

7/3/2025

Supervised Learning

Y : εξαρτημένη μεταβλητή / response / outcome

Προβλέψεις για την Y με βάση προηγούμενα δεδομένα

Πλαίσιο

$X = (X_1, \dots, X_p)$ ανεξάρτητες μεταβλητές / predictors / features

Y = outcome

$H_X(x)$: περιθώρια κατανομή του X

$H_{Y|X}(y|x) = P(Y \leq y | X=x)$ δοσμ. cdf $Y|X=x$.

$$f(x) = E(Y|X=x) = \int_x y dH(y|x)$$

Πρόβλημα : Δεδομένου $X=x \rightarrow$ Πρόβλεψη για $Y=?$

Εστω $p(x) =$ Πρόβλεψη των Y για $X=x$.

Κριτήριο (συνάρτηση) (απώλειας) : $\overset{\text{"min"}}{\rightarrow} L(Y, p(x))$: ποινή αν πρόβλεψη : $p(x)$
πραγμ. τιμή : Y

ιδιότητα $L \geq 0 \quad \forall y, p(x)$

$$L(y, y) = 0 \quad \forall y.$$

$$L(y, c) > 0 \quad \text{αν } c \neq y$$

Δεν έχει νόημα το πρόβλημα $\min_c L(y, c)$

$$\min_c L(y, c) = 0 \quad \text{για } c=y$$

Θεωρούμε $\min_c E_{Y|X=x} (L(Y, c))$
($Y=x$)

Ποια L ?

① Y : overfitting (regression)

$$L(y, c) = (y - c)^2$$

$$\min_c E_{Y|X} [(Y - c)^2] =$$

$$= \min_c \left\{ E_{Y|X} (Y^2 - 2Yc + c^2) \right\} =$$

$$= \min_c \left\{ E(Y^2|X=x) - 2c E(Y|X=x) + c^2 \right\}$$

$$\min \text{ για } \boxed{c = E(Y|X=x) = f(x)}$$

$$\min EL = E_{Y|X=x} [(Y - f(x))^2]$$

$$= E_{Y|X=x} \left[(Y - E(Y|X=x))^2 \right] =$$

$$= \text{Var}(Y|X=x) \leftarrow \text{irreducible error}$$

Πρόβλημα $H(Y|X)$: άγνωστο

$\Rightarrow$ $f(x)$: άγνωστη συνάρτηση

Χρειάζονται δεδομένα

Δεδομένα $\left[\begin{array}{l} \text{Σύνολο Εκπαίδευσης} \\ \text{(training set)} \end{array} \right]$

ακέραιο δείγμα μεγέθους n :

i	x_1	x_2	$\dots$	x_p	y
1	x_{11}	x_{21}	$\dots$	x_{p1}	y_1
2	$\vdots$	$x_{\cdot 2}$			
$\vdots$	$\vdots$				
n	x_{1n}	x_{2n}	$\dots$	x_{pn}	y_n

