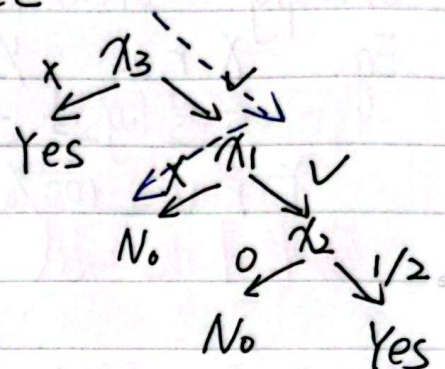


CS182: Introduction to Machine Learning

决策树 Decision Tree

x_1	x_2	x_3	y
✓	0	✓	No
X	1	✓	No
X	0	X	Yes
✓	1	✓	Yes
✓	2	X	Yes



想以 Feature x_1, x_2, x_3 判断 y , 可用决策树建立流程
上表格用来做 training set, 建右上图的树。则遇到:

$(x_1, x_2, x_3) = (X, 2, \checkmark)$ 样例, 决策流程如蓝笔所示
遇到 y 状态叶子结点, 流程停止

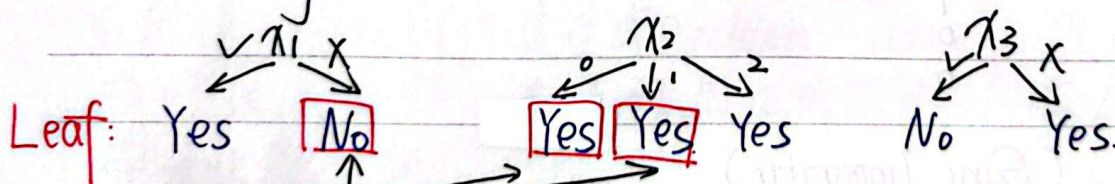
那么, 在树中, 依据为先 x_3 后 x_1 后 x_2 , 能换顺序么?



• Def: A splitting criterion is a function that measures how good or useful splitting on a particular feature is for a specified dataset.

Feature 顺序怎么选? 由 splitting criterion 量化
有以下常用标准:

① Training Error Rate:



Leaf:

判断 Leaf Node 为谁时: 填啥 error 最低. 就填啥; 一样则

Error Rate: $4/5$

$3/5$

$4/5$

则选 x_1/x_3 均可

均可



② Mutual Information Gain

Entropy: $H(X) = - \sum_{v \in V(X)} P(X=v) \log_2 [P(X=v)]$

Eg: X r.v.: $\frac{1}{2}$ 为 0, $\frac{1}{2}$ 为 1. 则称 X 的熵为:
 $-\frac{1}{2} \log_2 \frac{1}{2} - \frac{1}{2} \log_2 \frac{1}{2} = 1$

Y r.v.: 100% 为 0, $H(Y) = -1 \log_2 1 = 0$

感性来说: 越不确定, 则熵越大

而又 Def: $I(Y; X) = H(Y) - H(Y|X)$

$= H(Y) - \sum_{v \in V(X)} P(X=v) H(Y|X=v)$, i.e., $I(y; x_d)$

$= H(Y) - \sum_{v \in V(x_d)} \underbrace{f_v}_{\text{fraction of data points where } x_d=v} \underbrace{H(Y|x_d=v)}_{\text{set of all labels where } x_d=v}$

Eg: x_d y $I(y, x_d) = H(Y) - \sum_v H(Y|x_d=v) \cdot x \cdot f_v$
 $\begin{array}{cc} 1 & 1 \\ 0 & 0 \end{array}$
 $\begin{array}{cc} 0 & 0 \\ 0 & 0 \end{array}$
 x_d y
 $\begin{array}{cc} 1 & 1 \\ 0 & 1 \\ 1 & 0 \\ 0 & 0 \end{array}$
 $\begin{array}{cc} 0 & 0 \\ 0 & 0 \end{array}$

$= 1 - \frac{1}{2} H(Y|x_d=1) - \frac{1}{2} H(Y|x_d=0)$

$= 1$

$I(y, x_d) = H(Y) - \dots$

$= 1 - \frac{1}{2} H(Y|x_d=1) - \frac{1}{2} H(Y|x_d=0)$

$= 1 - \frac{1}{2} \cdot 1 - \frac{1}{2} \cdot 1 = 0$

则在: x_1 x_2 y 这样的情况: 选 x_1 作为第一个决策点

$\begin{array}{ccc} 1 & 1 & 1 \\ 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{array}$

③ 基尼指数 (Gini Impurity)

$Gini(D) = \sum_{k=1}^{|Y|} \sum_{k' \neq k} P_k P_{k'} = 1 - \sum_{k=1}^{|Y|} P_k^2$

$Gini_index = \sum_{v=1}^V \frac{|D^v|}{|D|} Gini(D^v) \Rightarrow a^* = \underset{a \in A}{\operatorname{argmin}} Gini_index(D, a)$



在选出第一个 feature 之后. 取 $X_i = D_i$ 的样本子集, 抛开 X_i 因素, 考虑选下一个 X_i

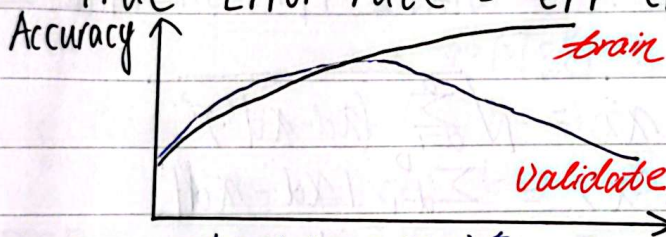
Decision Tree Pros: 可解释; 高效; 可用于分类/回归
但 Cons: splits consider immediate impact; **Overfit***

若 X_i 特征是 real-valued feature, 则可以考虑划分分类区间, 如: $(-\infty, a]$ - 类 $[a, b)$, $[b, +\infty)$... ($b > a$)

*: Training Error = $\text{err}(h, D_{\text{train}})$

Test Error = $\text{err}(h, D_{\text{test}})$

True Error rate = $\text{err}(h)$, h 为所有样品



Why overfitting? 有时划分过程重复, branch 过多

为了对付 Decision Tree 中的 overfitting, 可采用剪枝处理:

分为 prepruning & post-pruning

预剪枝: 对每个结点在划分前估计, 若当前结点的划分不能带来泛化性能提升, 则停止划分并当前结点标记为叶结点

后剪枝: 先从训练集生成 Decision Tree, 然后自底向上对非叶结点考察, 若该结点生成的子树换成叶子结点能带来性能升级, 则换。

在 CS182 PPT 中展示的仅为 post-pruning, 并且顺序是自上而下。以西瓜书为主。

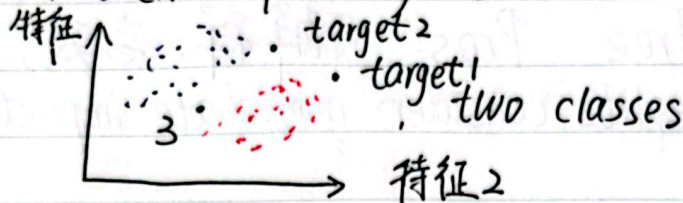
两策略 (pre & post) 均在讨论: 一个非叶子结点是否有必要划分出新特征; pre 是生成时便讨论 splitting feature (selected through criterion) 有必要么; post 是生成后自下而上讨论这一点。



KNN: K-Nearest Neighbors

— Duck test: If it looks like a duck, swims like a duck, and quacks like a duck, then it probably is a duck

考虑一个 Dataset:



那么新样本 target1 是 blue(B) or red(R)? 直观上是 R, 因为它距离红点“更近点”; target2 & 3 呢? 也许就要多看自己身边的样本点“们” label 是啥。

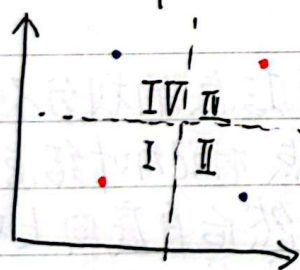
上二段有几个透露的关键:

① “近” $\Rightarrow$ Distance Metric! 如何衡量样本间在特征空间上的距离? Empirically:

$$\text{Euclidean: } d(x, x') = \sqrt{\sum_{d=1}^D (x_d - x'_d)^2}$$

$$\text{Manhattan: } d(x, x') = \sum_{d=1}^D |x_d - x'_d|$$

② “们” $\Rightarrow$ 看身边哪些点? Intuitively, k-nearest samples! KNN 算法 idea 便呼之欲出。



在左例中, 黑线划分 Feature Space 为 6 份, 构成 Voronoi Diagram。划分依据: e.g., 在 I 中的点, 身边最近的一个点是红点, 在 II 中, 是红点。 (KNN, $k=1$)

可视构成 Voronoi Diagram 的过程为 training, i.e., training 是无 error 的, 因为 $d(x, x) = 0$; 在 Test 时, 对于一个新 sample, 找 K-Nearest Points, 取 labels 然后 majority vote 决定它的 label。



既然有 majority vote, 在 Binary Classification 中, 我便希望 ties 不会发生, 因此可设为奇数。当然也可: ① 再看额外一个点

② 移去 k 个中最远的 ③ 距离 $\Rightarrow$ 权重 ④ another distance metric
在 K -NN 中, Distance 为欧式距离 $d: \mathbb{R}^M \times \mathbb{R}^M \rightarrow \mathbb{R}$

$$d(u, v) = \sqrt{\sum_{m=1}^M (u_m - v_m)^2}$$

当然其他 Distance 也可以。那么这个 algorithm complexity 为?
设 $x \in \mathbb{R}^M$, N samples for training, 则:

train time: $O(N)$

on average

Predict time (one sample): $O(MN)$

可用 k -d tree 优化, 则上述两者变为: $O(MN \log N)$; $O(2^M \log N)$

在 M 大时, GG! 实践中常用 stochastic approximation

那么 KNN Performance 如何呢? 有 Theoretical Guarantee:
(by Cover & Hart, 1967): $h(x)$ 为 KNN ($k=1$) 的二分类器:

$$\text{error true}(h) < 2 \times \text{Bayes Error Rate.}$$

BER 可近似理解为: the best you could probably do

★ Inductive Bias: (归纳偏差) 是算法在进行归纳学习时所持有的先验假设。如决策树的 Inductive Bias 就是尝试找到 (最小的) 决策树 s.t. 训练误差 $\downarrow$ 且 交互信息 $\uparrow$

Occam's Razor: try to find the 'simplest' classifier that explains the training dataset

那么 KNN 的 Inductive Bias 便有: $\left\{ \begin{array}{l} 1. \text{ similar point} \rightarrow \text{similar label} \\ 2. \text{ 所有维度 are created equally} \end{array} \right.$
第2点 $\Rightarrow$ 有一个大问题: 特征不同的 scale 对 Distance 有影响!
因此假设应含第2点; 当然也可人为定义权重以纠正。



(Voronoi Diagram中)

之前提到 k 的影响：那么一个重要现象是： $k \uparrow$, boundary 越平滑；ppt 中的例子十分直观！

Extension: Model Selection & Experiment Design

一个 model 可以设很多不同 setting：如 KNN, $k = ?$; Decision Tree 中, criterion? 因此与其说是一个 "model", 不如是 "一类" model.

setting 有参数, 也有超参。model training 的是为了什么?

