

STAT 391

5/25/23

Lecture 18

logistic regression

Regression examples

Clustering 2 lectures

HW 6 optional

Last HW deadline
Friday 5 pm ← ^{solutions} 5,6

T.B. Posted
- notes on linear
 & logistic regn.
- slides on
 Clustering
 + b. edited

logistic Regression

used for: Classification

$$x \in \mathbb{R}^d$$

$$y \in \{0, 1\} \text{ "Label"}$$

$$\mathcal{D} = \{(x^i, y^i)_{i=1:n}\}$$

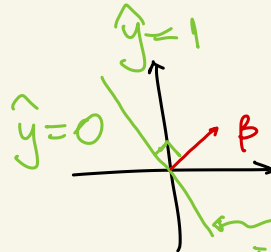
Model $f(x) = \ln \underbrace{\frac{P_{y|x}(y=1|x)}{P_{y|x}(y=0|x)}}_{\substack{\text{meaning} \\ \text{odds} \in (0, \infty) \\ \text{log odds} \in (-\infty, \infty)}} = \underline{\beta^T x} \quad \beta \in \mathbb{R}^d$

parametrization

Classification with $f = \beta^T x$ (β known)

$$\hat{y} = \frac{\text{sign } f(x) + 1}{2} \iff f(x) > 0 \iff \underbrace{P(y=1|x)}_{y|x} > \underbrace{P(y=0|x)}_{y|x}$$

Probabilistic classifier

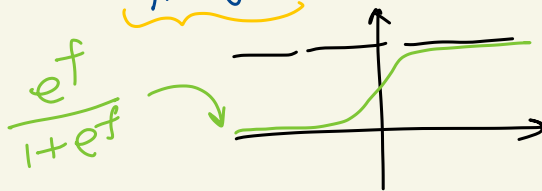


linear classifier $\iff$
decision boundary
linear

- $P_{y|x}[y=1|x] = p$

$$f = \ln \frac{p}{1-p} = \ln \frac{P_{y|x}(y=1|x)}{P_{y|x}(y=0|x)} \Rightarrow e^f = \frac{p}{1-p} \Rightarrow p = \frac{e^f}{1+e^f}$$

$f = \beta^T x$



$$p = \frac{e^f}{1+e^f} \quad e^{(0/1)}$$

$$1-p = \frac{1}{1+e^f}$$

$p + (1-p) = 1$

- Estimating β by Max likelihood

likelihood

$$L(\beta) = P[y^{1:n} | x^{1:n}, \beta]$$

$$= \prod_{i=1}^n P_{y|x}[y^i | x^i, \beta] = \prod_{i=1}^n \frac{e^{y^i \beta^T x^i}}{1 + e^{\beta^T x^i}}$$

$\downarrow$ $\downarrow$

P for $y^i=1$ $1-P$ for $y^i=0$

confidence
of
classification

$$P_{y|x}[y|x] = \frac{e^{y^T \beta}}{1 + e^{\beta^T x}}$$

log likelihood

$$l(\beta) = \sum_{i=1}^n [y^i \beta^T x^i - \ln(1 + e^{\beta^T x^i})]$$

maximize over β
by Gradient
Ascent

$$l(\beta) = \sum_{i=1}^n \left[y^i \overset{f}{\beta^T x^i} - \ln(1 + e^{\beta^T x^i}) \right]$$

Gradient $\frac{\partial l}{\partial \beta} \in \mathbb{R}^d$

$$\frac{\partial}{\partial \beta_j} = \frac{\partial}{\partial f} \cdot \frac{\partial f}{\partial \beta_j} = \sum_{i=1}^n \left[y^i - \frac{e^{f(x^i)}}{1 + e^{f(x^i)}} \right] x^i$$

$$\frac{\partial}{\partial \beta} \left[\frac{\partial l}{\partial \beta_j} \right]_{j=1:d} = \frac{\partial}{\partial f} \left[\frac{\partial f}{\partial \beta_j} \right]_{j=1:d}$$

$$\frac{\partial}{\partial \beta} = \frac{\partial}{\partial f} (\beta^T x) = x \leftarrow \text{Exercise } \Leftrightarrow \frac{\partial f}{\partial \beta_j} = x_j$$

