

11. Методи за обучение. Обучение с учител - параметрични и непараметрични подходи. Линеен дискриминант на Фишер. Обучение без учител – minimax, kMeans, Izodata [5, 12].

(Виж материалите страница ПР_Част1 44 – 52, , <http://www2.it.lut.fi/kurssit/04-05/010586001/lectures.html> - лекции 3,4 и 5)

Методи за обучение

Теорията на Бейс за статистическа класификация дава прави добра класификация само ако вероятностната плътност на класовете е известна.

Обикновено тя не е известна и тогава се налага да се направят оценки на базата на някакви данни и някакво обучение.

Методи за обучение на класификатора се делят на две групи – с учител и без учител.

- С учител
Ако имаме класове $C_1, C_2, \dots, C_n$ и сега от тях изберем чрез експерименти T_i – множество от елементи от C_i – тия клас. Като:

$$T = \bigcup_{i=1}^n T_i$$

T се нарича тестово множество. Освен това както вече споменахме:

$$T_i \subset C_i \subset W$$

За всеки пример x от тестовото множество можем да дефинираме 2-ка (x, i) , която да указва, че примера $x \in T_i$.

- Самообучение $x \in T \subset W$ и знаем само това на входа.

Интуитивни методи

1. k -те център – имаме k на брой еталонни обекта, някаква метрика например Евклидова или Статистическа(Мехаланобис) и представянето на x :

$$x_{rep} = \frac{1}{|T_i|} \sum_{x_i \in T_i} x_i$$

Изолиниите на гаусиана означават събития с еднаква вероятност. За разстоянието по Мехаланобис може да е изпълнено $D_M(x, \mu) = D_M(y, \mu)$, докато за евклидовото $D_E(x, \mu) \leq D_E(y, \mu)$. В зависимост от статистиката разстоянието по Мехаланобис разтегля или смачква евклидовото разстояние. Основната идея все пак си остава:

$$y = \text{Argmin } D(x, x_{rep})$$

2. NN

За всяко x и x_i , където $x_i \in T$ намираме $y = \text{Argmin } D(x, x_i)$.

Проблем би било, ако имаме такова разположение на класовете, че x да е по-близо до центъра на един клас, но да се намира в друг.



3. kNN (k-Nearest Neighbour)

Отново имаме представители от всички класове, но този път не само по един.

- Пресмятат се всички налични разстояния
- Сортират се
- Избират се k – на брой на близки разстояния

$$y = \text{Argmax } k_i, \text{ където } i = 1, 2, \dots, n$$

k_i е броят на елементите от класа T_i в избраните k най-близки. (Може да се зададе и по-сложна функция от чистия сбор)

Необходими данни:

- n -те точки за центровете (или други)
- обучаващите примери - $|T|$ на брой.

Методи за обучение с учител

При неясна вероятностна плътност, но знание за разпределението. Остава чрез настройване на параметрите и тяхна оценка да достигне до правдоподобна класификация.

При настройване на параметри казваме, че имаме така наречения параметричен подход, защото се опитваме да обучим системата да намери правилните параметри.

Нека да имаме да дефинираме модела на разпределение като $M(\theta_1, \theta_2, \dots, \theta_n)$ (за пример може да бъде даден гаусиан $N(\mu, \Sigma, x)$ – където μ и Σ могат да бъдат настроени според някакви тестови данни).

Един параметричен подход е подхода за максимално правдоподобие.

Какво имаме до момента – $p(x|C_i)$ има нормално разпределение, всеки един клас е описан от параметричен вектор от модела – на C_i – θ_i и тестово множество T .

Като тестовите примери са независими.