是当前 model structure、hyperparameter 下的最优 parameters.

因此 model selection 就是在众多 model 中挑最好, 而众多 model 不同的是超参! 如何挑超参呢? 需要实验!

Training exp: input: 超参 & T Data; output: Best parameters

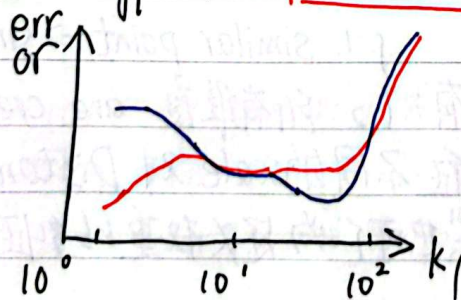
Hyper exp: input: T & V Data; Output: Best hyperparameters

关于 Training Exp, 有 Cross-Validate 特殊技巧: E.g., Data Fold $\{1, 2, 3\}$. 则可 $\{1, 2\}$ 训 $\{3\}$ 验, $\{2, 3\}$ 训 $\{1\}$ 验 ... 它们 Loss 的和可作为为实验设计训练 Loss 的重要参考!

挑超参作实验 规范流程:

- ① Data $\Rightarrow$ {train} | {val} | {test} 挑出来
- ② {train} 上训, 更换 hyper; {val} 上测的最好的 hyper
- ③ 以这个 hyper, 在 {train} + {val} 上训, 用 {test} 看表现

Eg: K-NN:



超参选择实验: Grid!

$\alpha, \beta, \dots$ 所有可能取值作组合

——实验。也可 Random 选 α, β , 重复多次。

Δ : 有限时间内, Random 比 Grid

更有可能找到好的 hyper, 尤其高维空间中



那么:如果 data 并不 linearly separable 呢? 那么 perceptron 不会收敛。但有定理发现: 在 examples 上训一遍, 依然能得出近似边界

Theorem: 设 $\langle (x_1, y_1), \dots, (x_m, y_m) \rangle$, $\|x_i\| \leq R$, 令 u 为任一单位向量, $\gamma > 0$. 定义每一个 example deviation 为:

$$d_i = \max \{0, \gamma - y_i (u \cdot x_i)\}$$

再定义: $D = \sqrt{\sum_{i=1}^m d_i^2}$, 则 Perceptron 在这个集上 mistakes 数量:

$$\leq \left(\frac{R+D}{\gamma}\right)^2$$

* Proof: Claim 1: $W_{t+1} \cdot W^* \geq W_t \cdot W^* + \gamma$ 在 W^* 上投影 点到 W^* 平面最短距离
 W^* 为 max-margin 权重向量, 且 $\|W^*\| = 1$

Claim 2: $\|W_{t+1}\|^2 \leq \|W_t\|^2 + R^2$, 因为 $W_{t+1} = W_t \pm x$, $\|x\| \leq R$

则 $W_{M+1} \cdot W^* \geq \gamma M$, $\|W_{M+1}\| \leq R\sqrt{M}$ 且 $W_{M+1} \cdot W^* \leq \|W_{M+1}\|$
 $\therefore \gamma M \leq R\sqrt{M}$, $M \leq \left(\frac{R}{\gamma}\right)^2$



Perceptron

Key idea: Try to learn a hyperplane

设 $\vec{a}$ 为列向量 (by default), $\vec{a}^T \vec{b} = \sum_{d=1}^D a_d b_d$

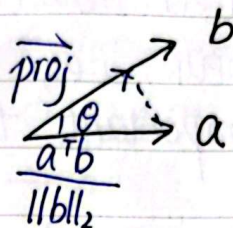
L_2 -norm: $\|\vec{a}\|_2 = \sqrt{\vec{a}^T \vec{a}}$

则 $\vec{a}$ 在 $\vec{b}$ 上的投影为?

$$\Delta: \vec{a} \cdot \vec{b} = \|\vec{b}\| \|\vec{a}\| \cos \theta$$

$$\text{而 } \vec{a}^T \vec{b} = \|\vec{b}\| \cdot \|\text{proj}\|, \quad \|\text{proj}\| = \frac{\vec{a}^T \vec{b}}{\|\vec{b}\|_2}$$

故: $\text{proj} = \frac{(\vec{a}^T \vec{b})}{\|\vec{b}\|_2^2} \vec{b}$, 因为方向为 $\vec{b}$ 方向



在 2D 中, $w_1 x_1 + w_2 x_2 + b = 0$ 为 line; 3D 中为 plane

4+D 中表示为 hyperplane, $w^T x + b = 0$. 值得注意的是:

w 向量指向总是 $w^T x + b > 0$ 的部分!

因此, hyperplane creates 2 halfspaces.

Perceptron 算法: 如何更新 w 与 b ?

初始化: $w = \vec{0}$, $b = 0$

$$\text{input: } x^{(t)} y^{(t)} = \{\pm 1\}, \quad \hat{y} = \text{sign}(w^T x + b) = \begin{cases} +1 & \text{if } w^T x + b \geq 0 \\ -1 & \text{otherwise} \end{cases}$$

若 misclassify - 个 +1 例子 ($y^{(t)} = +1, \hat{y} = -1$):

$$w \leftarrow w + x^{(t)} \quad b \leftarrow b + 1$$

若 misclassify - 个 -1 例子 ($y^{(t)} = -1, \hat{y} = +1$):

$$w \leftarrow w - x^{(t)} \quad b \leftarrow b - 1$$

为何这样更新? 若: $b + w^T x < 0$, $y^{(t)} = +1$, 欲 $w^T x \uparrow$.

$$w \leftarrow w + x^{(t)}, \quad w^T x = (w + x^{(t)})^T x = w^T x + \underbrace{x^{(t)T} x}_{>0} > 0, \uparrow$$

[反之亦然] 同时 b 也 $\uparrow$, $b + w^T x \uparrow$

而更新逻辑可简化为:

$$\begin{cases} w \leftarrow w + y^{(t)} x^{(t)} \\ b \leftarrow b + y^{(t)} \end{cases}$$

可见, w 可视为 $x^{(1)}, \dots, x^{(n)}$ 的线性组合



同时还有一个 trick: $\vec{x} = \begin{bmatrix} 1 \\ \vec{x} \end{bmatrix}$ $\vec{\theta} = \begin{bmatrix} b \\ \vec{w} \end{bmatrix}$. 则:

$$\hat{y} = \text{sign}(\vec{w}^T \vec{x} + b) = \text{sign}(\vec{\theta}^T \vec{x})$$

$$\theta \leftarrow \theta + y^{(t)} \vec{x}^{(t)}, \text{ where } \vec{x}^{(t)} = \begin{bmatrix} 1 \\ \vec{x}^{(t)} \end{bmatrix}$$

那么 Perceptron 的 Inductive Bias 是什么? Decision Boundary should be linear, i.e., hyperplane, 且 Recent mistakes are more important than older ones.

而当 θ, b 训练到不断产生正确预测时, 我们就认为感知机收敛 (converge) 了。同时, Perceptron 也可能受 overfitting 影响。之前介绍到, 学的其实是一个超平面; 在 binary 分类中, 如果两类点之间的分隔不是“面”而是“space”呢?



Def: 对二分类问题, 样本集 S is linearly separable if there exists a linear boundary that separate the points

“space”大小又引出下面定义: Def: The margin γ for dataset D is the greatest possible distance between a linear separator and the closest point in D to the linear separator

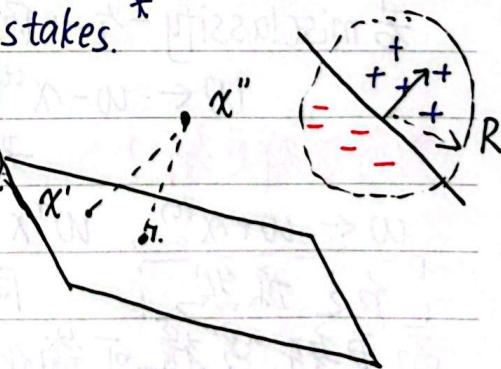
Theorem: 若 Data margin 为 γ . 所有点都在半径为 R 的球内, 则 online perceptron 算法 makes $\leq (R/\gamma)^2$ mistakes.*

那么 margin 如何计算? 相当于是 $(\vec{x}'' - \vec{x}')$

向量在 $\vec{w}/\|\vec{w}\|_2$ 上的投影长度 (i.e., 内积)

$$\text{则 } \left| \frac{\vec{w}^T (\vec{x}'' - \vec{x}')}{\|\vec{w}\|_2} \right| = \frac{|\vec{w}^T \vec{x}'' - \vec{w}^T \vec{x}'|}{\|\vec{w}\|_2}$$

$$= \frac{|\vec{w}^T \vec{x}'' + b|}{\|\vec{w}\|_2}$$



SVM

感知机中, inductive bias 中认为二分任务中的边界是线性的

但若 if not linearly separable 呢?