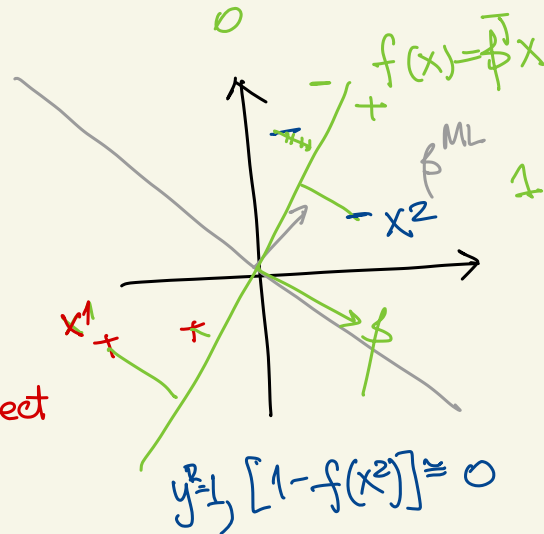
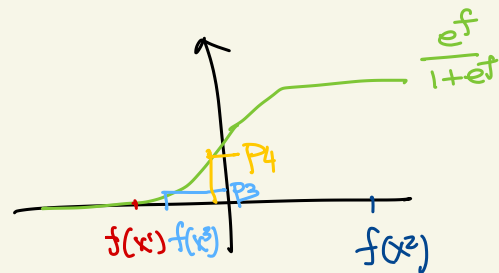
prove this

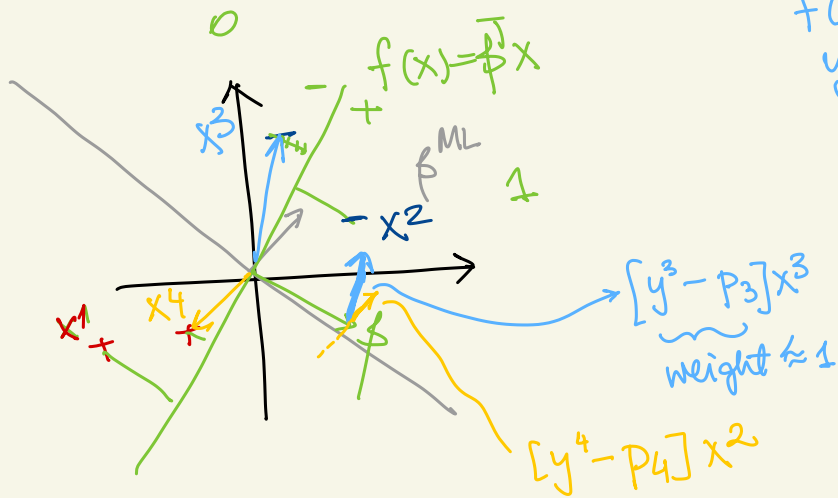
$$\frac{\partial l}{\partial f} = \sum_{i=1}^n \left[y^i - \underbrace{\frac{e^{f(x^i)}}{1 + e^{f(x^i)}}}_{p_j | x^i} \right] \in \mathbb{R}$$

by the model

$$0 \approx \frac{e^{f(x^i)}}{1 + e^{f(x^i)}} = p_1 \quad \begin{cases} y^i = 0 \\ |f(x^i)| \text{ large} \\ f(x^i) < 0 \text{ correct} \end{cases}$$

