorem 2) hold, taking $h \sim Cn^{-4/5}$

$$\text{MSE}(\hat{r}(x)) = \mathcal{O}(n^{-4/5})$$

resp

$$\text{MRE}(\hat{r}) = \mathcal{O}(n^{-4/5})$$

3) Choice of the smoothing parameter in practice. (Choice of h).

• A popular method : cross-validation.

Introduce $CV(h) = \sum_{i=1}^n (y_i - \hat{r}^{-i}(x_i))^2 w(x_i)$

where $\hat{r}^{-i}$ is computed from all the data excepted (x_i, y_i)

$$h_{cv} = \arg \min_{h \in \mathcal{H}_n} CV(h)$$

Remark.

$$\begin{aligned} CV(h) &= \sum_{i=1}^n (r(x_i) + \varepsilon_i - \hat{r}^{-i}(x_i))^2 w(x_i) && \text{arg} \in \text{the giving MSE} \\ &= \sum_{i=1}^n \varepsilon_i^2 w(x_i) + \sum_{i=1}^n (r(x_i) - \hat{r}^{-i}(x_i))^2 w(x_i) \\ &\quad + 2 \sum_{i=1}^n \varepsilon_i (r(x_i) - \hat{r}^{-i}(x_i)) w(x_i) && \text{by } E \varepsilon_i = 0. \end{aligned}$$

$\varepsilon_i = y_i - r(x_i)$

• Property. If condition of Theorem 2 hold and H_n contains values fulfilling (c) then (with technical conditions)

$$\frac{\inf_{R \in H_n} \text{MISE}(R)}{\text{MISE}(R_{\text{opt}})} \xrightarrow[n \rightarrow +\infty]{\text{as}} 1.$$

• R programming: data (pressure)

```
t = ...
p = ...
plot(t, p)
sm. regression(t, p, method = 'w')
```

• Simulated data:

```
n = 1000
x = rexp(1/2)
y = 10 * exp(-3 * x) + rnorm(n, 0, exp(-x/2))
plot(function(x) 10 * exp(-3 * x), 0, 3)
lines(x, y, col = 'red', type = 'p')
res = lm. regression(x, y, method = 'w', add = T, col = 'blue', eval.points = sort(x), aicc)
```

↪ time x change từ 0 → 3.

↪ vẽ thêm đ

```
e = rnorm(n, 0, 1)
u = rnorm(n, 2, 0.5)
y = cos(2 * pi * x) + 30 * exp(-3 * x) + 0.1 * e.
plot(x, y)
```

names(res) : xem res gồm = btrên gì.

cv=0

for (i in 1:n)

{ cv = cv + $\frac{1}{n} \left(y[i] - \text{lm.regression}(x[i], y[i], \text{method} = 'cv', \text{display} = 'none', \text{eval.points} = x[i]) \$ estimate \right)^2$ }

abline(0,1) : vẽ đường $y = 1x + 0$

• Discrimination: (bản phân loại)

Y takes value in $\mathcal{Y}$

X real valued.

• Estimated Bayes classifier:

$$\hat{y}_0 = \arg \max_{y \in \mathcal{Y}} P(Y=y(x)) \Big|_{x=x_0}$$

$$= \arg \max_{y \in \mathcal{Y}} \hat{r}_y(x_0)$$

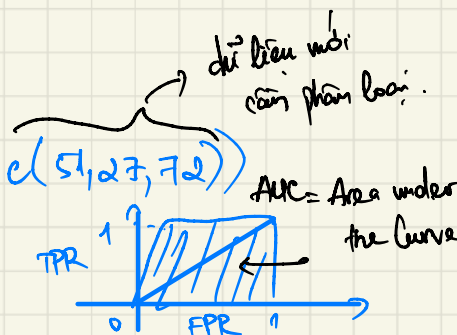
Reminder : Take $Z_y = \mathbb{1}_{\{Y=y\}}$

$$Z_y = \underbrace{r_y(x)} + \varepsilon_y.$$

$$= P(Y=y|X)$$

classif - NP (Perma, test, c(51, 27, 72))

	$\hat{y}$	-	+
$\hat{y}$	-	TN	FP
+	+	FN	TP



$$\hat{y} = \begin{cases} 1 & \text{if } \hat{P}(Y=y|x) \big|_{x=x_0} > \underbrace{0.5}_t \\ 0 & \text{else} \end{cases}$$

if $\hat{P}(Y=y|x) \big|_{x=x_0} > \underbrace{0.5}_t$
 else

AUC càng lớn thì phân loại càng tốt
 độ tin cậy phụ thuộc vào t

- $FPR(t) = \frac{FP(t)}{TN(t) + FP(t)}$
- $TPR(t) = \frac{TP(t)}{FN(t) + TP(t)}$

28.10.2022

Chapter 3.

From univariate, to multivariate and functional data.

1) Multivariate case.

$$X: \Omega \rightarrow \mathbb{R}^p$$

$$Y: \Omega \rightarrow \mathbb{R}$$

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

$$\hat{r}(x) = \begin{cases} \frac{\sum Y_i K\left(\frac{X_i - x}{h}\right)}{\sum K\left(\frac{X_i - x}{h}\right)} & \text{if } \sum K\left(\frac{X_i - x}{h}\right) \neq 0 \\ \bar{Y} & \text{else.} \end{cases}$$

Ex: 1) $\hat{K}(\vec{t}) = K_1(t_1) \cdot K_2(t_2) \dots K_p(t_p)$, where $K_i: \mathbb{R} \rightarrow \mathbb{R}$.