Въвежда ме функция на правдоподобие:

$$p(x|\theta) = p(x_1, x_2, \dots, x_n|\theta) = \prod_{k=1}^N p(x_k|\theta)$$

Прилагам стандартната стъпка за подхода на максималното правдоподобие логаритмуване на гаусиана и получаваме:

$$\frac{\sum_{x \in T_i} (x_1, x_2, \dots, x_n)}{|x_i|} = (\mu_1, \mu_2, \dots, \mu_n)$$

За матрицата Σ имаме:

$$\Sigma = \left| |\sigma_{ij}| \right| = \sum_{x \in T_i} (\mu_i - x_i)(\mu_j - x_j)$$

След това използваме класическия подход на Бейс, при който приемаме, че $p(\theta)$ има фиксирано разпределение.

$$\sigma_{ij} = \frac{\Sigma}{|T_i|-1} - \text{също състоятелна, но отместена оценка.}$$

Непараметрични

Когато не знаем дори функцията на разпределението и нямаме как да приложим параметричен подход. Тогава прилагаме така наречения непараметричен подход.

При него единственото, което знаем са тестовите примери.

Можем да приложим Хистограмен подход – в двумерното пространство разделяме абсисата на равни интервали с дължина h .

Вероятността за всяко една парче от хистограмата е:

$$P \approx \frac{k}{N}$$

Където k е броя на примерите вътре в парчето от хистограмата а N броя на всички примери.

Необходимо допускане е, че функцията на плътността е еднаква в рамките на всеки правоъгълник от хистограмата. Така получаваме:

$$\hat{p}(x) \equiv \hat{p}(\hat{x}) \approx \frac{k}{hN}, |x - \hat{x}| \leq \frac{h}{2}$$

Оценката изглежда добра при малко h и голям брой примери N .

Parzen Windows

Същия подход само, че за многомерно пространство. Тук вместо правоъгълници със страна h , ще използваме хиперкубове със страна h и обем V_l – където l измерението на пространството.

При избор на x построяваме хиперкуба с център x и изчисляваме елементите в него:

$$n = |T_l|, k_i \in T_l,$$

$$p(x) = \frac{k_i}{nV_l}$$

n – броя на всички експерименти

k_i – броя на тези, които са вътре в куба

V_l – обема на куба

Ако имаме празни клетки то имаме нужда от още експерименти.

Нека дефинираме функцията $\varphi(x) = \begin{cases} 1, & |x - t_i| \leq h/2 \\ 0, & \text{иначе} \end{cases}$, където t_i е елемент на T_l .

$\Rightarrow k_i = \sum_{j=1}^n \varphi\left(\frac{x_j - x}{h}\right)$, тъй като φ е пръксната това води до прекъснатост на $p(x)$, затова можем да дефинираме φ по-добре:

$\Rightarrow \varphi(x) \geq 0$ и $\int \varphi(x) dx = 1$, това е легитимна оценка на функция на вероятностната плътност. Такива функции са известни като kernel функции или Parzen Windows. Типичен пример е Гаусиан – $N(0, I)$.

Избора на h е доста сериозен проблем при ограничен брой тестови данни.

За това може да се използва подход като при k -nearest neighbor:

Вместо да се броят примерите в хиперкубове се избират k на брой най-близки примери и се изчислява обема на получената фигура.

Така в по-гъстите участъци обема ще е по-малък, а в по-редките по-голям.

$$p(x) = \frac{k}{nV(x)}$$

Линеен дискриминант на Фишер

За да се определи един класификатор, първо трябва да изберем подходяща

проекционна права, която да разделя класовете по възможно най-добрия начин.

Линейната дискриминатна функция, която отговаря за проектирането може да бъде представена във вида:

$$y = w^T x$$

Където дължината на вектора w е ограниче до $\|w\|=1$. За да измери разграничеността на два класа Фишер въвежда:

$$J(w) = \frac{S_B^2}{S_W^2}$$

Където S_B е между класова дисперсия, а S_W е вътре класова дисперсия.

Намирайки проекционната права L ние можем да построим построим средство дискриминантната равнина между класовете.