这里介绍 kernel method (当然方法不限于此)

Def: A kernel K is a legal def of dot-product, i.e., 有一个隐式映射: Φ , s.t., $K(x, y) = \Phi(x) \cdot \Phi(y)$.

Eg. $K(x, y) = (x \cdot y + 1)^d$, $\Phi: n\text{维} \rightarrow n^d\text{维}$

这个定义看起来没啥用; 但 perceptron 中因 $x \cdot z$ 是跳不出 linear 的, 而 $x \cdot z$ 换为 $K(x, z)$ 就能拥抱高维

Formal Definition: $K(\cdot, \cdot)$ is a kernel if it can be viewed as a legal definition of inner product

• $\exists \Phi: X \rightarrow \mathbb{R}^N$ s.t. $K(x, z) = \Phi(x) \cdot \Phi(z)$, range of Φ : Φ -space
但 Φ 是隐式的! 这仅是为了方便 view kernel as inner product.

例: $K(x, z) = (x \cdot z)^d$ 对应: $\Phi(x) = (x_1^2, x_2^2, \sqrt{2}x_1x_2)$

可见 $\Phi: \mathbb{R}^2 \rightarrow \mathbb{R}^3$, $\Phi(x) \cdot \Phi(z) = (x_1^2, x_2^2, \sqrt{2}x_1x_2) \cdot (z_1^2, z_2^2, \sqrt{2}z_1z_2)$
 $= (x_1z_1 + x_2z_2)^2 = (x \cdot z)^2 = K(x, z)$

✧: 同一个 kernel, Φ 也许并不唯一! Eg. $(x, x) \rightarrow \Phi(x) = (x_1^2, x_2^2, x_1x_2, x_2x_1)$
应避免显式地扩到高维, i.e., Φ ! 因为 feature space 会长的飞快!!

因此可以尝试用 $K(x, z)$ 代替 $x \cdot z$ 。例如 Perceptron 中:

$w_t = \alpha_{i1} x_{i1} + \dots + \alpha_{ik} x_{ik}$ (可视为未正确分类的 sample 的线性组合).

对于第 t 的错误: mistake on positive: $\alpha_{it} \leftarrow 1$, store x_{it}

mistake on negative: $\alpha_{it} \leftarrow -1$, store x_{it}

$w_t \cdot x = \alpha_{i1} K(x_{i1}, x) + \dots + \alpha_{ik} K(x_{ik}, x)$



这个方法在 margin 恰是 Φ 维空间中 linear boundary 时表现很好!

可知: 若 margin 在 Φ -space 中, 则 Perceptron makes $(\frac{R}{\gamma})^2$ mistakes

Kernel Example:

Linear: $K(x, z) = x \cdot z$

Polynomial: $K(x, z) = (x \cdot z)^d$ or $K(x, z) = (1 + x \cdot z)^d$

Gaussian: $K(x, z) = \exp \left[-\frac{\|x - z\|^2}{2\sigma} \right]$

Laplace: $K(x, z) = \exp \left[-\frac{\|x - z\|}{2\sigma} \right]$

Properties of Kernel: K is kernel iff: ① K is symmetric

② For any training points $x_1, \dots, x_m$, and $\forall a_1, \dots, a_m \in \mathbb{R}$:

$$\sum_{i,j} a_i a_j K(x_i, x_j) \geq 0$$

i.e., $K = (K(x_i, x_j))_{i,j=1,\dots,n}$ 是半正定的: $a^T K a \geq 0$

若 $K_1(\cdot, \cdot)$, $K_2(\cdot, \cdot)$ 均为 kernels, 则 $\forall c_1, c_2 \geq 0$, $K(x, z) = c_1 K_1(x, z) + c_2 K_2(x, z)$ 也是 kernel! 且 $K_1(\cdot, \cdot)$, $K_2(\cdot, \cdot)$ 也是!

Proof: $\phi(x) = (\sqrt{c_1} \phi_1(x), \sqrt{c_2} \phi_2(x))$, $\phi(x) \phi(z) = c_1 \phi_1(x) \phi_1(z) + c_2 \phi_2(x) \phi_2(z)$

$\phi(x) = (\phi_{1,i}(x), \phi_{2,j}(x))_{i \in \{1, \dots, n\}, j \in \{1, \dots, m\}}$

$$\phi(x) \phi(z) = \sum_{i,j} \phi_{1,i}(x) \phi_{2,j}(x) \phi_{1,i}(z) \phi_{2,j}(z)$$

$$= \sum_i \phi_{1,i}(x) \phi_{1,i}(z) \left(\sum_j \phi_{2,j}(x) \phi_{2,j}(z) \right)$$

综合上述所有铺垫, SVM, Support Vector Machine 问世!

Def: Margin of example x w.r.t. linear separation w is the distance from x to plane $w \cdot x = 0$

Def: Margin γ_w of a set S is min margin over $x \in S$, w.r.t. a linear separator w .



Def: Margin γ of set S is the maximum γw over all linear sep. w

因此 SVM 目标为: Input: $S = \{(x_1, y_1) \dots (x_m, y_m)\}$

Find: some w and largest maximum γ where:

① $\|w\|^2 = 1$ ② For all i , $y_i w \cdot x_i \geq \gamma$

Output: Maximum marginal separator

等效为: $y_i \cdot \frac{w}{\gamma} \cdot x_i \geq 1$, 令 $w' = w/\gamma$. 可见是 $y_i w' \cdot x_i \geq 1$
满足前提下: minimize $\|w'\|^2$, i.e.,

$$\operatorname{argmin}_w \|w\|^2, \text{ s.t., } \forall i, y_i w \cdot x_i \geq 1$$

Dual Problem* (Equivalent to:)

$$\begin{aligned} \max_{\alpha} \sum \alpha_i - \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m \alpha_i \alpha_j y_i y_j \phi(x_i)^T \phi(x_j) \\ = \max_{\alpha} \sum \alpha_i - \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m \alpha_i \alpha_j y_i y_j K(x_i, x_j) \end{aligned}$$

$$\text{s.t. } \sum_{i=1}^m \alpha_i y_i = 0, \alpha_i \geq 0$$

最终: classifier 为: $w = \sum_i \alpha_i y_i x_i$

* 这一坨式子咋得到的?

在最大化间隔同时, 允许部分样本不满足约束 (尽可能少)

$$\min_{w, b} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \ell_{\alpha_1}(y_i (w^T x_i + b) - 1), \quad C: \text{constant}$$

$$\ell_{\alpha_1}(z) = \begin{cases} 1 & \text{if } z < 0 \\ 0 & \text{otherwise} \end{cases}, \text{ 取 hinge loss: } \ell_{\text{hinge}}(z) = \max(0, 1 - z)$$

$$\min_{w, b} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \max(0, 1 - y_i (w^T x_i + b)), \text{ 引入 slack 变量 } \xi_i \geq 0:$$

$$= \min_{w, b, \xi_i} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \xi_i \quad ①$$

三个约束条件

$$\text{s.t. } y_i (w^T x_i + b) \geq 1 - \xi_i, \quad \xi_i \geq 0$$

②

③

KOKUYO



则构造 Lagrange: $L(w, b, \alpha, \xi, \mu) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \xi_i$
 $+ \sum_{i=1}^m \alpha_i (1 - \xi_i - y_i (w^T x_i + b)) - \sum_{i=1}^m \mu_i \xi_i$

$$\frac{\partial L}{\partial w} = 0 \Rightarrow w = \sum_{i=1}^m \alpha_i y_i x_i \quad (1)$$

$$0 = \sum_{i=1}^m \alpha_i y_i \quad (2) \quad C = \alpha_i + \mu_i \quad (3) \quad \text{代回去:}^*$$

Dual Problem: $\max_{\alpha} \sum_{i=1}^m \alpha_i - \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m \alpha_i \alpha_j y_i y_j x_i^T x_j$

s.t. $\sum_{i=1}^m \alpha_i y_i = 0$, $0 \leq \alpha_i \leq C$, $i=1, 2, \dots, m$ $= K(x_i, x_j)$

这是 $K(x, z) = x \cdot z$ 时; 若为 $\Phi(x)$, 则 $\Phi(x_i)^T \Phi(x_j) = \Phi(x_i) \cdot \Phi(x_j)$

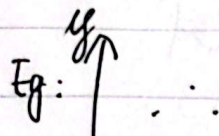
Δ : 附: 使用拉格朗日乘子法得到的便是“对偶问题”

$$\begin{aligned} * : \inf_{w, b} L(w, b, \alpha) &= \frac{1}{2} w^T w + \sum_{i=1}^m \alpha_i - \sum_{i=1}^m \alpha_i y_i w^T x_i - \sum_{i=1}^m \alpha_i y_i b \\ &= \frac{1}{2} w^T \sum_{i=1}^m \alpha_i y_i x_i - w^T \sum_{i=1}^m \alpha_i y_i x_i + \sum_{i=1}^m \alpha_i - b \sum_{i=1}^m \alpha_i y_i \\ &= -\frac{1}{2} w^T \sum_{i=1}^m \alpha_i y_i x_i + \sum_{i=1}^m \alpha_i \\ &= -\frac{1}{2} \sum_{i=1}^m \alpha_i y_i x_i^T \sum_{j=1}^m \alpha_j y_j x_j + \sum_{i=1}^m \alpha_i \\ &= \sum_{i=1}^m \alpha_i - \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^m \alpha_i \alpha_j y_i y_j x_i^T x_j \end{aligned}$$



