

4-4-2025

(Shrinkage methods)

Ridge Regression

$$\min \left\{ \underbrace{(y - X\beta)^T (y - X\beta)}_{\sum_{i=1}^n (y_i - \beta_0 - \dots - \beta_p x_{ip})^2} + \lambda \underbrace{\beta^T \beta}_{\sum_{j=1}^p \beta_j^2} \right\} \quad \beta = \begin{pmatrix} \beta_0 \\ \vdots \\ \beta_p \end{pmatrix}$$

$$\nabla_{\beta} = 0$$

$$\Rightarrow 2(X^T X) \beta - 2X^T y + 2\lambda \beta = 0 \quad \downarrow$$

$$\Rightarrow (X^T X + \lambda I) \beta = X^T y \Rightarrow \boxed{\hat{\beta}_{\text{ridge}} = (X^T X + \lambda I)^{-1} X^T y}$$

$$(\text{για } \lambda=0 : \hat{\beta} = (X^T X)^{-1} X^T y = \hat{\beta}_{\text{LSE}})$$

Παράδειγμα

$$\left. \begin{array}{l} x_1 = \text{εσοδμητα (1000 ευρώ)} \\ x_2 = \text{ανόσιαν (m)} \end{array} \right\} \Rightarrow \begin{array}{l} \hat{b}_1 \\ \hat{b}_2 \end{array}$$

$$\left. \begin{array}{l} \tilde{x}_1 (\text{σε ευρώ}) : \tilde{x}_1 = 1000 \cdot x_1 \\ \tilde{x}_2 (\text{σε cm}) : \tilde{x}_2 = 10^{-3} \cdot x_2 \end{array} \right\} \Rightarrow \begin{array}{l} \tilde{b}_1 = 10^{-3} \cdot \hat{b}_1 \\ \tilde{b}_2 = 10^3 \cdot \hat{b}_2 \end{array}$$

$$\hat{y}_j = \tilde{y}_j$$

Οι πραγματικοί μετ/οί των μεταβλητών επηρεάζουν τη σχετική μεταβολή των συντελεστών (διδ. τη σχετική σπινδαιότητα των μεταβλητών στο μοντέλο) σε αντίθεση με την λαμβάνειν. Εφακ. τετραγώνων.

Κανονικοποίηση των X_i πριν την εφαρμογή της μεθόδου

$$\tilde{X}_i = \frac{X_i - \bar{X}_i}{s(X_i)}$$

① Έστω $X_2 = X_1^2$, $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$

$$\tilde{X}_1 = \frac{X_1 - \bar{X}_1}{s(X_1)}, \quad \tilde{X}_2 = \frac{X_2 - \bar{X}_2}{s(X_2)}$$

$$\bar{X}_2 = \overline{X_1^2} \neq (\bar{X}_1)^2$$

② Έστω training set $\bar{X}_1 = 5$, $s(X_1) = 2$
test set $\bar{X}_1 = 5.5$, $s(X_1) = 1.5$