2) $K(\vec{t}) = K_0(\|\vec{t}\|)$ $K_0: \mathbb{R} \rightarrow \mathbb{R}$.

$$\underline{\text{vd:}} \quad K\left(\frac{X_i - x}{h}\right) = \prod_{j=1}^p K_j\left(\frac{X_{ij} - x_j}{h_j}\right)$$

$$K\left(\frac{X_i - x}{h}\right) = K_0\left(\frac{\|X_i - x\|}{h}\right)$$

We can get, under some conditions:

$$\text{USE}(\hat{r}(x)) = B^2(x) \cdot h^4 + V(x) \cdot \frac{1}{nh^p} + o\left(h^4 + \frac{1}{nh^p}\right)$$

$$\text{MISE}(\hat{r}) = B^2 \cdot h^4 + V \cdot \frac{1}{nh^p} + o\left(h^4 + \frac{1}{nh^p}\right)$$

Hence the optimal smoothing parameter :

$$h_{\text{opt}} = \left[\frac{V_p}{4B^2} \right]^{\frac{1}{p+4}} n^{-\frac{1}{p+4}}$$

Then optimal AMISE $\sim C \cdot n^{-\frac{4}{p+4}}$ and hence the cv rate $C \cdot n^{-2/p+4}$

For a.s convergence, the optimal rate $C \cdot \left(\frac{n}{\ln(n)} \right)^{-\frac{2}{4+p}}$

Important result. Stone 1981-1982.

if f_x and r are C^2 then the optimal convergence rate for any estimator:

1) $C \cdot n^{-\frac{k}{2k+p}}$ for the L_q cv $q > 0$.

2) $C \cdot \left(\frac{n}{\ln(n)} \right)^{-\frac{k}{2k+p}}$ for the L_∞ cv.

An explanation. the curse of dimensionality.

$n = 1\,000\,000$ values $\mathcal{U}([0,1]^{\otimes p})$

Take $h = 0.05$: $P(\max |x_i - u_i| \leq h) \leq \frac{0.1^p}{1} = 0.1^p$

Let N_p be the number of values in $B_h^p(x) = \{t, \max |t_i - u_i| \leq h\}$

$$N_p \sim \text{Bin}(n, p_h(x))$$

$$E[N_p] = n \cdot p_h(x) \leq 10^6 \times (0.1)^p = 10^{6-p}$$

for $p=1$

10^5

$p=2$

10^4

$p=3$

10^3

$p=4$

10^2

$$p=5$$

$$p=6$$

$$10'$$

$$1$$

. Remark The problem does not come from the estimator method but from the nature of the problem.

. We may consider semi-parametric model:

. Additive model:

$$r(x) = \mu + \sum_{i=1}^n r_i(x_i)$$

. Single index model:

$$r(x) = r(\langle x, \theta \rangle)$$

. Practice on R : faithful data.

ss' điểm chia.

$$e_0 = \text{seq}(\min(\text{eruptions}), \max(\text{eruptions}), \text{len} = 100)$$

$$w_0 = \text{seq}(\min(\text{waiting}), \max(\text{waiting}), \text{len} = 100)$$

$$\text{sm.density}(\text{cbind}(\text{eruptions}, \text{waiting}), \text{eval.points} = \text{cbind}(e_0, w_0), \text{eval.grid} = T, \text{structure} = 'common')$$

$$h_1 = h_2 = h : 'common'$$

$$h_1 \neq h_2 : 'separate'$$

$$h_1 = \sigma_1 h : 'scaled'$$

. Pirate data.

t
p
x
y

$$\text{sm.regression}(\text{cbind}(x, y), p, \text{structure} = 'common', \text{eval.points} = \text{cbind}(x, y), \text{method} = 'cv')$$

2. Functional case.

$X: \Omega \rightarrow (\mathcal{F}, d)$ d is a pseudo-distance ($2d^2 \neq \text{nhau có thể}$
is không cần $= 0$)

$$\hat{r}(x) = \begin{cases} \frac{\sum_{i=1}^n Y_i K(d(x_i, x)/h)}{\sum_{i=1}^n K(d(x_i, x)/h)} & \text{if } \sum K(d(x_i, x)/h) \neq 0 \\ \bar{Y} & \text{else.} \end{cases}$$

Remarks

1) for NP models, the curse of dimensionality cannot be avoided.

$$Y = r(X) + \varepsilon, E[\varepsilon|X] = 0$$

$$r(x) = E[Y|Z, d(X, Z) = 0] : \text{maybe } \neq r(\tilde{x})$$

2) A kind of semi-parametric model.

3) Idea: use the pseudo-distance as a tool to extract relevant information on Y from X

Example of pseudo-distance

1) Given the basis $\Psi_1, \dots, \Psi_k, \dots$: orthonormal basis.

$$X = \sum_{j=1}^{+\infty} \lambda_j \Psi_j$$

$$x = \sum_{j=1}^{+\infty} \alpha_j \Psi_j$$

$$d(x, x) = \sum_{i=1}^M (\lambda_i - \alpha_i)^2$$

2) Project on M first PCA scores and do the same.

3) L^2 distance on derivatives:

$$d(x, \alpha) = \int \left(\dot{x}_t^{(k)} - \dot{\alpha}_t^{(k)} \right)^2 dt$$