$x_{\cdot n}$

$\Rightarrow$ εκτίμηση
 $\hat{f}(x)$

συνάρτηση πρόβλεψης

Για νέες παρατηρήσεις με $X=x_0$
η προβλεπτή θα είναι $\hat{y} = \hat{f}(x_0)$

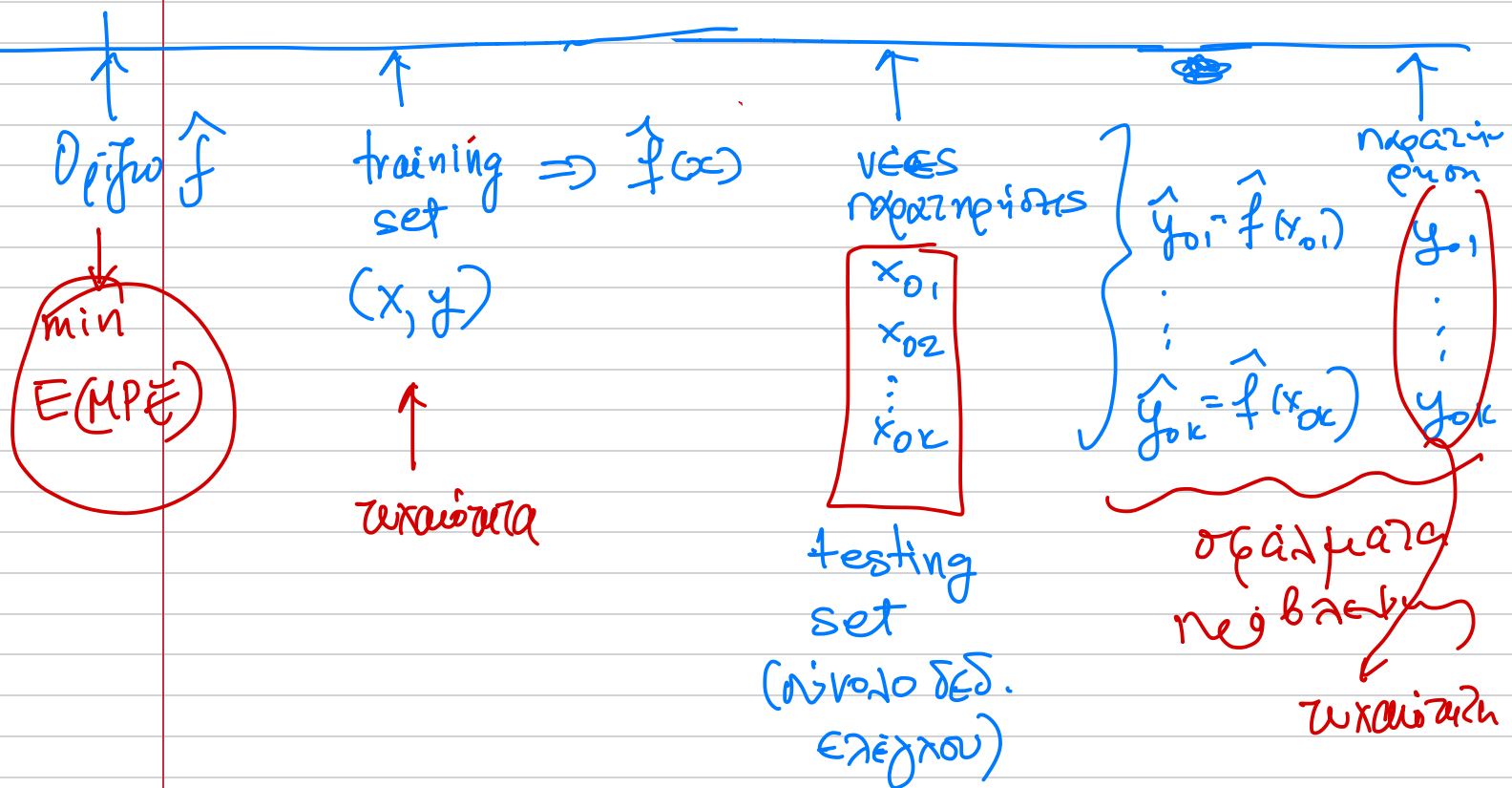
Πώς υπολογίζεται η $\hat{f}(x)$?

Ανάλυση με τη μέθοδο/μοντέλο που χρησιμοποιούμε (π.χ. regression, regularized regr, neural nets, random forests,...)

Ελαχιστοποίηση του Loss function απέναντι στο απείριστο

$$\min_{\hat{f}} E_{\gamma} [L(\gamma, \hat{f}(x))] = \text{MPE}$$

Mean Prediction Error



π.χ. Regression πα $P=1$ $X=X_1$

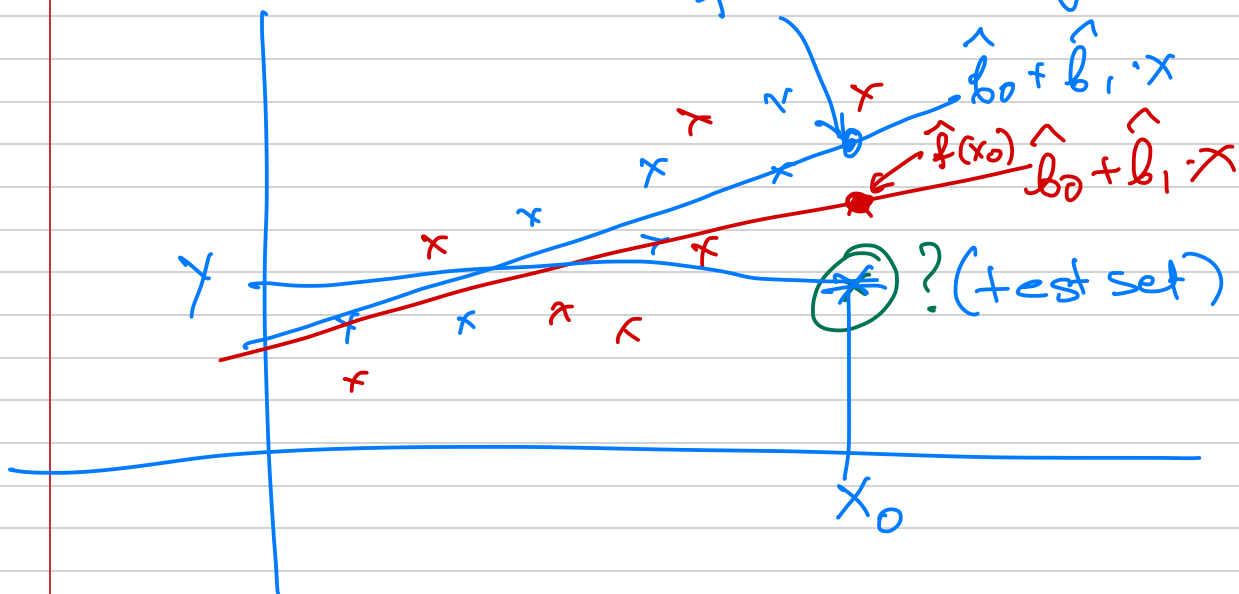
$$\hat{f}(x) = \hat{b}_0 + \hat{b}_1 \cdot x$$