Training : Training set $\tilde{X}_1 = \frac{X_1 - 5}{2}$

Test set $\tilde{X}_1 = \frac{X_1 - 5.5}{1.5}$

Training : $\hat{\beta}_0, \hat{\beta}_1$: $\forall i \in \text{Training}$

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 \cdot \tilde{x}_{i1}$$

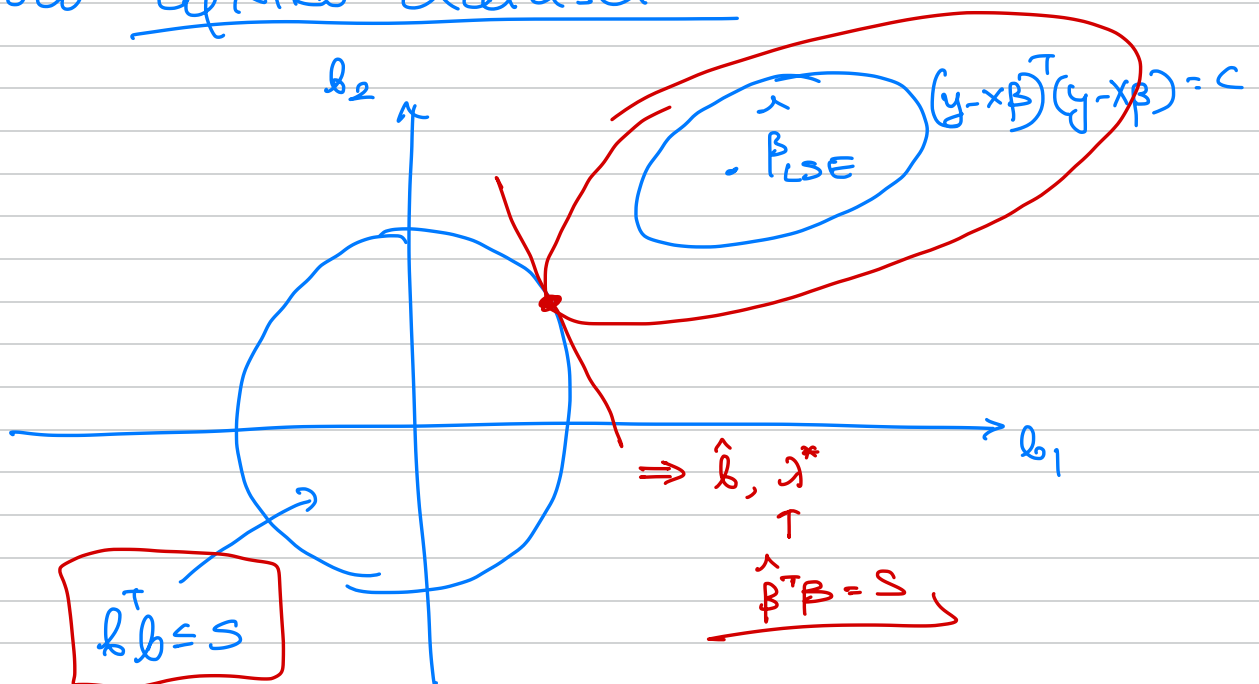
$$= \hat{\beta}_0 + \hat{\beta}_1 \cdot \frac{x_{i1} - 5}{2}$$