Избираме за модели гаусиани, които проектираме върху L. Така задачата от n-мерно пространство се свежда до 1-мерно. Метода е средно между линеен дискриминантен анализ и вероятностен подход.

$$S_B^2 = w_1(x - \mu_1)^2 + w_2(x - \mu_2)^2$$

$$S_w^2 = \sigma_1^2 + \sigma_2^2$$

Критерия за оптималност на Фишер е $J(w)$ да е максимално.

При еднаква априорна вероятност за класовете $\Rightarrow w_1 = w_2$.

$$\Rightarrow J(w) = \frac{(\mu_1 - \mu_2)^2}{\sigma_1^2 + \sigma_2^2}$$

Като разпишем за n-мерното пространство получаваме:

$$J(w) = \frac{w^T S_B w}{w^T S_w w}$$

Кое е максимално при $W = (\mu_1 - \mu_2) S_w^{-1}$, а ако са известни ковариационни матрици проблема се свежда до проблем за собствените стойности.

Непараметрични методи за определяне броя на класовете в T

Minimax

Дадено ни е само множеството T и някаква метрика ($D(N_0, N_1)$ – евклидово разстояние).

- 1) Избираме произволно ядро – N_0 е T (потенциално ядро на клас C_0)
- 2) N_1 е T, такова че $D(N_1, N_0) = \max$ (потенциално 2-ро ядро)
- 3) Разделяме T на два класа $\{N_0\}$ и $\{N_1\}$ по $\min D$.
- 4) N_2 е T, такова че като изберем по един кандидат от C_0 и C_1 , то N_{20} е C_0 – кандидат 1 е $D(N_{20}, N_0) = \max$.
 N_{21} е C_1 – кандидат 2 е $D(N_{21}, N_1) = \max$.
 $N_2 = \min \{N_{21}, N_{20}\}$
- 5) Разпределяме T в k класа по $\min D(N_i, x)$ за $i=1,2,\dots,k$.
- 6) За N_{k+1} имаме $N_{k+1} = \text{Argmin}_{\{N_i\}} (\text{Argmax } D(N_i, N_{k+1}))$
- 7) Остава само да кажем кога спираме – когато средното разстояние м/у класовете е под някакъв праг:

$$G_k = \frac{1}{\frac{k(k-1)}{2}} \sum_{j < i < k}^k D(N_i^k, N_j^k)$$

Спираме при $G_{k+1} < \frac{1}{2} G_k$.

Метод на k-те среди

Нека имаме вече k на брой ядра.

- 1) Правим разпределение в T на k класа $C_i = \{N_i\}$ за $i=1,2,\dots,k$ по минимално разстояние.
- 2) За всяко C_i изчисляваме новите центрове $\tilde{N}_i = \frac{1}{|N_i|} \sum_{x_i \in C_i} x_i$
- 3) Или спираме или се връщаме на 1 с новите центрове.

IZODATA

Може не само да увеличаваме, но и да намаляваме ядрата.

Имаме k ядра и 2 вида стъпки – за нарастване на ядрата и за съкращаването им.

- 1) Преразпределяне на T на k класа $C_i = \{N_i\}$ за $i=1,2,\dots,k$ по минимално разстояние.
- 2) Стъпка за намаляване на k – обединение на 2 класа.

- a. Намираме двата най-близки класа $\min D(N_i, N_j)$
 - b. Ако $D(N_i, N_j) < \epsilon$ – обединяваме иначе не, като е да кажем, че някаква допустима дисперсия
- 3) Разбиване на класа C_i на два класа – C_{i0} и C_{i1}
- Имаме статистиката на класа и ковариантната матрица. Можем да я приведем в диагонална форма и да подредим сигмите по големина:

$$\begin{matrix} \sigma_{i1} \\ \sigma_{i2} \\ \sigma_{in} \end{matrix}$$

Ако се окаже, че в тази редица има голям скок в разликите м/у сигмита ще е основателно да направим разделението иначе няма смисъл.