

$$y = X\beta + \varepsilon(1) \quad \varepsilon \sim \mathcal{N}(0, \sigma^2 I_n), \quad X \in \mathbb{R}^{n \times (m+1)}, \quad \beta \in \mathbb{R}^{m+1}, \quad y \in \mathbb{R}^n$$

$\downarrow$
sparse

Εύρεση ενεργών μεταβλητών:

$$\hat{\beta} = \arg \min_{\beta} \left\{ \frac{1}{p} \|X\beta - y\|_p^p + \lambda \frac{1}{q} \|\beta\|_q^q \right\}, \quad 0 < p, q < 2$$

(2)

$\lambda > 0$

$\|X\beta - y\|_p \rightarrow$ error term

$\|\beta\|_q \rightarrow$ regularization term

Αρα β sparse δέλω q : να ενισχύει το sparsity.

p : $\|X\beta - y\|_p \rightsquigarrow$ measures the residual

▷ Επιλογή p : ανάλογα με τον τύπο του θορύβου

• Λευός Γκαουσιανός $\rightarrow p=2$

• salt and pepper $\rightarrow 0 < p < 1$ πιο robust
to outliers

▷ Επιλογή q : επιδέχεται ώστε να ενισχυθεί το sparsity

• $q=1$ προσέγγιση ως ℓ_0 -νόρμα απ' την ℓ_1 (η ℓ_0 νόρμα μετράει το πλήθος των μη μηδενικών εσόδων αλλά: minimization ως ℓ_0 -νόρμα διασύνταξης $\rightarrow$ δύσκολο $\Rightarrow$ προσέγγιση ως ℓ_0 -νόρμα απ' την ℓ_1).

• $0 < q < 1$ προσεγγίζουν την ℓ_0 -νόρμα καλύτερα: The smaller q the better approximation of ℓ_0 -norm by ℓ_q -norm.

Όμως $0 < q < 1 \rightarrow$ non convex

• Το λ : ελέγχει το trade-off data fitting and sparsity

▷ Επιλογή λ : • μικρό $\lambda \rightarrow$ οδηγεί σε dense solution αλλά

• μεγάλο $\lambda \rightarrow$ sparse solution αλλά μπορεί να μην έχω ακριβή ανάκρουση των εσόδων y .

Χρησιμοποιούμε το μοντέλο παλινφύσεως (2) για να λύσουμε το πρόβλημα (1).

Σκοπός: έρευνα ενεργών μεταβλητών μέσω των μη μηδενικών εσόδων του θ που παίρνουμε απ' τη λύση του $\arg \min_{\theta} \left\{ \frac{1}{p} \|X\theta - y\|_p^p + \lambda \frac{1}{q} \|\theta\|_q^q \right\}$

▷ Στο Real Data model (it has been studied in almost every paper that deals with analysis of SSDs):

- Επειδή ο δόρυβος δεν είναι Gaussian $\rightarrow$ Το $p=2$ όχι καλή επιλογή
- Με βάση ερευνητικές μεθόδους: $p < 1$ και πιο συγκεκριμένα: $p=0.8$
- Για ενίσχυση sparsity: $q=0.1$

Το $p \rightarrow$ data fidelity

Το $q \rightarrow$ sparsity

Το $\lambda \rightarrow$ regularization weight $\rightarrow$ ισορροπεί (i) το error $X\theta - y$
(ii) το sparsity $\|\theta\|_q$

$y = X\theta + \varepsilon$: Το θ_j δείχνει: (i) πόσο συσβάλλει η στήλη X_j στο y

(ii) αν συσβάλλει θετικά ή αρνητικά

(iii) την ένταση / των οποίων συσβάλλει.

$\theta_j \leq 0 \Rightarrow$ η X_j δεν επηρεάζει το $y \Rightarrow$ μη ενεργός παράγοντας

$\theta_j \neq 0 \Rightarrow$ η X_j είναι εμφανής $\Rightarrow$ ενεργός παράγοντας

$X \rightarrow$ στήλες $\gg$ γραμμές $\rightarrow$ το θ είναι η πιο αραιή

λύση που εξηγεί σωστά τα δεδομένα

Ευρίσκω p, q για περάματα: $p = \{0.8, 1, 1.5, 2\}$

$q = \{0.1, 0.3, 0.5, 0.8, 1\}$

Παράδειγμα

Έστω $X = \begin{bmatrix} 1 & -1 & 1 \\ -1 & 1 & -1 \\ 1 & -1 & -1 \end{bmatrix}, y = \begin{bmatrix} 2 \\ 1 \\ 1 \end{bmatrix}$

και τρέχω τη μέθοδο l_p-l_q για $\lambda = [0.01, 0.1, 1, 10]$
και έστω ότι πήρα τα παρακάτω αποτελέσματα (υποθέτουμε
SEN TO ΕΡΕΣΕ)

λ	b_1	b_2	b_3
0.01	1.95	0.98	0.05
0.1	1.90	0.95	0.01
1	1.70	0.85	0.00
10	1.20	0.60	0.00

Παρατηρήσεις: $b_1, b_2 \rightarrow$ μεγαίνει για $\lambda = 0.01, 0.1, 1$
 $b_3 \rightarrow$ μικρό καθώς $\lambda \uparrow$

Θέλω: sparse $b \Rightarrow$ οι μη εμφανείς συντελεστές να $\rightarrow 0$
και οι εμφανείς (b_1, b_2) να παραμένουν αρκετά μεγάλες

Άρα $\lambda = 0.1$ φαίνεται καλή επιλογή

($\lambda = 0.01 \rightarrow$ όλα τα $b_i \neq 0 \rightarrow$ παραμένει όλο το

$\lambda = 1 \rightarrow$ οι b_1, b_2 αρχίζουν να $\downarrow$

$\lambda = 10 \rightarrow$ οι b_1, b_2 μειώνονται πολύ $\Rightarrow$ χάνεται πληροφορία)

- Μικρό λ καλό για αριθμ. fitting αλλά μπορεί να μη δώσει sparse b
- Μεγάλο λ καλό για επιλογή ανεξαρτητών αλλά μπορεί να χάνει πληροφορία