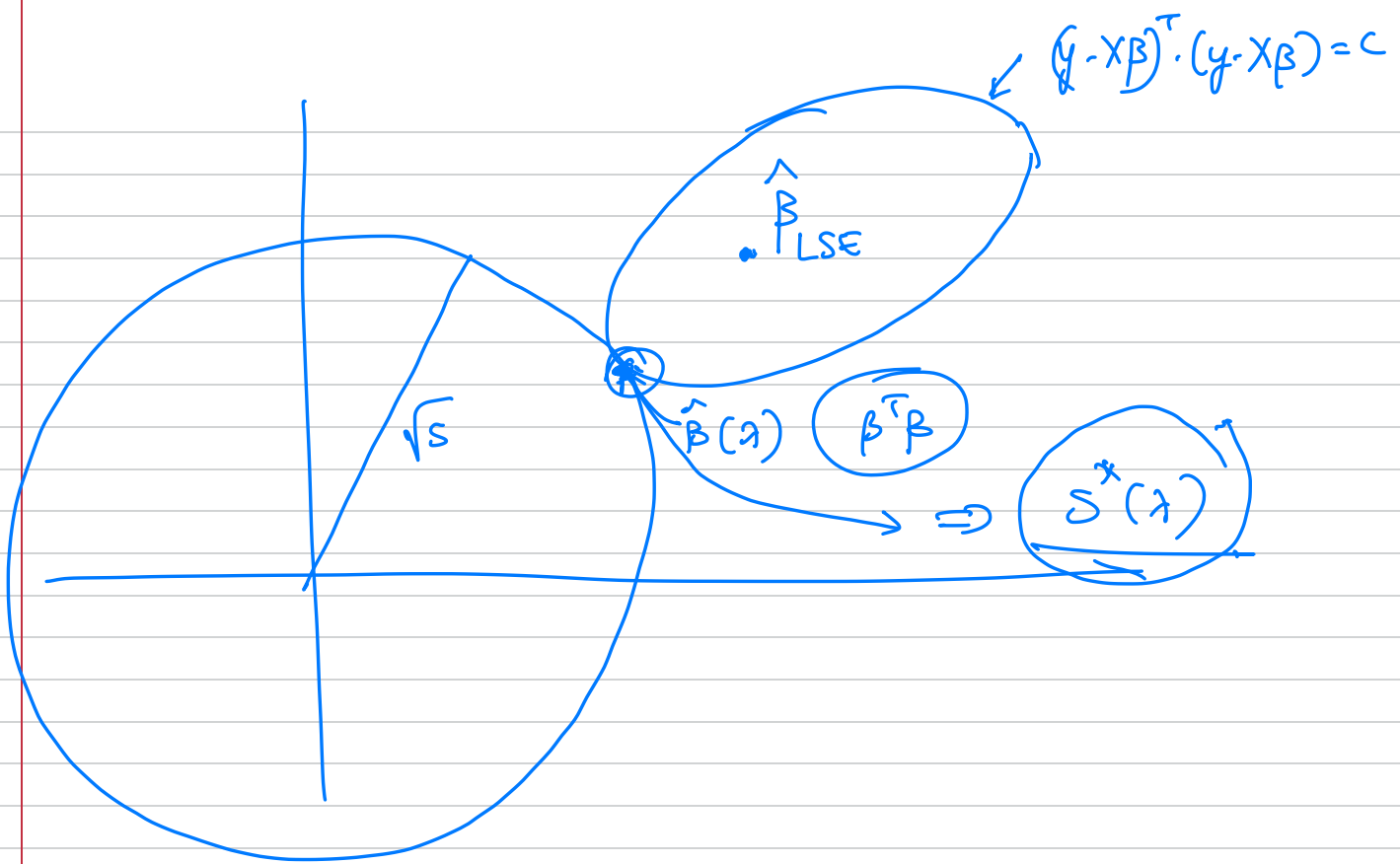
Αν στο testing set χρησιμοποιήσουμε

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 \underbrace{\tilde{x}_{i1}}_{\text{test}} = \hat{\beta}_0 + \hat{\beta}_1 \cdot \frac{x_{i1} - 5.5}{1.5} \neq$$

Είπαμε γάλας να χρησιμοποιήσουμε χωριστά
ωστούνουν στο test set !!

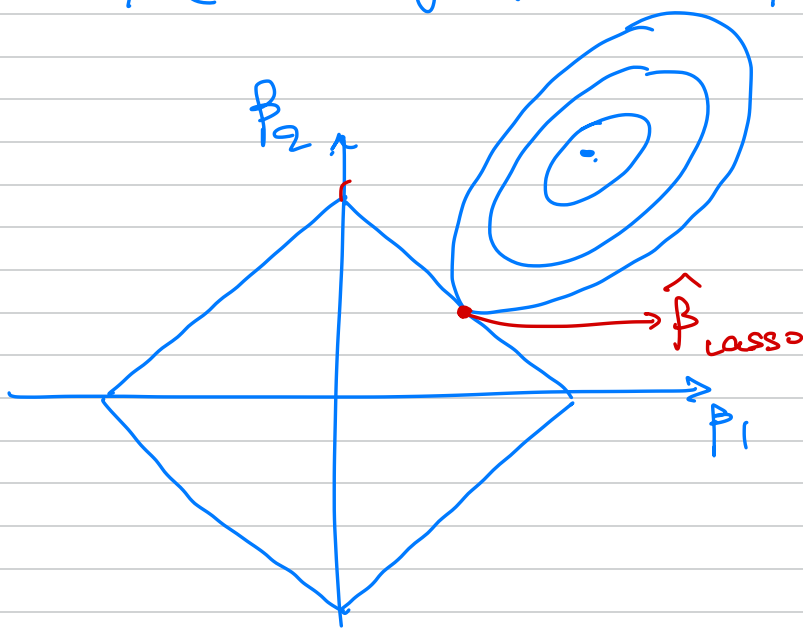
Ευαγγελισμός : Αν τα training + test
sets προέρχονται από ένα αρχικό dataset
n ωστούνουν μισή να γίνει εφ' αχίς
στο αρχικό dataset