Regression 回归

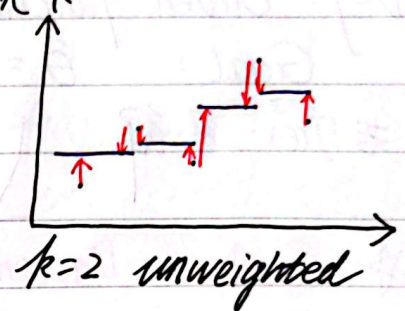
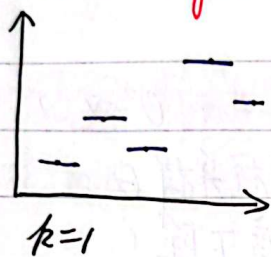
Goal: Given $\{(\vec{x}, y)\}$ dataset, learn a function:
 $y' = h(\vec{x})$, 应与 y 接近

Eg:  找函数拟合各点曲线

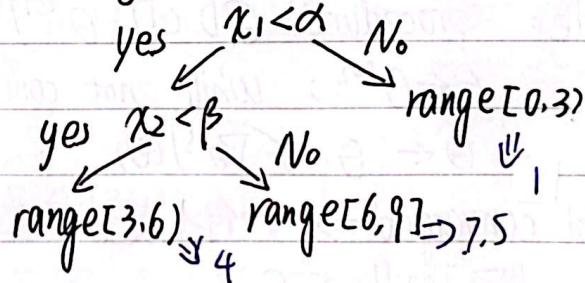
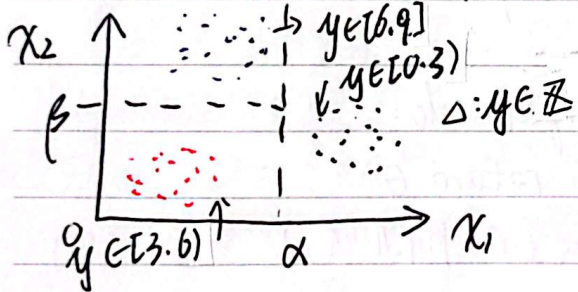
Method ①: k -NN 回归

$k=1$: pick the nearest x , return its y

or: $k=2$: weighted / unweighted 效果见下



Method ②: Decision Tree Regression



划分 x_i ; 然后 $\{y_i\}$ 聚为 $\sim$ 类, 看 mean (Decision Tree 分支到最后时)

Method ③: Linear Functions, Residuals & MSE

Δ Linear functions $\neq$ Linear decision Boundaries

Def: Regression is predicting real-valued outputs

$$D = \{(\vec{x}^{(i)}, y^{(i)})\}_{i=1}^n, \quad \vec{x}^{(i)} \in \mathbb{R}^M, y^{(i)} \in \mathbb{R}$$

Eg: $y = w^T x + b$

Eg: $y = \text{sign}(w^T x + b)$



目标

Key idea: Find the linear function h (with w, b) that minimizes the squares of residual for training set.

Residual: $e_i = |y^{(i)} - h(w^T x^{(i)} + b)|$ 不是点向直线
作垂线

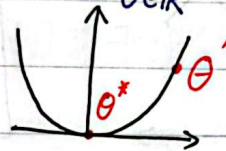
$$\text{MSE: Loss} = \frac{1}{n} \sum_{i=1}^n e_i^2$$

有MSE有什么用? The big picture: Optimization for ML

Def: Given function: $J(\theta)$, $J: \mathbb{R}^m \rightarrow \mathbb{R}$

Goal: $\hat{\theta} = \underset{\theta \in \mathbb{R}^m}{\operatorname{argmin}} J(\theta)$

形象理解:



在 θ' 如何一步步到 θ^* 呢?

Naive one: 随机采样 θ

★ Gradient Descent:

$$\nabla J(\vec{\theta}) = \begin{bmatrix} \frac{\partial J(\theta)}{\partial \theta_1} \\ \vdots \\ \frac{\partial J(\theta)}{\partial \theta_m} \end{bmatrix}$$

梯度下降!

$$\vec{\theta} \leftarrow \vec{\theta} - \delta t \nabla J(\vec{\theta}), \text{ where } \delta t \text{ is step size, } \in \mathbb{R}$$

Pseudocode: procedure GD($J, \theta^{(0)}$):

$\theta \leftarrow \theta^{(0)}$; while not converged do:

$\theta \leftarrow \theta - \gamma \nabla_{\theta} J(\theta)$; return θ

那么 not converged $\Rightarrow$ 如何判定 θ 已收敛至 θ^* 附近呢?

$$\Delta: \|\nabla_{\theta} J(\theta)\|_2 \leq \epsilon$$

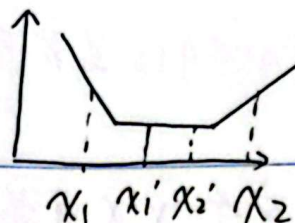
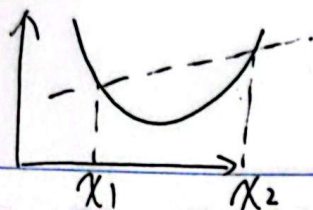
而 MSE Loss: $J^{(i)}(\theta) = (y^{(i)} - \theta^T x^{(i)})^2$, $J(\theta) = \frac{1}{n} \sum_{i=1}^N J^{(i)}(\theta)$

$$\text{则 } \frac{d}{d\theta_k} J^{(i)}(\theta) = 2(\theta^T x^{(i)} - y^{(i)}) \frac{d}{d\theta_k} (\theta^T x^{(i)} - y^{(i)})$$

$$= 2(\theta^T x^{(i)} - y^{(i)}) \frac{d}{d\theta_k} \left(\sum_{j=1}^K \theta_j x_j^{(i)} - y^{(i)} \right) = 2(\theta^T x^{(i)} - y^{(i)}) x_k^{(i)}$$

$$\therefore \frac{d}{d\theta_k} J(\theta) = \sum_{i=1}^N (\theta^T x^{(i)} - y^{(i)}) x_k^{(i)}, \quad \nabla_{\theta} J(\theta) = \begin{bmatrix} \frac{d}{d\theta_1} J(\theta) \\ \vdots \\ \frac{d}{d\theta_m} J(\theta) \end{bmatrix} = \sum_{i=1}^N (\theta^T x^{(i)} - y^{(i)}) x^{(i)}$$





No.

Date

Optimization for ML:

Convexity: A function $f: \mathbb{R}^D \rightarrow \mathbb{R}$, convex if:
 $\forall x^{(1)} \in \mathbb{R}^D, x^{(2)} \in \mathbb{R}^D, c \in [0, 1]$, 有:

$$f(cx^{(1)} + (1-c)x^{(2)}) \leq cf(x^{(1)}) + (1-c)f(x^{(2)})$$

上述为 not strictly convex^{*}; 下述为 strictly^{*}: 如 fig 1

$c \in (0, 1)$, 有: $f(cx^{(1)} + (1-c)x^{(2)}) < cf(x^{(1)}) + (1-c)f(x^{(2)})$

^{*}: 如 fig 2, 取 x_1', x_2' , 上面这行实现不了!

对于凸函数来说: global min: $f(x^*) \leq f(x), \forall x \in \mathbb{R}^D$

local min: $\exists \varepsilon$ s.t. $f(x^*) \leq f(x)$, 且 s.t. $\forall x, \|x - x^*\|_2 < \varepsilon$

Closed form Optimization:

key: 找到 θ^* , s.t. $\nabla J(\theta^*) = 0$, 其中 J 凸

若无法确保 J 凸, 则找 $\{\theta_i\}$, s.t. $\nabla J(\theta_i) = 0$ for $i \in \text{range}$

并确认它们是极大 or 小值

Given $D = \{(x^{(n)}, y^{(n)})\}_{n=1}^N$, 令: $X = \begin{bmatrix} 1 & x_1^{(1)} & \dots & x_D^{(1)} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_1^{(N)} & \dots & x_D^{(N)} \end{bmatrix} \in \mathbb{R}^{N \times D+1}$

而 $y = [y^{(1)}, \dots, y^{(N)}] \in \mathbb{R}^N$ 是目标:

$$J(\theta) = \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{2} (y^{(i)} - \theta^T x^{(i)})^2 \right) = \frac{1}{2N} \sum_{i=1}^N (x^{(i)T} \theta - y^{(i)})^2$$

$$= \frac{1}{2N} (X\theta - y)^T (X\theta - y) = \frac{1}{2N} (\theta^T X^T X \theta - 2\theta^T X^T y + y^T y)$$

$$\nabla_{\theta} J(\theta)^* = \frac{1}{2N} (2X^T X \theta - X^T y) = 0 \quad \therefore X^T X \hat{\theta} = X^T y$$

$$\hat{\theta} = (X^T X)^{-1} X^T y$$

^{*}: 矩阵求导!!

^{*}: $\frac{1}{2}$ 无伤大雅; $J(\theta)$ 含义为: 最小 Mean Square Error!!

Δ : 原: $w^T x + b$, 令 $x = \begin{bmatrix} 1 \\ x \end{bmatrix}$, $\theta = \begin{bmatrix} b \\ w \end{bmatrix}$, 则 $\Rightarrow \theta^T x$

KOKUYO



扫描全能王 创建

No. *: 直觉上: $N=2$, 那么超平面显然有很多种划分

Date

$\hat{\theta} = (X^T X)^{-1} X^T y$, 这个 $X^T X$ 可逆么? 一般认为 $N \gg D+1$ 时, 几乎总是 full-rank 的*. 但是关键在于计算 $X^T X$ 逆时间复杂度多少? 算 $X^T X: O(ND^2)$; 算它的逆: $O(D^3)$.

那么, 有没有不那么精确但复杂度能够接受的近似求法呢?

