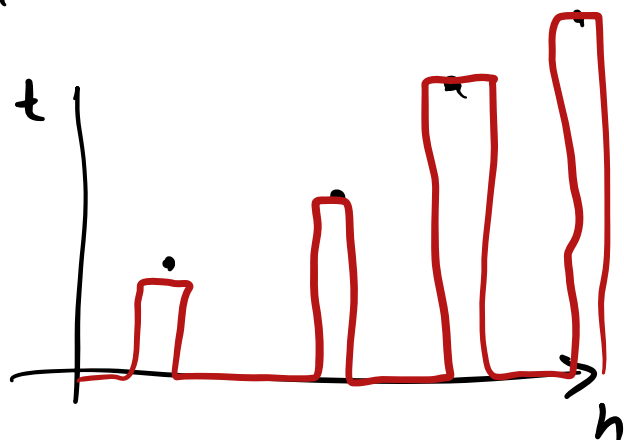


Χρήσιμος πηγάδι



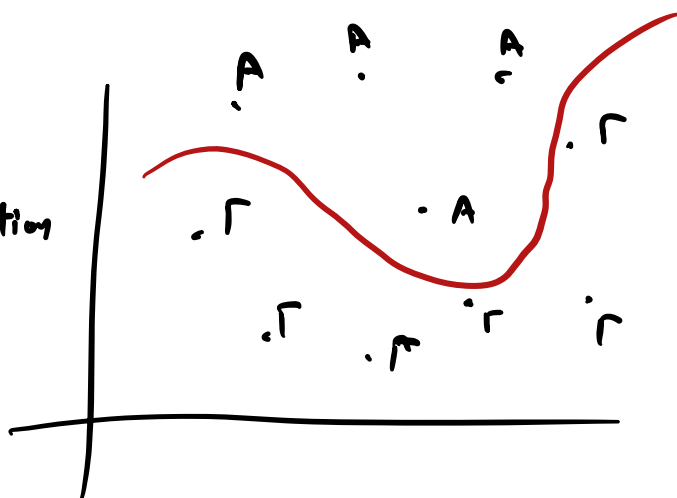
Regression

Φωτογραφίες

Αντρας / Γυναίκα

0.7	0.1
0.3	0.9

Classification



Μοντελοποίηση

$$(x, y) \sim D$$

features

$$x \in \mathcal{X} \quad (\text{domain})$$

$$y \in \mathcal{Y}$$

Στόχος είναι να μάθουμε συνάρτηση

$h: X \rightarrow Y$ ώστε να ελαχιστοποιείται
το $E[\ell(h(x), y)]$
 $(x, y) \sim D$

Για κάποια συνάρτηση κόστους ℓ

- Παραδείγματα:

Για regression

$$L(\hat{y}, y) = (\hat{y} - y)^2$$

Για classification

$$L(\hat{y}, y) = 1_{\{\hat{y} \neq y\}} = \begin{cases} 1 & \text{αν } \hat{y} \neq y \\ 0 & \text{αλλιώς} \end{cases}$$

Training set

$$S = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$$

$$\text{όπου } (x_i, y_i) \sim D$$

Κλάση υποθέσεων

H

$$\rightarrow \min_{h \in H} \underbrace{E_{(x,y) \sim D} [\rho(h(x), y)]}_{L_D(h)}$$

Empirical Risk Minimization

$$\min_{h \in H} \underbrace{E_{(x,y) \sim S} [\rho(h(x), y)]}_{L_S(h)} \rightarrow h_S \quad \text{ERM}$$

$$h(x) = \begin{cases} y & \text{av unäppl} \\ & (x,y) \in S \\ 0 & \text{alltids} \end{cases}$$

overfitting

Ανaly περίπτωση

$$L_D(h^*) = 0$$
$$L_D(h) \in [0, 1]$$

$$l(h(x), y) = 1[h(x) \neq y]$$

Θέλω να βρω $h \in H$

$$\Pr_{(x, y) \sim D} [h(x) \neq y] \leq \epsilon$$

πραγματικό loss

Έστω H_B όλες οι συναρτήσεις με

$$\Pr_{(x, y) \sim D} [h(x) \neq y] > \epsilon$$

Για κάθε $h \in H_B$

H πιθανότητα να μην την έχω
απορρίψει μετά από m δείγματα

είναι το πολύ $(1-\epsilon)^m \leq e^{-m\epsilon}$

Αν $m = \frac{\log(1/\delta)}{\epsilon}$ απορρίπτω την
κακή h με πιθανότητα $\geq 1-\delta$