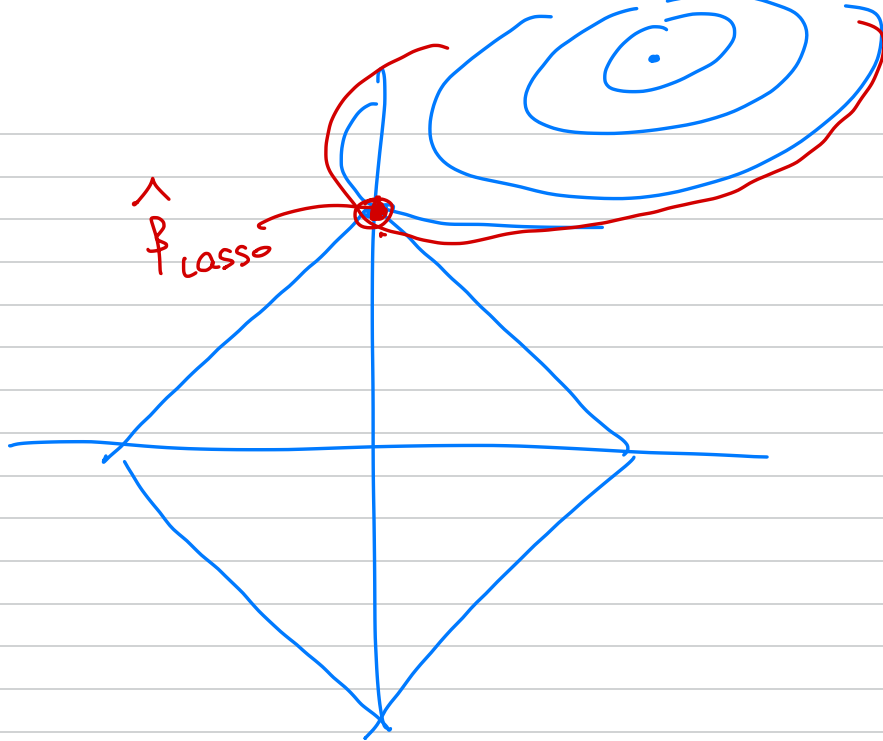




Lasso Regression

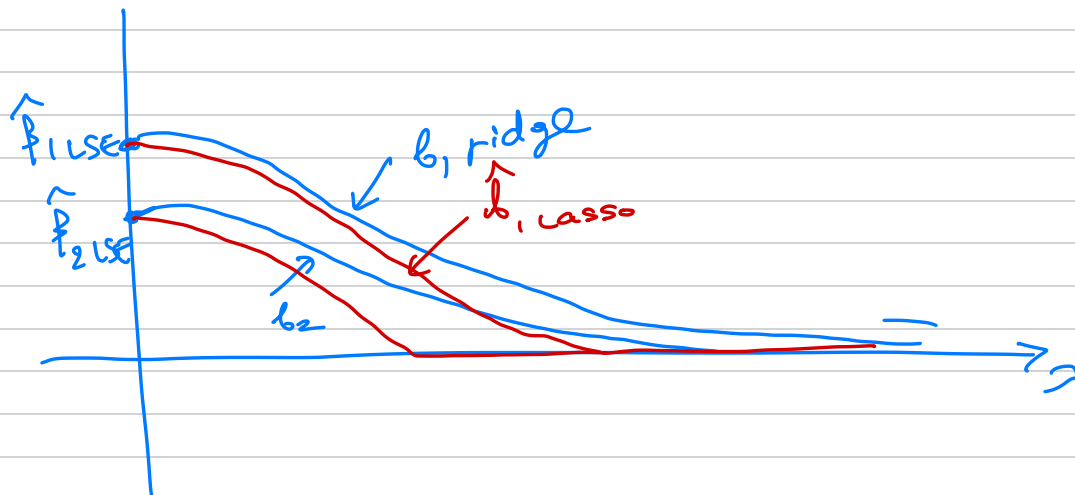
$$\min \left\{ (y - X\beta)^T (y - X\beta) : \|\beta\|_1 \leq s \right\} \quad (\text{quadratic programming})$$

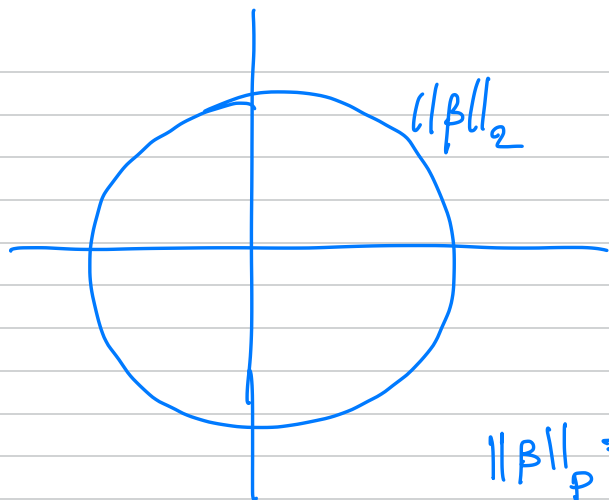
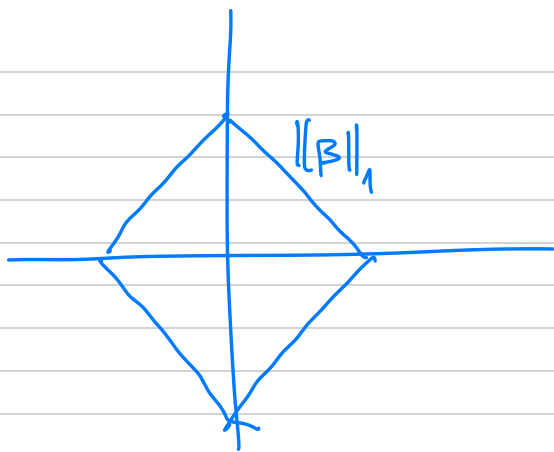