GD & SGD:

Gradient Descent:

$\theta \leftarrow \theta^{(0)}$
while not converged do
 $\theta \leftarrow \theta - \gamma \nabla J(\theta)$

N 项, $\theta \in \mathbb{R}^M$

return θ

而 Stochastic Gradient Descent:

$\theta \leftarrow \theta^{(0)}$
while not converged do
 $i \sim \text{uniform}(\{1, 2, \dots, N\})$
 $\theta \leftarrow \theta - \gamma \nabla_{\theta} J^{(i)}(\theta)$
return θ

$\Delta: J(\theta) = \frac{1}{N} \sum_{i=1}^N J^{(i)}(\theta)$

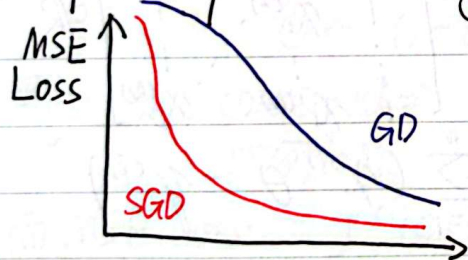
实质: 原来用全部样本, 现在一步用一个

why SGD 能 work? 因为 $E[\nabla_{\theta} J^{(i)}(\theta)] = \sum_{i=1}^N p(x^{(i)}, y^{(i)}) \nabla_{\theta} J^{(i)}(\theta) = \frac{1}{N} \sum_{i=1}^N \nabla_{\theta} J^{(i)}(\theta) = \nabla_{\theta} J(\theta)$



黑: GD
蓝: SGD

Def: An epoch is a single passing through training data



GD: One update per epoch

SGD: N update per epoch

step to convergence

Computation per step

Theoretical Comparison:

GD

$O(\log 1/\epsilon)$

$O(NM)$

SGD

$O(1/\epsilon)$

$O(M)$

理论上 SGD 收敛更慢, 但实践中更快

那么在上-面的 Linear Regression 情景中, $J(\theta) = \frac{1}{N} \sum_{i=1}^N (y^{(i)} - \theta^T x^{(i)})^2$

则 SGD:

$\theta \leftarrow \theta^{(0)}$
while not converged do:

为了保证 N 次可过一个 epoch

故是 range $(1, N)$ shuffle

以模拟 Uniform.

GD: $\theta \leftarrow \theta - \gamma \frac{1}{N} \sum_{i=1}^N (\theta^T x^{(i)} - y^{(i)}) x^{(i)}$
Campus $x^{(i)}$
return θ



扫描全能王 创建

MLE MAP

M 回顾 Likelihood: 有 N 个 i.i.d 样本: $D = \{x^{(1)}, \dots, x^{(N)}\}$ of r.v. X

L 若 X 离散, PMF 为 $p(x|\theta)$, 则 Likelihood of D is:

$$E \quad L(\theta) = \prod_{n=1}^N p(x^{(n)}|\theta)$$

若 X 连续, PDF 为 $f(x|\theta)$, 则:

$$L(\theta) = \prod_{n=1}^N f(x^{(n)}|\theta)$$

考虑 toss coin: $P(\text{Head}) = \theta$ $P(\text{Tail}) = 1 - \theta$

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} P(D|\theta) = \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^N P(x_i|\theta)$$

$$= \underset{\theta}{\operatorname{argmax}} \theta^{d_H} (1-\theta)^{d_T} \prod_{i=1}^N f(x_i|\theta)$$

$$\partial \prod_{i=1}^N f(x_i|\theta) / \partial \theta = 0 \Rightarrow \hat{\theta}_{MLE} = \frac{d_H}{d_H + d_T}$$

考虑 $X \sim \text{Gaussian}(\mu, \sigma)$, 欲 $\underset{\mu, \sigma}{\operatorname{argmax}}: \left(\frac{1}{\sqrt{2\pi}\sigma}\right)^N e^{-\sum_{i=1}^N (x_i - \mu)^2 / 2\sigma^2}$

$$\frac{\partial J}{\partial \mu} = \sum_{i=1}^N (x_i - \mu) / \sigma^2 = 0; \quad \frac{\partial J}{\partial \sigma^2} = -n + \frac{1}{2\sigma^2} = 0$$

$$\therefore \hat{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^N x_i \quad \hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^N (x_i - \hat{\mu})^2$$

△ 附: MLE 推出的 σ^2 有 bias! $\hat{\sigma}_{unbiased}^2 = \frac{1}{n-1} \sum_{i=1}^N (x_i - \hat{\mu})^2$

why? 因为 $\hat{\sigma}_{MLE}^2$ 中用的是 sample mean 而非 true mean

M : The Bayesian Way :

$$A \quad P(\theta|D) = \frac{P(D|\theta) P(\theta)}{P(D)} \propto \underbrace{P(D|\theta)}_{\text{Likelihood}} \underbrace{P(\theta)}_{\text{prior}}$$

P

$$\text{欲 } \underset{\theta}{\operatorname{argmax}} P(D|\theta) P(\theta) \Leftrightarrow \underset{\theta}{\operatorname{argmax}} [\log P(D|\theta) + \log P(\theta)]$$

考虑 flip coin: $P(D|\theta) = \binom{n}{d_H} \theta^{d_H} (1-\theta)^{d_T}$, if $P(\theta) \sim \text{Beta}(\beta_H, \beta_T)$

$$\text{则 } P(\theta|D) \sim \text{Beta}(\beta_H + d_H, \beta_T + d_T), \quad \hat{\theta}_{MAP}^* = \frac{d_H + \beta_H - 1}{d_H + \beta_H + d_T + \beta_T - 2}$$

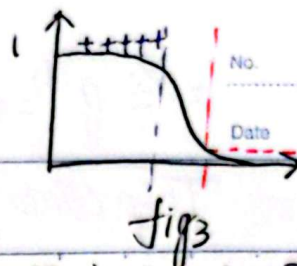
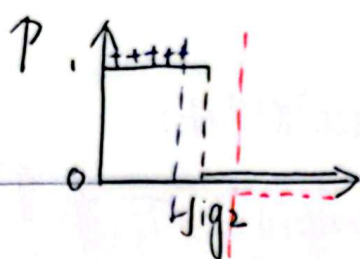
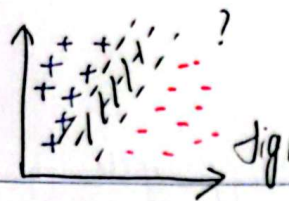
$$(\hat{\theta}_{MAP} = \underset{\theta}{\operatorname{argmax}} P(\theta|D))$$

总结: Frequentist: Sample 少时表现不好; Bayesians: 不同 prior 不同 answer

KOKUYO



扫描全能王 创建



Logistic Regression

Recall: Perceptron 是简单二分模型, 但它无法处理不可分的数据; 同时, Perceptron 无法建模 uncertainty: 如 fig1, 超平面划分有很多种, 但若有点在“ $\cdot$ ”区域呢? **Perceptron 被迫给出 ± 1 , 但直观上, 给出概率可能更合理!**

如 fig2 & fig3: 同样如 fig 中的 sample point, 在红蓝中间处, perceptron 只能硬性找一个地方然后作分界; 但感性上, 希望如 fig3 这样建模

用怎样的函数能很好建模 fig3 曲线? **Sigmoid / Logistic 函数**

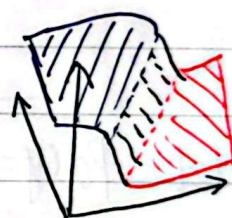
$$\sigma(z) = \frac{1}{1 + \exp(-z)}$$

$\Downarrow$

$$g(x) = \sigma(w^T x + b) = \frac{1}{1 + \exp\{-(w^T x + b)\}}$$

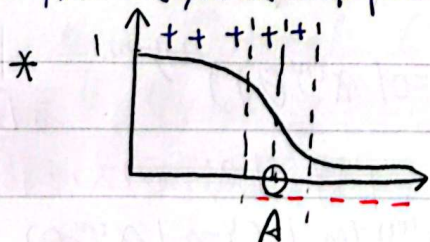
如何预测: predict +1 if probability > 0.5
 $\Rightarrow \sigma(w^T x + b) > 0.5$

$$\Rightarrow \frac{1}{1 + \exp\{-(w^T x + b)\}} > 0.5 \Rightarrow \exp\{-(w^T x + b)\} < 1$$



$\Rightarrow \underline{w^T x + b > 0}$, 可见依然是要训 w, b 参数, 预测时, $\Rightarrow +1$

但! Sigmoid 建模能引入 Uncertainty, 并且在数据不可分时, 建模仍有质量保障。*



虽然 A 点处有两种 sample point
 但因为给出的是 probability distribution
 因此有 quality guarantees

How do we learn a classifier?



$$g^{(i)} = \sigma(w^T x^{(i)} + b)$$

$$p(\text{data}) = \prod_i p(\text{data point } i) = \prod_i \begin{cases} g^{(i)} & \text{if } y^{(i)} = +1 \\ 1 - g^{(i)} & \text{else} \end{cases}$$

$$= \prod_i \left\{ (g^{(i)})^{1\{y^{(i)}=+1\}} (1-g^{(i)})^{1\{y^{(i)} \neq +1\}} \right\}$$

$$\Rightarrow \text{Loss} = \frac{1}{n} \sum_{i=1}^n - (1\{y^{(i)}=+1\} \log g^{(i)} + 1\{y^{(i)} \neq +1\} \log (1-g^{(i)}))$$

取负对数 (negative log likelihood loss) (g for guess, a for actual)

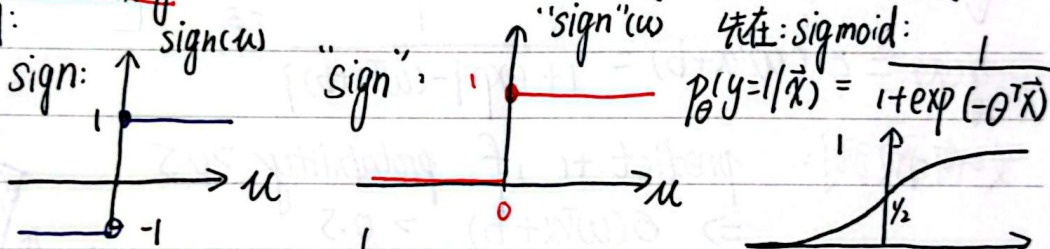
$$-L_{\text{NLL}}(g, a) = 1\{a=+1\} \log(g) + 1\{a \neq +1\} \log(1-g)$$

目标便是: $J(w, b) = \frac{1}{n} \sum_{i=1}^n L_{\text{NLL}}(\sigma(w^T x^{(i)} + b), y^{(i)})$

找 $\arg\min_{w, b} J(w, b)$, where: $-L_{\text{NLL}}(g, a) = 1\{a=+1\} \log g + 1\{a \neq +1\} \log(1-g)$

Back to Learning: Data: $D = \{\vec{x}^{(i)}, y^{(i)}\}_{i=1}^N, \vec{x} \in \mathbb{R}^M, y \in \{0, 1\}$.