Αν το H είναι πεπερασμένο

$$\Pr[\exists h \in H_B \text{ που επιβιώνει}]$$

$$\leq \sum_{h \in H_B} \Pr[h \text{ επιβιώνει}]$$

$$\leq |H_B| \delta \leq \underbrace{|H| \delta}_{\delta'}$$

$$m = \frac{\log(|H|/\delta')}{\epsilon}$$

τότε η $h_S \in H \setminus H_B$
με πιθανότητα $1 - \delta'$

Γενίκευση με ϵ Uniform Convergence

Στόχος είναι να μάθουμε καλή συνάρτηση:

$$L_D(h_S) - L_D(h^*) \leq \epsilon$$

Uniform Convergence

$$\max_{h \in H} |L_S(h) - L_D(h)| \leq \epsilon$$

$\equiv$ ερουμε

$$\mathbb{E}_S [L_S(h)] = L_D(h) \quad \forall h$$

Γιατί είναι καλό το Uniform
Convergence;

$$\begin{aligned} L_D(h_S) - L_D(h^*) &= L_D(h_S) - L_S(h_S) \\ &\quad + L_S(h_S) - L_S(h^*) \\ &\quad + L_S(h^*) - L_D(h^*) \end{aligned}$$

$$\leq |L_D(h_S) - L_S(h_S)|$$

$$+ \underbrace{L_S(h_S) - L_S(h^*)}_{\leq 0}$$

$$+ |L_S(h^*) - L_D(h^*)|$$

$$\leq 2 \max_{h \in H} |L_D(h) - L_S(h)| \leq 2\varepsilon$$

Περαστική είναι η

Με m δείγματα

Για κάθε $h \in H$

$$\Pr[|L_S(h) - L_D(h)| > \varepsilon] \leq e^{-m\varepsilon^2}$$

(And Hoeffding inequality)

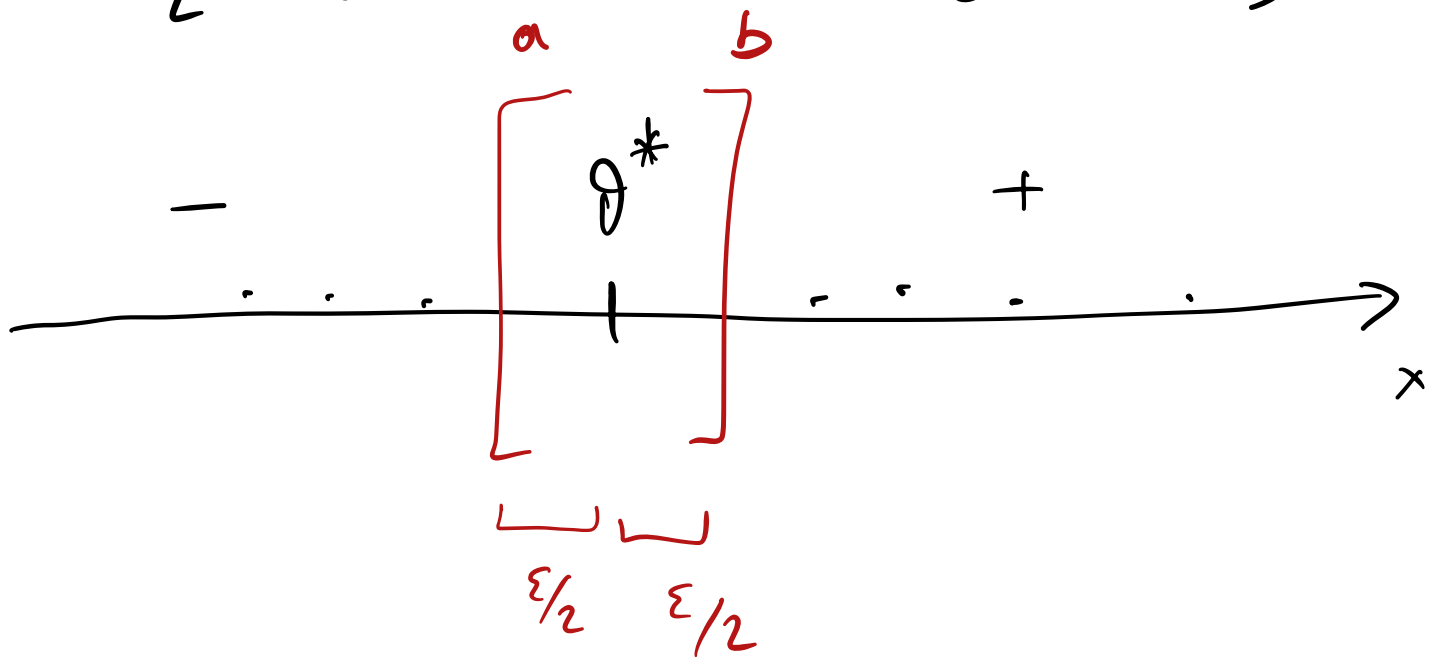
$$\Pr[\exists h \in H : |L_S(h) - L_D(h)| > \varepsilon]$$

$$\leq |H| e^{-m\varepsilon^2} \quad \text{αν} \quad m = \frac{\log(|H|/\delta)}{\varepsilon^2}$$

$$\leq \delta$$

Threshold Euclidean

$$H = \{ h_\theta : h_\theta(x) = \text{sign}(x - \theta) \}$$



a t.w.