$$\min_{\beta} \left\{ (Y - X\beta)^T (Y - X\beta) + \lambda \|\beta\|_1 \right\}$$

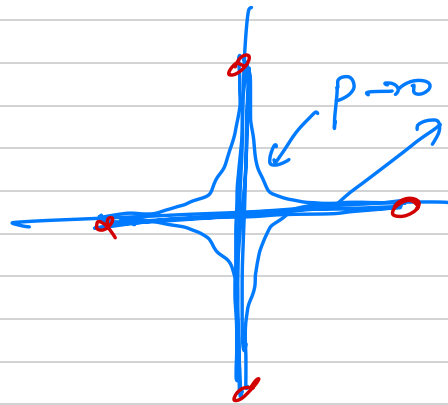
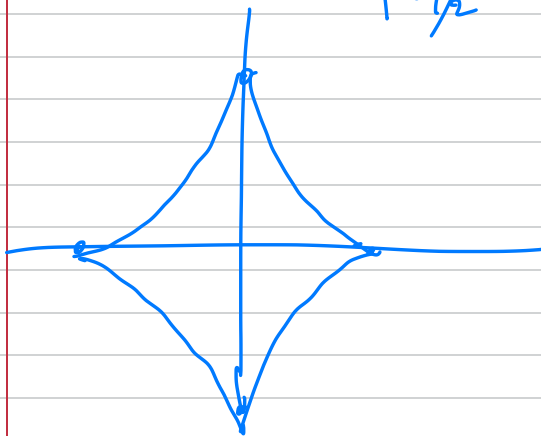
Ridge





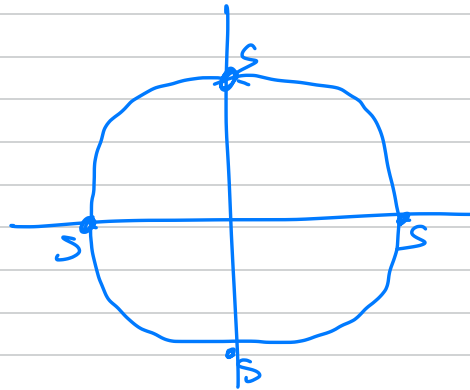
$$\|\beta\|_p = \left(\beta_1^p + \beta_2^p \right)^{1/p}$$