原先预测用:



Model: $p_{\theta}(y=1|\mathbf{x}) = \frac{1}{1+\exp(-\theta^T \mathbf{x})} = g(\mathbf{x}) \Rightarrow p(y|\mathbf{x}, \theta) = \begin{cases} g(\mathbf{x}) & \text{if } y=1 \\ 1-g(\mathbf{x}) & \text{if } y=0 \end{cases}$

Learning: objective:

$$J(\theta) = \frac{1}{n} \sum_{i=1}^n -\log p(y^{(i)}|\mathbf{x}^{(i)}, \theta)$$

做 $\theta^* = \arg\min_{\theta} \ell(\theta)$, $\ell(\theta) = -\frac{1}{N} \log \prod_{n=1}^N p(y^{(n)}|\vec{x}^{(n)}, \theta)$

$$= -\frac{1}{N} \log \prod_{n=1}^N p(y=1|\mathbf{x}^{(n)}, \theta)^{y^{(n)}} (p(y=0|\mathbf{x}^{(n)}, \theta))^{1-y^{(n)}} \quad *$$

*: 若 $y^{(n)}=1$, 欲 $p(y=1|\mathbf{x}^{(n)}, \theta)$ 大; 反之, 欲 $p(y=0|\mathbf{x}^{(n)}, \theta)$ 大; 对象不同.

$$= -\frac{1}{N} \sum_{n=1}^N y^{(n)} \log p(y=1|\mathbf{x}^{(n)}, \theta) + (1-y^{(n)}) \log p(y=0|\mathbf{x}^{(n)}, \theta)$$

$$= -\frac{1}{N} \sum_{n=1}^N y^{(n)} \theta^T \mathbf{x}^{(n)} - \log(1 + \exp(\theta^T \mathbf{x}^{(n)}))$$



Gradient: Target: $J(\theta) = -\frac{1}{N} \sum_{n=1}^N y^{(n)} \theta^T x^{(n)} - \log(1 + \exp(\theta^T x^{(n)}))$

$$\begin{aligned} \nabla_{\theta} J(\theta) &= -\frac{1}{N} \sum_{n=1}^N [y^{(n)} \nabla_{\theta} (\theta^T x^{(n)}) - \nabla_{\theta} \log(1 + \exp(\theta^T x^{(n)}))] \\ &= -\frac{1}{N} \sum_{n=1}^N \left[y^{(n)} x^{(n)} - \frac{\exp(\theta^T x^{(n)})}{1 + \exp(\theta^T x^{(n)})} x^{(n)} \right] \\ &= \frac{1}{N} \sum_{n=1}^N x^{(n)} (p(y=1 | x^{(n)}, \theta) - y^{(n)}) \end{aligned}$$

$\frac{1}{1 + \exp(-\theta^T x^{(n)})}$

Then for learning with Gradients given, we can GD/SGD
 ☆: No close form solution! (无 closed form 解 for MLE 参数)

附: 以 LMS (Least Mean Square), $h_{\theta}(x) = \theta^T x$,

则 $J_{\theta} = \frac{1}{n} \sum_{i=1}^n (\theta^T x^{(i)} - y^{(i)})^2$, stochastically, (针对 1 sample)

$$\nabla_{\theta} J_{\theta}^{(i)*} = (\theta^T x^{(i)} - y^{(i)}) x^{(i)}, \text{ 则 } \theta_k \leftarrow \theta_k + \nabla_{\theta} J_{\theta}$$

$$= \theta_k + \lambda (h_{\theta}(x^{(i)}) - y^{(i)}) ; \nabla_{\theta} J_{\theta}^{(i)*} = x^{(i)} (p(y=1 | x^{(i)}, \theta) - y^{(i)})$$

发现 Linear Regression 的 SGD 也是: $\theta_k \leftarrow \theta_k + \lambda (h_{\theta}(x^{(i)}) - y^{(i)})$

*△: 再次 recall: SGD: for $i \in \text{shuffle}(\{1, 2, \dots, N\})$ do:

$$\theta \leftarrow \theta - \gamma \nabla_{\theta} J^{(i)}(\theta) \quad \text{☆ 注意这里是 } -, \text{ 且 } \gamma > 0!$$

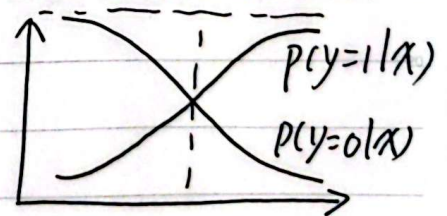
可见, 之前以 $\theta_k \leftarrow \theta_k + \lambda (h_{\theta}(x^{(i)}) - y^{(i)})$ 中的 λ 应是负数

Bayes Optimal Classifier:

若 $p^*(y | x)$ 这个先验信息被告知呢?

如: 测核酸时, $P(\text{阳}) \ll P(\text{阴})$

$$\text{则: } \ell(y, \hat{y}) = \begin{cases} 1000000, & y \neq \hat{y}, y = \text{阳} \\ 1000, & y \neq \hat{y}, y = \text{阴} \\ 0, & y = \hat{y} \end{cases}$$



则可考虑: Bayes Decision Rule: $\hat{y} = h(x) = \begin{cases} 1, & \text{if } P(y=1|x) \geq \alpha \\ 0, & \text{otherwise} \end{cases}$

若 $\alpha = 0.5$, 则 $1(y, \hat{y}) = \mathbb{I}(y \neq \hat{y})$

这样在 $\alpha \neq 0.5$ 下, 可以设计 asymmetric loss

带这种思想, 能用于修改 Log-Loss:

原: $J(\theta) = \frac{1}{n} \sum_{i=1}^n y^{(i)} \log P(Y=1|x^{(i)}, \theta) + (1-y^{(i)}) \log P(Y=0|x^{(i)}, \theta)$

如例中, 若 $y=1$, 但 $P(Y=1|x^{(i)}, \theta)$ 很小, 则 Loss 更大

则赋予权重: $J(\theta) = \frac{1}{n} \sum_{i=1}^n w_1 y^{(i)} \log P(Y=1|x^{(i)}, \theta) + w_2 (1-y^{(i)}) \log P(Y=0|x^{(i)}, \theta)$

$w_1/w_2 > 1$, 则 $y=1$ 预测错后果更严重! i.e., $\nabla_{\theta} J(\theta) \uparrow \Delta \theta \uparrow$

Recommendation Systems + Boosting

关于前面算法介绍, 详见另一文件

(unconstrained!) * 啥意思? 后面将会知道

Matrix Factorization (用于 Latent factor methods) (MF)

有许多方式可将一个 matrix 表示为多个 matrix 的乘法, 比如 SVD
而 MF 就是一个套入了梯度下降式学习的方法:

define model $\rightarrow$ define 目标 func $\rightarrow$ optimize with SGD

High level idea: rating matrix (R) $\Rightarrow U \times V$

where: $U \rightarrow$ users, $V \rightarrow$ items

$R \in \mathbb{R}^{m \times n}$, and R has rank $k \ll \min(m, n)$, then: Insight.
 $\exists U \in \mathbb{R}^{m \times k}$ and $V \in \mathbb{R}^{n \times k}$, s.t., $R = UV^T$

若 R rank 为 l , $l > k$ 呢? 则是 $\exists U (m \times k) V (n \times k)$, $R \approx UV^T$

这就是低秩分解的优点! (low rank matrix's rank k : $k \ll \min(m, n)$)

很幸运的是, User-Item 矩阵 R 被普遍认为低秩! 那么如何找 U, V ?

Unconstrained MF: Let $E = R - UV^T$

欲 E 接近零矩阵! 故定义 Objective Function:

$$J(U, V) = \frac{1}{2} \|E\|_F^2 = \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^n E_{ij}^2 = \frac{1}{2} \sum_{i=1}^m \sum_{j=1}^n (R_{ij} - U_{j \cdot} \cdot V_{\cdot i}^T)^2$$

问题在于, 这便要 R_{ij} 均已知; 但现实是: User 只给极少 item 打分!

即: 有 set: $Z = \{(i, j) : R_{ij} \text{ is known}\}$, $J = \frac{1}{2} \sum_{(i,j) \in Z} (R_{ij} - U_{j \cdot} \cdot V_{\cdot i}^T)^2$

则: $\nabla_{U_{j \cdot}} J_{ij}(U, V) = -E_{ij} V_{\cdot i}^T + \lambda U_{j \cdot}$ $+\frac{1}{2}(\|U\|_F^2 + \|V\|_F^2)$

$\nabla_{V_{\cdot i}^T} J_{ij}(U, V) = -E_{ij} U_{j \cdot} + \lambda V_{\cdot i}^T$ 正则化

$U_{j \cdot} \leftarrow U_{j \cdot} - \eta \nabla_{U_{j \cdot}} J_{ij}(U, V)$; $V_{\cdot i}^T \leftarrow V_{\cdot i}^T - \eta \nabla_{V_{\cdot i}^T} J_{ij}(U, V)$

上述便是用 SGD 优化! 每一次, 挑一个 $(i, j) \in Z$, 以更新参数



Trick: 若知 U, V^T 中一个, 则 solve 是十分简单的! 回想 $J(\theta) = \frac{1}{2} \sum_{i=1}^N (y^{(i)} - \theta^T x^{(i)})^2$

则可考虑: 初始化 U, V^T 后, while not converged:

• Fix V^T and solve U

• Fix U and solve V^T

SVD for collaborative filtering

SVD: $A = U \Lambda V^T$, where Λ : diagonal, U, V : orthogonal

考虑: 对于 User-item R matrix, 若已知:

$$R = Q \Sigma P^T$$

then truncate each of Q, Σ and P s.t.: Q, P has k columns.