$\hat{b}_0, \hat{b}_1$: συν.
ελατ. τερ.

ωχαίες μεταβλητές

εξαρτώνται από

$\hat{f}(x_0)$ training set



Μέτρο σφάλμα πρόβλεψης

$$MPE(\hat{f}) = E_{X,Y} \left[\left(Y - \hat{f}(X) \right)^2 \right]$$

Training set :

- (a) $(x_1, \dots, x_n)$: fixed σε δύο προτάσεις
σχεδιασμένες
(controlled studies)

π.χ. $\underline{x} = (x_1, x_2, \dots, x_n)$ where $x_1 = 25, x_2 = 25, x_3 = 25, x_4 = 25, x_5 = 25, x_6 = 30, x_7 = 30, \dots, x_n = y_n$

weights for training set: $w_j, j=1, \dots, n$

⑥ $(x_1, \dots, x_n)$ weights, $(y : weights)$

HELPS approximation
(observational studies)

$$MPE = E_{Y,X} \left((Y - \hat{f}(X))^2 \right)$$

$X, Y : \text{test set} !!$

training set $\Rightarrow \hat{f} : \text{weights}$

$\hat{f}(X) \rightarrow X : \text{weights (test set)}$

$\hat{f} : \text{weights (weights and training set)}$

Παράδειγμα

Εσω δατα (training) $n=20$

$$X = (x_1, \dots, x_n) \rightarrow \text{fixed } (\in (0, 10))$$

$$Y = (y_1, \dots, y_n)$$

$$Y = 2 + 3.6x + 1.5x^2 - 0.15x^3 + \varepsilon$$

$$f(x)$$

$$\varepsilon \sim N(0, \sigma^2), \sigma = 5$$

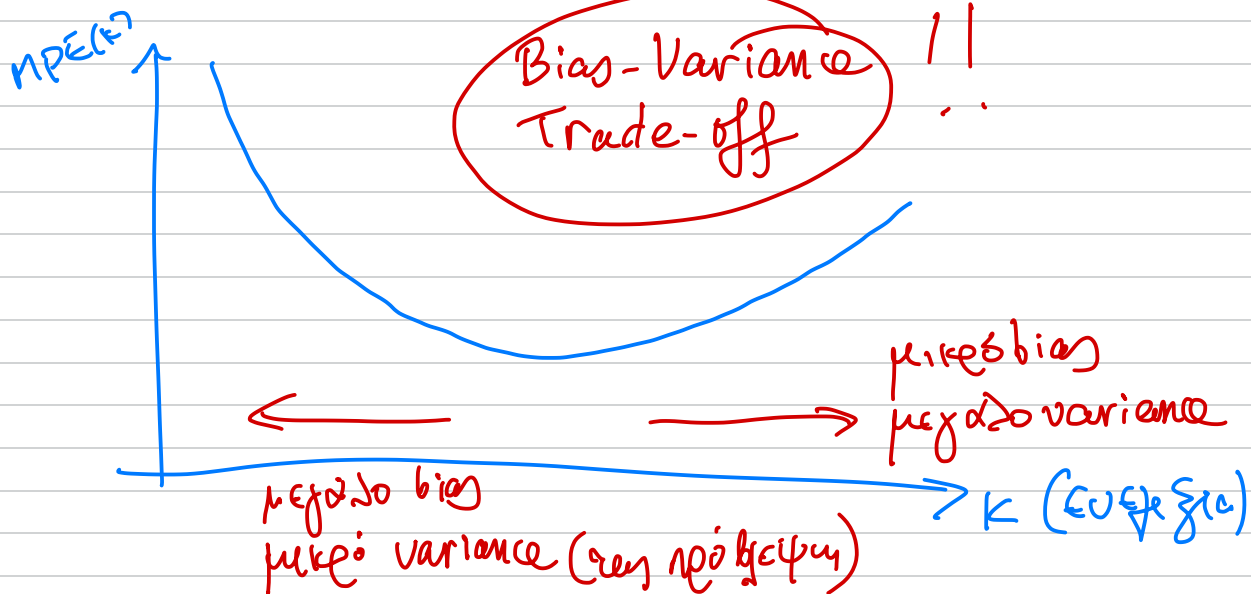
Συμμοίχεται
από την
μεθοδολογία
regression