$$\|\beta\|_{1/2} = \left(\beta_1^{1/2} + \beta_2^{1/2} \right)^{1/2}$$

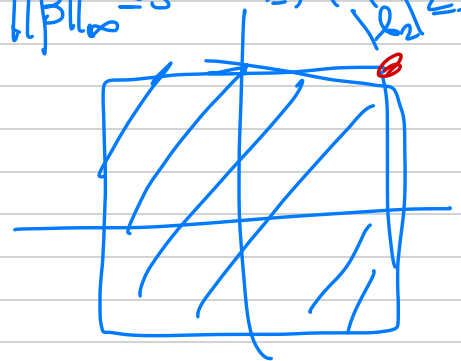


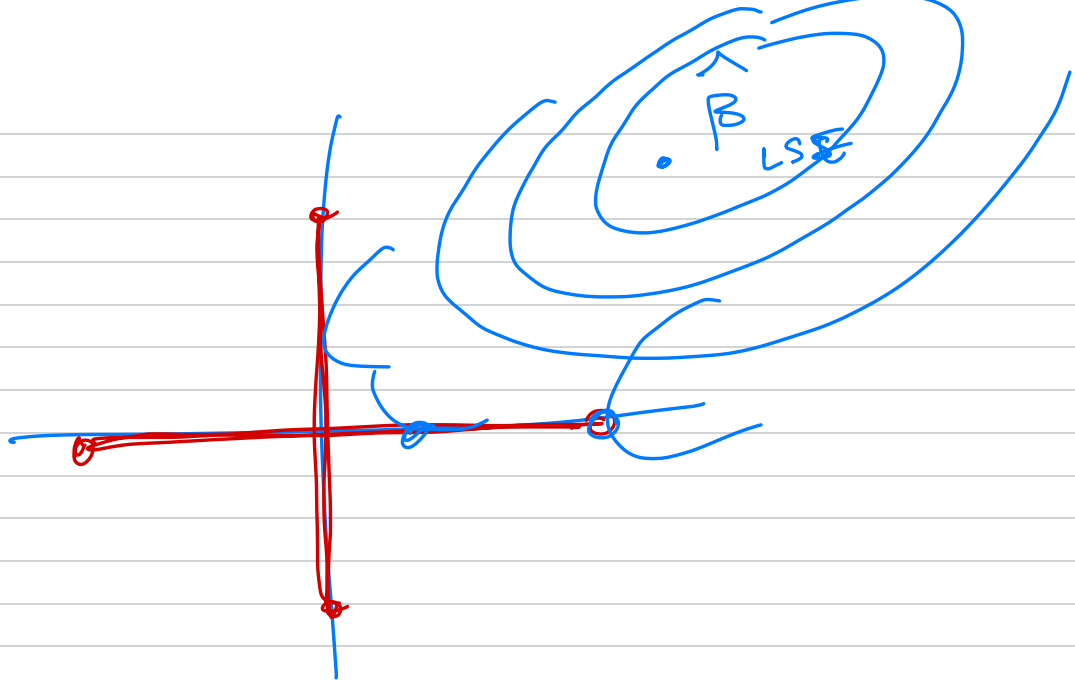
variable selection

$$\|\beta\|_3 \leq s$$



$$\|\beta\|_\infty \leq s \Rightarrow \max(|\beta_1|, |\beta_2|) \leq s \Rightarrow |\beta_1| \leq s, |\beta_2| \leq s$$





Elastic net ($\Rightarrow$ (glmnet) library R)

$$\min \left\{ (Y - X\beta)^T (Y - X\beta) + \alpha \lambda_1 \|\beta\|_1 + (1-\alpha) \lambda_2 \|\beta\|_2^2 \right\}$$

Ridge-Lasso εφαρμόζεται κ' για $p > n$

Μέθοδοι μείωσης διασποράς (του X)

$$Y = b_0 + b_1 X_1 + \dots + b_p X_p$$

(π.χ. $X_1, \dots, X_p$ βαθμοί σε 36 μαθήματα
αποφοίτου)

y = μισθός 5 ετών μετά την αποφοίτηση)

$X_1, \dots, X_p$ έχουν συσχέτιση

Αν θέλαμε να πάρουμε ένα σκέρ (συνάρτηση
των $X_1, \dots, X_p$)