$$0 > [y^i - p_1] \approx 0$$



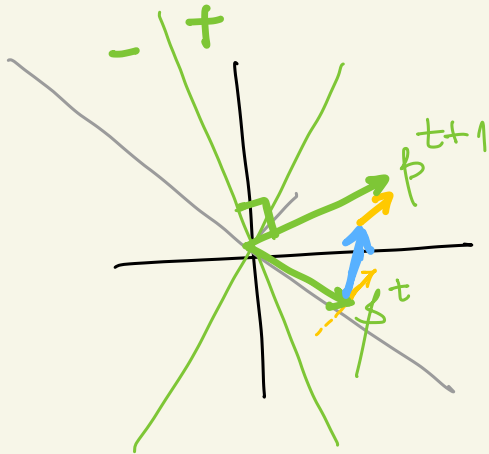


$$f(x^3) < 0 \Rightarrow y^3 - p_3 > 0 \approx 1$$

$$y^3 = 1$$

$$f(x^4) \leq 0 \Rightarrow y^4 - p_4 \approx -\frac{1}{2}$$

$$y^4 = 0$$



$$\beta^{t+1} = \beta^t + \eta \frac{\partial}{\partial \beta}$$

$$= \beta^t + \eta \sum_{i=1}^n [y^i - p^i] x^i$$

linear combination
of x^i

β_{yx} wanted

classification $y \in \{0,1\}$ or discrete

(logistic)

Clustering: wanted clusters = groups in data

Unsupervised

no output

Ex: dimension reduction PCA

$x = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$ $x_1 \xrightarrow{\text{causes}} x_2$

x_1 = treatment, x_2 = effect, health
smoking cancer

Lecture Notes IX – Clustering

Marina Meilă
`mmp@stat.washington.edu`

Department of Statistics
University of Washington

May, 2023

Paradigms for clustering ←

+ what is clustering

Parametric clustering algorithms (K given)

Cost based / hard clustering

K-means clustering and the quadratic distortion

Model based / soft clustering

Issues in parametric clustering

Selecting K

Reading: Ch. 18

What is clustering? Problem and Notation

- ▶ **Informal definition Clustering** = Finding groups in data
- ▶ **Notation**
 - $\mathcal{D}$ = $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ a **data set**
 - n = number of **data points**
 - K = number of **clusters** ($K \ll n$)
 - Δ = $\{C_1, C_2, \dots, C_K\}$ a partition of $\mathcal{D}$ into disjoint subsets
 - $k(i)$ = the **label** of point i
 - $\mathcal{L}(\Delta)$ = cost (loss) of Δ (to be minimized)
- ▶ **Second informal definition Clustering** = given n **data points**, separate them into K **clusters**
- ▶ Hard vs. soft clusterings
 - ▶ **Hard** clustering Δ : an item belongs to only 1 cluster
 - ▶ **Soft** clustering $\gamma = \{\gamma_{ki}\}_{k=1:K}^{i=1:n}$
 γ_{ki} = the **degree of membership** of point i to cluster k

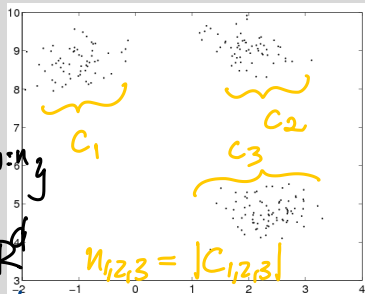
$$\sum_k \gamma_{ki} = 1 \quad \text{for all } i$$

(usually associated with a probabilistic model)