Μέθοδος πρόβλεψης (Πολυωνομική Παμμεθόδευση
Ελαχίστων Τετραγώνων)

$$\hat{f}(x) = \hat{b}_0 + \hat{b}_1 x + \hat{b}_2 x^2 + \dots + \hat{b}_k x^k \quad (k: \text{τάξη μοντέλου})$$

$\hat{b}_0, \dots, \hat{b}_k$: εκτιμ. ελαχ. τετραγώνων

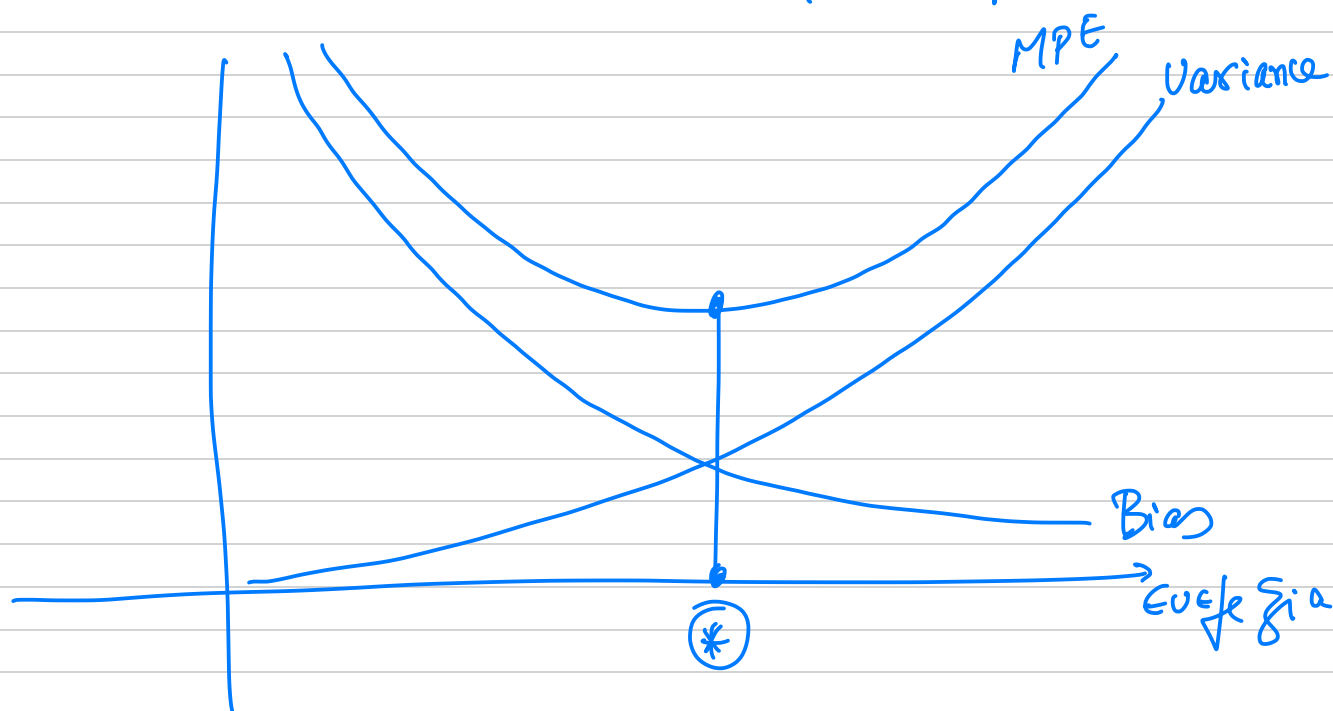
test : $x_0 = 6.9$



κ : ευελιξία του μοντέλου (flexibility)
 (αρ. παραμέτρων που εξετάζονται ($b_0, b_1, \dots, b_k$))