που να ^(δίνεται) σχετίζεται με την y

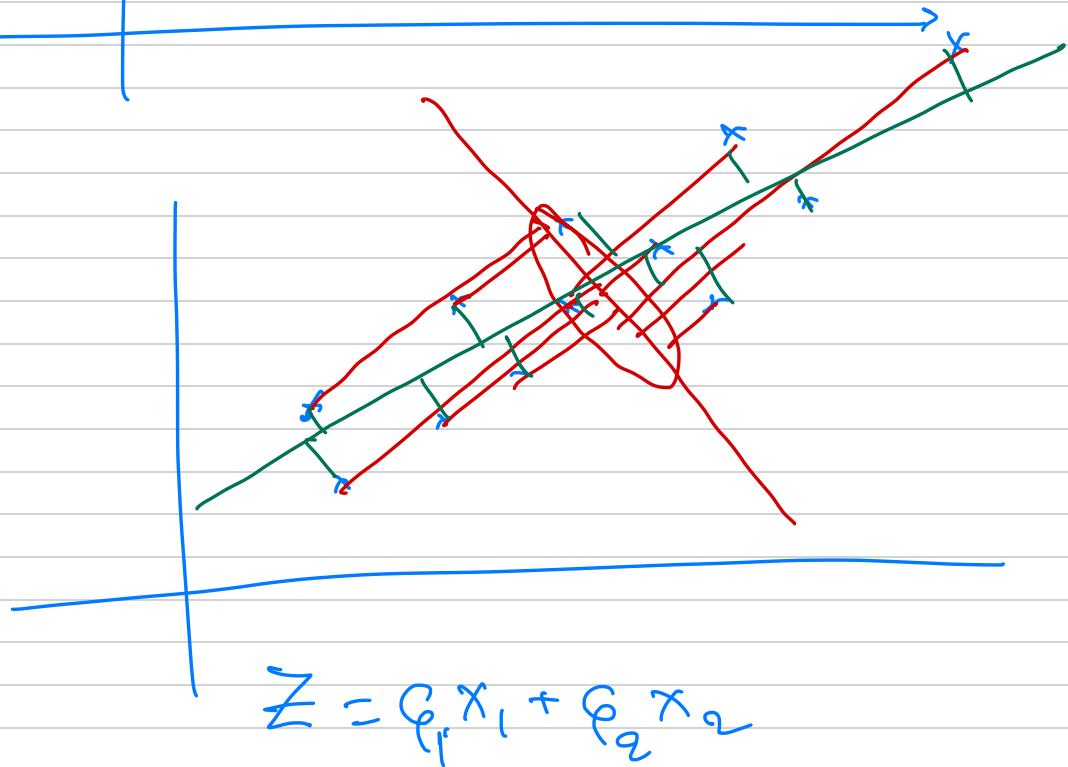
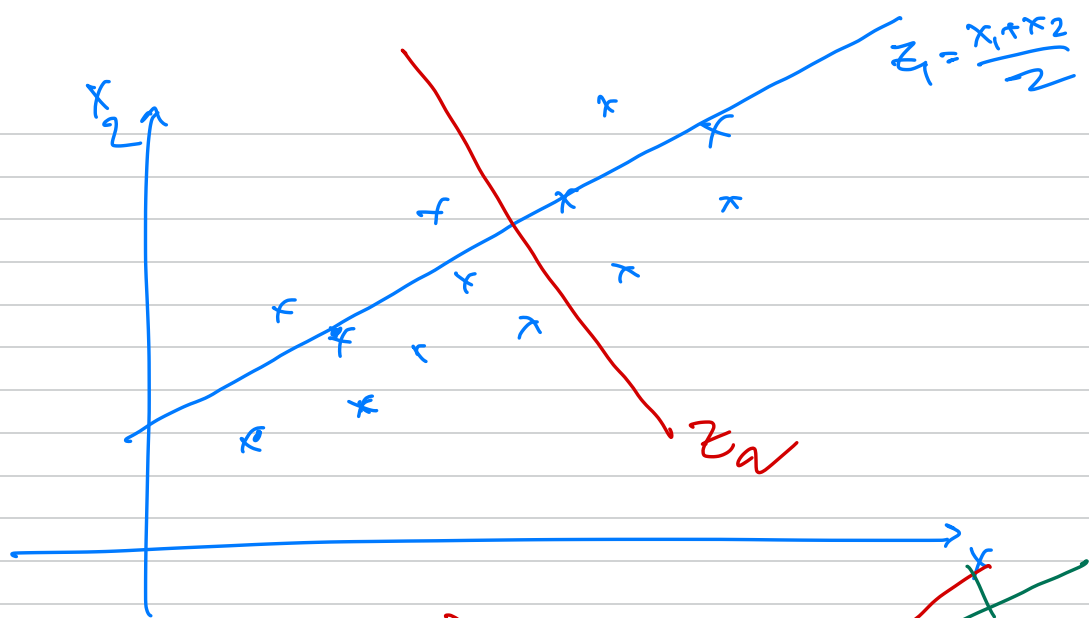
(π.χ. μέσος όρος μωχίου) = Z