$$\mathcal{D} = \{x^i\}_{i=1}^n$$

$$x \in \mathbb{R}^d$$

$$n = |\mathcal{D}|$$



$$n_{1,2,3} = |C_{1,2,3}|$$

$$C_{1,2,3} = \text{clusters}$$

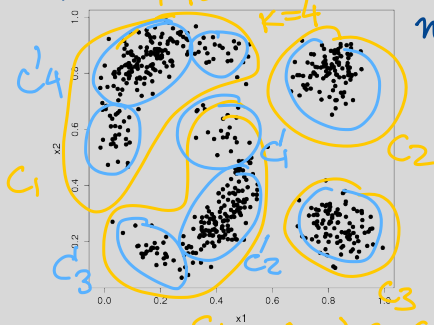
$$\mathcal{D} = C_1 \cup C_2 \cup C_3$$

$$n = n_1 + n_2 + n_3$$

clustering

$$\Delta = \{C_1, C_2, C_3\} \Rightarrow k=3$$

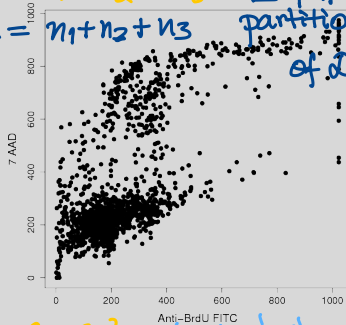
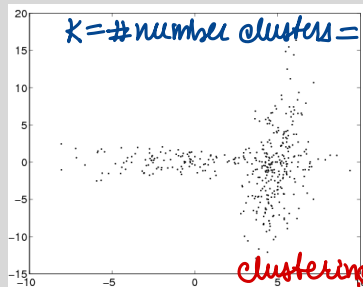
partition of $\mathcal{D}$

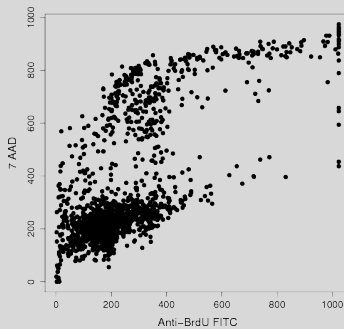
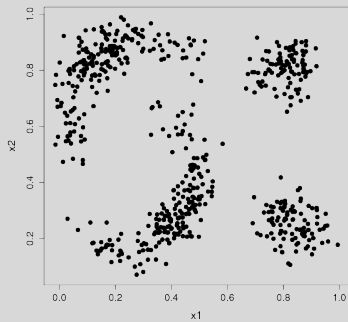
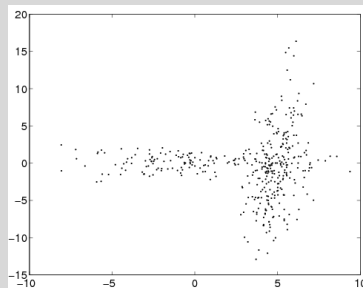
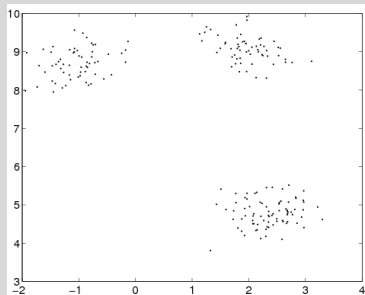


$$C_4$$

$$\Delta = \{C_1, C_2, C_3, C_4\}$$

$$\Delta' = \{C'_1, C'_2, \dots\}$$





(from [Nugent and Meila, 2010])

Paradigms

← each method has own definition of cluster

Depend on type of data, type of clustering, type of cost (probabilistic or not), and constraints (about K , shape of clusters)

► Data = vectors $\{x_i\}$ in $\mathbb{R}^d$

Parametric

(K known)

Cost based [hard]

Model based [soft]

Non-parametric

(K determined

by algorithm)

Dirichlet process mixtures [soft]

Information bottleneck [soft]

Modes of distribution [hard]

Gaussian blurring mean shift [Carreira-Perpinan, 2007] [hard]

► Data = similarities between pairs of points $[S_{ij}]_{i,j=1:n}$, $S_{ij} = S_{ji} \geq 0$ **Similarity based clustering**

Graph partitioning

spectral clustering [hard, K fixed, cost based]

typical cuts [hard non-parametric, cost based]

Affinity propagation

[hard/soft non-parametric]

Classification vs Clustering

	Classification	Clustering
Cost (or Loss) $\mathcal{L}$	Expected error	many! (probabilistic or not)
	Supervised	Unsupervised
Generalization	Performance on new data is what matters	Performance on current data is what matters
K	Known	Unknown
“Goal”	Prediction	Exploration <i>Lots of data to explore!</i>
Stage of field	Mature	Still young