Σε παραμετρικά μοντέλα

πχ. $\hat{f}(x) = b_0 + b_1 x \leftarrow$ γραμμικό μοντέλο (b)
 $= b_0 + b_1 \sin(b_2 x)$ μη γραμμικό (b)



$MSE = \text{Bias} + \text{Variance}$

$$MSE = E_{x,y} \left(y - \hat{f}(x) \right)^2$$

x, y (test set)
 (νέες παρατηρήσεις)
 $\hat{f}$: από training set

πχ. Έστω $x = x_0 = \text{σταθερό}$

$$y \sim H(y | x_0) = E(y) = f(x_0)$$

$\hat{f}$: also training set $\leftarrow \begin{matrix} x_1, \dots, x_n : \text{συνθετά} \\ y_1, \dots, y_n \text{ και } y_i \sim H(\cdot | x_i) \end{matrix}$

$\hat{f}(x_0)$: και μεταβίβει

$\hat{f}(x_0)$ ανεξαρτητά $Y \sim H(Y|x_0)$
 $\downarrow$ $\downarrow$
 training set test

$$E(Y - \hat{f}(x_0))^2 =$$

$$= E_{tr, Y} \left(\underbrace{Y - \hat{f}(x_0)}_{H(Y|x_0)} + \underbrace{\hat{f}(x_0) - E(\hat{f}(x_0))}_{\text{απόσπασμα προφανώς μέση τιμή του } Y|X=x_0}}_{\text{απόσπασμα προφανώς μέση τιμή του } Y|X=x_0}} - \underbrace{\hat{f}(x_0)}_{\text{απόσπασμα προφανώς μέση τιμή του } Y|X=x_0}} \right)^2$$

$\underbrace{E(\hat{f}(x_0))}_{\downarrow \text{ (also training set)}}$ = μέση τιμή ως προβλεψή για $X=x_0$

$$\begin{aligned} \text{MPE} &= E \left((Y - \hat{f}(x_0)) - (\hat{f}(x_0) - E(\hat{f}(x_0))) + (\hat{f}(x_0) - E(\hat{f}(x_0))) \right)^2 \\ &= E \left[(Y - \hat{f}(x_0))^2 + (\hat{f}(x_0) - E(\hat{f}(x_0)))^2 + (\hat{f}(x_0) - E(\hat{f}(x_0)))^2 \right. \\ &\quad \left. - 2(Y - \hat{f}(x_0)) \cdot (\hat{f}(x_0) - E(\hat{f}(x_0))) \right. \\ &\quad \left. + 2(Y - \hat{f}(x_0))(\hat{f}(x_0) - E(\hat{f}(x_0))) \right] \end{aligned}$$

$$- 2 \left(\hat{f}(x_0) - E(\hat{f}(x_0)) \cdot (f(x_0) - E(\hat{f}(x_0))) \right)$$

Παρατήρηση

$$E \left[\underset{\uparrow}{(Y - f(x_0))} \cdot \underset{\uparrow}{(\hat{f}(x_0) - E(\hat{f}(x_0)))} \right]$$

$$= \underbrace{E(Y - f(x_0))}_0 \cdot \underbrace{E(\hat{f}(x_0) - E(\hat{f}(x_0)))}_0 = 0$$

$$E \left((Y - f(x_0)) (f(x_0) - E(\hat{f}(x_0))) \right)$$

$$= \left(f(x_0) - E(\hat{f}(x_0)) \right) \cdot \underbrace{E(Y - f(x_0))}_0 = 0$$

οπότε $E \left[\hat{f}(x_0) - E(\hat{f}(x_0)) \cdot \underbrace{(f(x_0) - E(\hat{f}(x_0)))}_{\text{σταθ.}} \right]$

$$= 0$$

$$\Rightarrow \text{MPE} = E \left((Y - f(x_0))^2 \right)$$

$$+ E \left[(\hat{f}(x_0) - E(\hat{f}(x_0)))^2 \right]$$

$$+ \left(E(\hat{f}(x_0)) - f(x_0) \right)^2$$

$$Y \sim H(y|x_0)$$

$$E(Y) = f(x_0)$$

$$= \text{Var}(y|X=x_0) \leftarrow (\text{irreducible error})$$

$$+ \text{Var}(\hat{f}(x_0))$$

$$+ \text{Bias}^2$$

$$\text{MPE} = \text{Bias}^2 + \text{Variance} + \text{Irred. Error}$$

!!

↓ ↗
 avg. $\hat{f}$