$$Z = \varphi_1 X_1 + \dots + \varphi_p X_p$$

$$y = \theta_0 + \theta_1 \cdot Z$$

Αν προσελάμε κ' μια δεύτερη Z_2

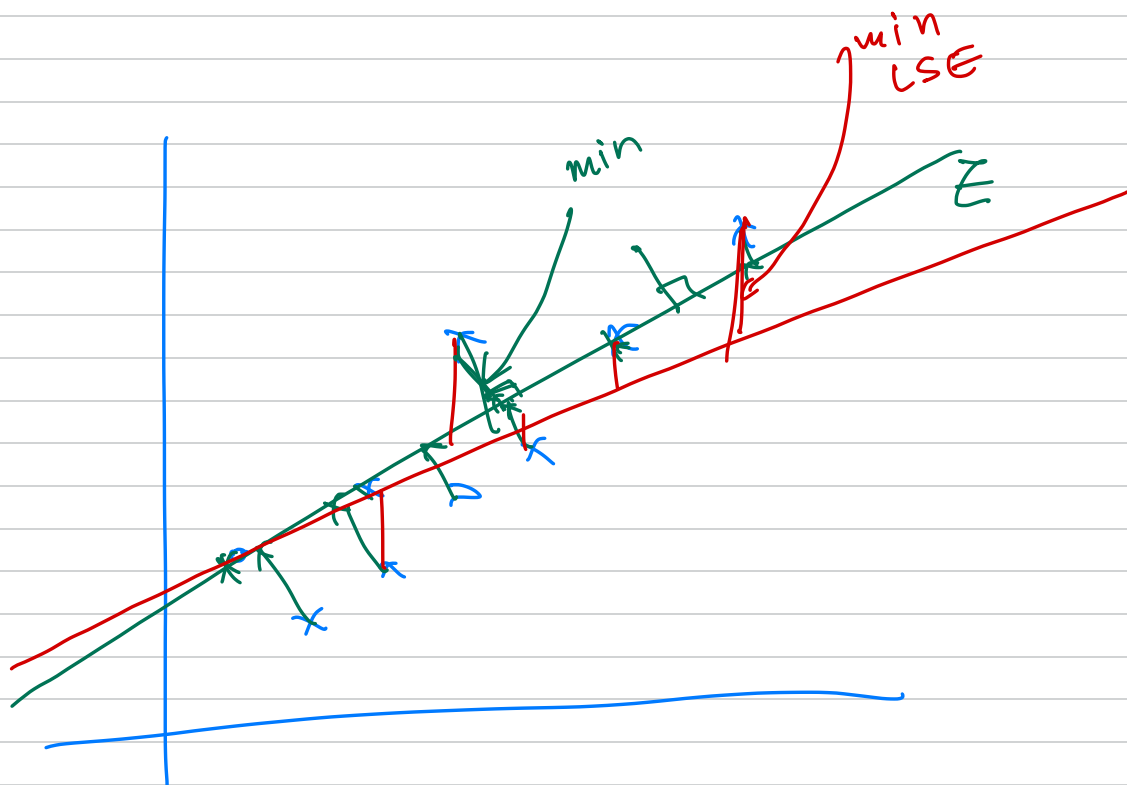
(θα επηρεάσει καίρια συσχέτιση με τα Z_1)



Ποια είναι τα c_1, c_2 που
δείχνουν το σωστό καίριο ρεόνο
τα $(x_1, x_2) \neq$ παραμένουν;

$$(x_1, x_2) \rightarrow z = c_1 x_1 + c_2 x_2$$

x_1	x_2	z
x_{11}	x_{12}	z_1
x_{21}	x_{22}	$\vdots$
$\vdots$		$\vdots$
x_{n1}	x_{n2}	z_n



Z : οι προβολές των παρατηρήσεων
 να έχουν τη μικρότερη δυνατή διασπορά



ελαχιστοποίηση των αποστάσεων (μέσω προβολή)
 από των εγδοία

↓ ↓

I^n κύρια συνιστώσα εν $(X_1, X_2, \dots, X_p)$

$$Z_1, Z_2, \dots, Z_M \quad (M \leq p)$$

$$Z_m = \sum_{j=1}^p \phi_{jm} X_j \quad m=1, \dots, M$$

ϕ_{jm} : loadings

① $y = b_0 + b_1 X_1 + \dots + b_p X_p \quad \leftarrow$ αρχικό
 $= b_0 + \sum_{j=1}^p b_j X_j$

② $y = \theta_0 + \theta_1 Z_1 + \dots + \theta_M Z_M \quad \leftarrow$ μεωμένη
διαταξη

$$= \theta_0 + \sum_{m=1}^M \theta_m Z_m$$

$$= \theta_0 + \sum_{m=1}^M \theta_m \sum_{j=1}^p \phi_{jm} X_j =$$

$$= \theta_0 + \sum_{j=1}^p \underbrace{\sum_{m=1}^M \theta_m \phi_{jm}}_{b_j} X_j$$

$$b_j = \sum_{m=1}^M \theta_m \phi_{jm} \quad \forall j=1, \dots, p$$

(2) $\Leftrightarrow$

$$y = b_0 + b_1 x_1 + \dots + b_p x_p$$

v.d.

$$b_j = \sum_{m=1}^n \theta_m$$

$$\phi_{jm}$$

$\forall j=1, \dots, p$

$\uparrow$
doorwa