b t.w.

$$\Pr [[a, \theta^*]] = \frac{\epsilon}{2}$$

$$\Pr [[\theta^*, b]] = \frac{\epsilon}{2}$$

$$\Pr [x_1, \dots, x_m \notin [a, \theta^*]] = \left(1 - \frac{\epsilon}{2}\right)^m \leq e^{-m\epsilon/2} \leq \delta/2$$

$$m = \frac{2}{\epsilon} \log(2/\delta)$$

VC-dimension

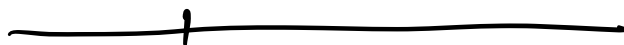
Shattering: Ένα σύνολο σημείων X είναι shattered από μια κλάση H αν για κάθε labelling γ του X υπάρχει $h \in H$ που να δίνει αυτό το labelling
 S.A.S $\forall i \quad h(x_i) = \gamma_i$

VC-dimension: Το μεγαλύτερο ^{μέγεθος ενός} συνόλου σημείων που είναι shattered από την H .

Παραδείγματα :

Thresholds

+



2 σημεία

+	+	✓
-	-	✓
-	+	✓
+	-	x



VC-dim = 1

Thresholds μc $\varphi o p a$
 η intervals

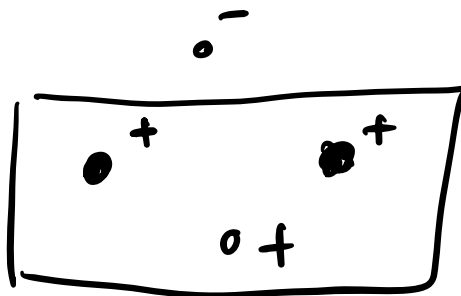
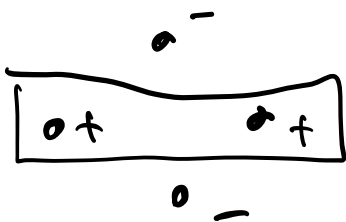


2 $\sigma \eta \mu \epsilon \iota \alpha \checkmark$
 3 $\sigma \eta \mu \epsilon \iota \alpha \times$

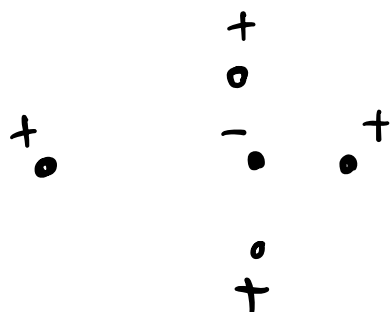


$VC-dim = 2$

Ορθογώνια (Axis-Aligned)



4 $\sigma \eta \mu \epsilon \iota \alpha \checkmark$



5 $\sigma \eta \mu \epsilon \iota \alpha \times$
 $VC-dim = 4$

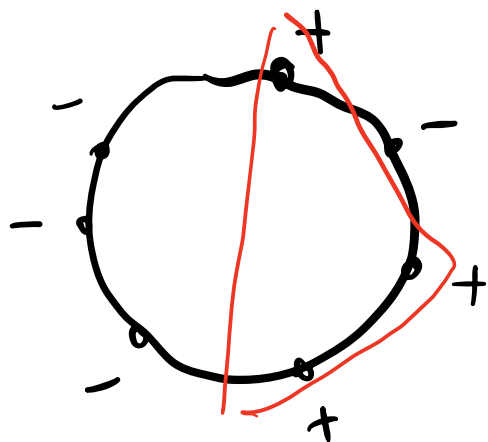
Βασικό Θεώρημα του Statistical Learning
Αν έχουμε μια κλάση H με $VC\text{-dim} = d$

Τότε για binary Classification
χρειαζόμαστε

$$m = \Theta\left(\frac{d + \log(1/\delta)}{\epsilon^2}\right)$$

δείγματα για uniform convergence.

$H = \{ \text{Όλα τα κύρια υποσύνολα του } \mathbb{R}^2 \}$



$$VC\text{-dim}(H) = \infty$$

Αν H έχει μέγεθος K
τότε

$$VC\text{-dim}(H) \leq \log K$$

Για Realizable Learning

$$m \leq O\left(\frac{d \log(1/\epsilon) + \log(1/\delta)}{\epsilon}\right)$$