Σ has $k \times k$ entries: $R \approx Q_k \Sigma_k P_k^T$

Then let $U \triangleq Q_k \Sigma_k, V \triangleq P_k$

$$\Rightarrow U, V = \underset{U, V}{\operatorname{argmin}} \frac{1}{2} \|R - UV^T\|_F^2$$

Theorem: R 是 fully observed 时, SVD 优化出的 U, V 与 MF 的, 是一样的

Non-negative MF:

通常公司收集的 R 矩阵中是无负数的! (如: 打星: 1~5 颗, 无负!)

故问题变成了: $U, V = \underset{U, V}{\operatorname{argmin}} \frac{1}{2} \|R - UV^T\|_F^2$

s.t. $U_{ij} \geq 0, V_{ij} \geq 0$

与先前介绍的 Unconstrained MF 不同, 此处有限制!

Ensemble method: 集成学习

Weighted Majority Algorithm: (example of ensemble method)

Given: pool of binary classifier

make new prediction

Goal: design a new learner that uses predictions of pool to

则加权 majority vote 流程如下:



初始化 classifier weights: $\alpha_t = 1$, $\forall t \in \{1, \dots, T\}$

对于每个样本 $(\vec{x}, y)$, do:

$$\hat{h}(x) = \text{sign} \left(\sum_{t=1}^T \alpha_t h_t(x) \right)$$

if $\hat{h}(x) \neq y$:

for each classifier $h_t \in \{1, \dots, T\}$ do:

if $h_t(x) \neq y$, $\alpha_t \leftarrow \beta \alpha_t$ hyperparameter

根据个体学习器生成方式, 目前 ensemble learning 可分两类:

- ① 个体学习器间存在依赖关系, 必须串行生成 $\Rightarrow$ boosting 家族
- ② $\dots$ 不存在强依赖关系, 也可并行化式 $\Rightarrow$ Random Forest

Adaboost: Δ 只作 = 分类 基模型个数

流程: Dataset $D = \{(x_1, y_1), \dots, (x_m, y_m)\}$, 基学习算法 $\mathcal{L}$

$$D_1(x) = y/m$$

$$x = [\vec{x}_1, \vec{x}_2, \dots, \vec{x}_m]$$

for $t = 1, \dots, T$ do:

$D_t(x)$: 样本权值分布

(寻找当前) $h_t = \mathcal{L}(D_t, D)$ 用 D_t 分布从 D 训练 t 个弱分类器 h_t

误差最小基 $\epsilon_t = P_{x \sim D_t}(h_t(x) \neq f(x)) \rightarrow$ h_t 误差估计

分类器为 h_t) (if $\epsilon_t > 0.5$, break) (若 > 0.5 , 则比随机还差!)

$$\alpha_t = \frac{1}{2} \ln \left(\frac{1 - \epsilon_t}{\epsilon_t} \right) \quad \text{规范化, 确保 } D_{t+1}(x) \text{ 也是分布}$$

$$D_{t+1}(x) = \frac{D_t(x)}{Z_t} \times \begin{cases} \exp(-\alpha_t) & \text{if } h_t(x) = f(x) \\ \exp(\alpha_t) & \text{if } \neq \end{cases}$$

$$= \frac{D_t(x)}{Z_t} \exp(-\alpha_t h_t(x) f(x))$$

end for

输出: $F(x) = \text{sign} \left(\sum_{t=1}^T \alpha_t h_t(x) \right)$

*: 用 D_t 分布, 可重采样出一个新数据集, 在这个数据集上, 优化 h_t 参数以 minimize loss. (题目中, 此步可能不进行)



理解: "加性模型". $H(x) = \sum_{t=1}^T \alpha_t h_t(x)$

↳ minimize 指数损失函数: $\text{lexp}(H|D) = \mathbb{E}_{x \sim D} [e^{-f(x)H(x)}]$

推导: $\frac{\partial \text{lexp}(H|D)}{\partial H(x)} = -e^{-H(x)} P(f(x)=1|x) + e^{H(x)} P(f(x)=-1|x).$

$\Delta: \frac{\partial e^{-H(x)}}{\partial H(x)} = -e^{-H(x)} \quad \frac{\partial e^{H(x)}}{\partial H(x)} = e^{H(x)} \quad = 0$

$$\begin{aligned} \text{lexp}(H|D) &= \sum_{i=1}^{|D|} D(x_i) (e^{-H(x_i)} \mathbb{I}(f(x_i)=1) + e^{H(x_i)} \mathbb{I}(f(x_i)=-1)) \\ &= \sum_{i=1}^{|D|} (e^{-H(x_i)} P(f(x_i)=1|x_i) + e^{H(x_i)} P(f(x_i)=-1|x_i)) \end{aligned}$$

$$\begin{aligned} \Rightarrow H(x) &= \frac{1}{2} \ln \frac{P(f(x)=1|x)}{P(f(x)=-1|x)}, \quad \text{sign}(H(x)) = \begin{cases} 1, & P(f(x)=1|x) > P(f(x)=-1|x) \\ -1, & \text{otherwise} \end{cases} \\ &= \underset{y \in \{-1, 1\}}{\operatorname{argmax}} P(f(x)=y|x). \end{aligned}$$

$$\begin{aligned} \text{补: } Z_t &= \sum_{i=1}^n D_t(i) \exp(-\alpha_t y_i h_t(x_i)) \\ &= e^{\alpha_t} \sum_{i=1}^n D_t(i) \mathbb{I}(y_i \neq h_t(x_i)) + e^{-\alpha_t} \sum_{i=1}^n D_t(i) \mathbb{I}(y_i = h_t(x_i)) \\ &= e^{\alpha_t} \epsilon_t + e^{-\alpha_t} (1 - \epsilon_t) \end{aligned}$$

训后: 最终:

$$\epsilon = \frac{1}{n} \sum_{i=1}^n \begin{cases} 1, & \text{if } y_i \left(\sum_{t=1}^T \alpha_t h_t(x_i) \right) \leq 0 \\ 0, & \text{else} \end{cases} \leq \frac{1}{n} \sum_{i=1}^n \exp\left(-y_i \left(\sum_{t=1}^T \alpha_t h_t \right)\right)$$

(iii) $D_{T+1}(i) = \frac{D_T(i)}{Z_T} \exp(-\alpha_T y_i h_T(x_i)).$ 左右连乘:

$$\epsilon \leq \frac{1}{n} \sum_{i=1}^n \exp\left(-y_i \left(\sum_{t=1}^T \alpha_t h_t(x_i) \right)\right) = \prod_{t=1}^T Z_t \cdot \sum_{i=1}^n D_{T+1}(i) = 1.$$

$\therefore \epsilon \leq \prod_{t=1}^T Z_t$, Upper bound $\downarrow$, 又 $Z_t = e^{\alpha_t} \epsilon_t + e^{-\alpha_t} (1 - \epsilon_t).$

$\therefore \frac{\partial Z_t}{\partial \alpha_t} = 0$ ($\frac{\partial Z_t}{\partial \alpha_t} > 0$, convex) $\Rightarrow \alpha_t = \frac{1}{2} \log \frac{1 - \epsilon_t}{\epsilon_t}.$

Campus

再代回 Z_t : $Z_t = 2\sqrt{\epsilon_t(1 - \epsilon_t)}$



Bagging & Random forest

Bagging: Bootstrap AGGregation: ①: Sample bagging
 Key: Repeatedly sample (with replacement!) a collection of training examples and train a model on that sample

流程: for $t = 1, \dots, T$ do
 for $s = 1, \dots, S$ do
 $i_s \sim \text{Uniform}(1, \dots, N)$
 $S_t = \{(x^{(i_s)}, y^{(i_s)})\}_{s=1}^S$ $\leftarrow$ bootstrap sample.
 $h_t = \text{train}(S_t)$ $\leftarrow$ classifier
 return $\hat{h}(x) = \text{aggregate}^*(h_1, \dots, h_T)$

注意, 是对每一个 h_t 都重
 ↓ 采样一份 S_t

* for classification: majority vote
 for regression: average

② Feature Bagging key: select subset of feature as well

for $t = 1, \dots, T$ do
 for $s = 1, \dots, S$ do
 $m_s \sim \text{Uniform}(1, \dots, M)$
 for $i = 1, \dots, N$ do
 $\hat{x}^{(i)} = [x_{m_1}^{(i)}, x_{m_2}^{(i)}, \dots, x_{m_s}^{(i)}]^T$ $\leftarrow$ sample feature
 $D_t = \{(\hat{x}^{(i)}, y^{(i)})\}_{i=1}^N$ $\leftarrow$ subspace
 $h_t = \text{train}(D_t)$
 return $\hat{h}(x) = \text{aggregate}(h_1, \dots, h_T)$

Random Forest: Key: Combine prediction of many diverse decision trees to reduce variability

If B r.v. all have variance σ^2 , Then $\frac{1}{B} \sum_{b=1}^B x^{(b)}$'s variance is: σ^2/B !



So we can combine (sample) bagging and a specific variant of feature bagging to train decision trees

基决策树

流程: ↓ 重复: 抽一个样本 (with replacement)

→ 对树的每个结点, 先从该结点属性集合中随机选一个子集, 再在子集中选最优属性。设|集合| = d , |子集| = k , 若 $k=d$: 则

vanilla decision tree splitting node; 若 $k=1$, 则随机选一个作结点

建议: $k = \log_2 d$

上述算法中的采样还有一个优点: 每个基学习器只用了 dataset 中 63.2% 的样本, 剩下 36.8% 可用作 validation set 来对泛化能力作 '包外估计' (out of bag estimate); 这便可用于超参优化!

Random forest 简单易实现, 且 performance 很不错, 而且收敛性与 Bagging 相似

$$*: \lim_{m \rightarrow \infty} \left(1 - \frac{1}{m}\right)^m \approx 0.368$$

Expectation Maximization (EM)

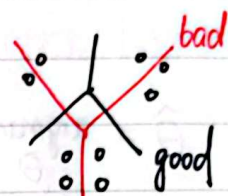
Unsupervised Learning: Kmeans & Gaussian Mixture Models (GMM)

考虑下述任务:

Clustering (informal goals): Goal: data 为 unlabeled, 将它们分成数组, 每组中点相似。

这个任务极为常用!

问题在于: 如何建模出 partitions 的好坏?

Input: unlabeled data: $D = \{x^{(i)}\}_{i=1}^N$, $x^{(i)} \in \mathbb{R}^M$ $\Rightarrow$ cluster centers: $C = [c_1, \dots, c_k]$ $c_j \in \mathbb{R}^M$ cluster assignments: $z = [z^{(1)}, z^{(2)}, \dots, z^{(N)}]$, $z^{(i)} \in \{1, \dots, k\}$

$$\text{Objective: } \hat{C} = \underset{C}{\operatorname{argmin}} \sum_{i=1}^N \min_j \|x^{(i)} - c_j\|_2^2$$

$$= \underset{C}{\operatorname{argmin}} \sum_{i=1}^N \min_{z^{(i)}} \|x^{(i)} - C_{z^{(i)}}\|_2^2$$

$$\Leftrightarrow \hat{C}, \hat{z} = \underset{C, z}{\operatorname{argmin}} \sum_{i=1}^N \|x^{(i)} - C_{z^{(i)}}\|_2^2 \quad \mathcal{J}(C, z)$$

Algorithm: 1) Initialize $C = \{c_1, \dots, c_k\}$ * 随机k个样本作初始均值向量

2) Repeat until Convergence:

a) for i in $\{1, \dots, N\}$:

$$z^{(i)} \leftarrow \underset{j}{\operatorname{argmin}} (\|x^{(i)} - c_j\|_2)^2 \leftarrow \text{每个 } x_i \text{ 选最近 } c_j$$

b) for j in $\{1, \dots, k\}$:

$$c_j \leftarrow \underset{C_j}{\operatorname{argmin}} \sum_{i: z^{(i)}=j} (\|x^{(i)} - c_j\|_2)^2$$

$$\text{此优化 i.e.: } \mu_i = \frac{1}{|C_i|} \sum_{x \in C_i} x, \quad C_i \leftarrow \mu_i$$

* 不止随机, 也可 Furthest Point Heuristic: 先随机选一个, 之后每一个选取尽可能离先前选的点们远! (但 outlier 处理是问题)

也可 K-means++: 先随机选 c_1 , 则 $j = 2, \dots, K$, pick c_j from distribution:

$$P(c_j = x^{(i)}) \propto \min_{j' < j} \|x^{(i)} - c_{j'}\|^2, \text{ where } x^{(i)} \text{ are not picked before}$$



EM

与Kmeans不同, 高斯混合聚类采用概率模型表达聚类。先介绍EM:

考虑 $p(X, Z | \theta) = \prod_{n=1}^N p(x_n, z_n)$, 其中:

observed: $X = \{x_n\}_{n=1}^N$; latent: $Z = \{z_n\}_{n=1}^N$

Goal: Estimate θ via MLE (or MAP).

z :

$$\hat{\theta} = \operatorname{argmax}_{\theta} \log p(X | \theta) = \operatorname{argmax}_{\theta} \log \sum_z p(X, z | \theta) \quad (\text{discrete})$$

$$= \operatorname{argmax}_{\theta} \log \int_z p(X, z | \theta) dz. \quad (\text{continuous})$$

$$\log p(X | \theta) = \log \sum_z p(X, z | \theta) = \log \sum_z q_n(z) \frac{p(X, z | \theta)}{q_n(z)}$$

$$\geq \sum_z q_n(z) \log \frac{p(X, z | \theta)}{q_n(z)} \quad (\text{since } \log(\cdot) \text{ is concave; Jensen})$$

$$= \sum_z q_n(z) \log p(X, z | \theta) - \sum_z q_n(z) \log q_n(z) \quad \text{constant! since irrelevant with } \theta$$

(if $q_n(z) = p(z | X, \theta)$, 则为等号)

故令 $q_n(z) = p(z | X, \theta)$ (独立): $\log p(X | \theta) = [\mathbb{E}[\log p(X, z | \theta)]] + \text{const}^*$

* lower-bound factor! EM就是最大化它!

EM流程: E: 以当前 θ^t 推后验 $P(z | X, \theta^t)$, 计算: 后验 $\uparrow$ 先验 $\uparrow \theta$

$$Q(\theta | \theta^t) = \mathbb{E}_{z | X, \theta^t} [\log p(z, X | \theta)] = \sum_z p(z | X, \theta^t) \log p(X, z | \theta)$$

M: $\theta^{t+1} = \operatorname{argmax}_{\theta} \{Q(\theta | \theta^t) + [\log p(\theta)]\}$ 若加: MAP; 不加: MLE

latent variable, cluster中, $|z|=k$ $\sum_{k=1}^K \pi_k = 1$ k 类

GMM: 假设生成 N 个点: $x_n, n = 1, 2, \dots, N$

首先从 K 个高斯中选一个。此选择基于混合成分先验概率 π 的:

$z_n \sim \text{Multinomial}(z_n | \pi)$ (多项式分布)

然后生成点服从 $N(x_n | \mu_k, \Sigma_k)$, 其中 μ_k 是第 k 个高斯的均值,

Σ_k 是第 k 个高斯的协方差 (suppose $z_n = k$)

Δ : 说白了就是在 $[\pi_i]_{i=1}^K$ 的概率下, 抽一个 z_i



上述描述了在有 π_i, μ_i, Σ_i 下, 如何抽 n 个点出来。

那么学 GMM 呢? 考虑:

observed: $X = \{x_1, \dots, x_N\}$ latent: $Z = \{z_1, \dots, z_N\}$

$z_i \in \{1, \dots, K\} \Rightarrow N$ 点分 K 类

则 complete-data log-likelihood:

$$\log p(X, Z | \theta) = \sum_{i=1}^N \log \pi_{z_i} + \log N(x_i | \mu_{z_i}, \Sigma_{z_i})$$

则 E-step: posterior: $\gamma_{ik} = p(z_i = k | x_i, \theta^{\text{old}})$

$$\gamma_{ik} = \frac{p(x_i | z_i = k, \theta^{\text{old}}) p(z_i = k | \theta^{\text{old}})}{p(x_i | \theta^{\text{old}})}$$

Bayes

$$= \frac{\pi_k^{\text{old}} N(x_i | \mu_k^{\text{old}}, \Sigma_k^{\text{old}})}{\sum_{j=1}^K \pi_j^{\text{old}} N(x_i | \mu_j^{\text{old}}, \Sigma_j^{\text{old}})}$$

Bring in Gaussian

M-step: $Q(\theta, \theta^{\text{old}}) = \mathbb{E}_{Z|X, \theta^{\text{old}}} [\log p(X, Z | \theta)]$

$$= \sum_{i=1}^N \sum_{k=1}^K \gamma_{ik} [\log \pi_k + \log N(x_i | \mu_k, \Sigma_k)]$$

先验

后验

, constraint: $\sum_k \pi_k = 1$

$$\mathcal{L} = Q(\theta, \theta^{\text{old}}) + \lambda (1 - \sum_{k=1}^K \pi_k)$$

$$\frac{\partial \mathcal{L}}{\partial \pi_k} = \sum_{i=1}^N \frac{\gamma_{ik}}{\pi_k} - \lambda = 0 \Rightarrow \pi_k \propto \sum_{i=1}^N \gamma_{ik}$$

$$\text{又 } \sum_k \pi_k = 1, \text{ 故 } \sum_{k=1}^K \sum_{i=1}^N \gamma_{ik} = 1$$

$$\begin{cases} \pi_k^{\text{new}} = \frac{1}{N} \sum_{i=1}^N \gamma_{ik} \\ \mu_k^{\text{new}} = \frac{\sum_{i=1}^N \gamma_{ik} x_i}{\sum_{i=1}^N \gamma_{ik}} \\ \Sigma_k^{\text{new}} = \frac{\sum_{i=1}^N \gamma_{ik} (x_i - \mu_k^{\text{new}})(x_i - \mu_k^{\text{new}})^T}{\sum_{i=1}^N \gamma_{ik}} \end{cases}$$

$$\text{同理: } \frac{\partial Q}{\partial \mu_k} = \sum_{i=1}^N \gamma_{ik} \sum_k^{-1} (x_i - \mu_k) = 0$$

$$\frac{\partial Q}{\partial \Sigma_k^{-1}} = \frac{1}{2} \sum_{i=1}^N \gamma_{ik} [\Sigma_k - (x_i - \mu_k)(x_i - \mu_k)^T] = 0$$



附: PCA: 主成分分析 $\rightarrow \in \mathbb{R}^d$

Input: $D = \{x_1, x_2, \dots, x_m\}$, 低维空间维数 d'

Algorithm: 中心化: $x_i \leftarrow x_i - \frac{1}{m} \sum_{i=1}^m x_i$

计算协方差矩阵: XX^T

特征值分解

取最大 d' 个特征值对应的 eigenvector $w_1, w_2, \dots, w_{d'}$

$\Rightarrow$ Output $W^* = (w_1, w_2, \dots, w_{d'}) \in \mathbb{R}^{d' \times d}$

$x_i \cdot W^{*T}$ 后便进入低维空间

PCA 是最常用的降维方法!

