



1

Vari. Aleatorias Continuas

Fun. de Densidad de Probabilidad

$$f_X(x)$$

Propiedades

$$f_X(x) \geq 0$$

$$\int_{-\infty}^{\infty} f_X(x) dx = 1$$

$$P(X \leq a) = F_X(a) = \int_{-\infty}^a f_X(x) dx$$

$$P(x_1 < x < x_2) = \int_{x_1}^{x_2} f_X(x) dx$$

$$f_X(x) = \frac{dF_X(x)}{dx}$$

2

Fun. de Den. de Prob. Conjunta

$$f_{X,Y}(x,y)$$

$$f_X(x) = \int_{-\infty}^{\infty} f_{X,Y}(x,y) dy$$

$$f_Y(y) = \int_{-\infty}^{\infty} f_{X,Y}(x,y) dx$$

$$f_{X|Y}(x|y) = \frac{f_{X,Y}(x,y)}{f_Y(y)} \quad f_Y(y) > 0$$

$$f_{Y|X}(y|x) = \frac{f_{X,Y}(x,y)}{f_X(x)} \quad f_X(x) > 0$$

$$f_{Y|X}(y|x) = \frac{f_{Y|X}(x|y) f_Y(y)}{\int_{-\infty}^{\infty} f_{Y|X}(x|y) f_Y(y) dy} \quad (\text{Bayes}) \quad f_X(x) \neq 0$$

Promedios Estadísticos

$$E\{g(x)\} = \int_{-\infty}^{\infty} g(x) f_X(x) dx$$

$$E\{g(x,y)\} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} g(x,y) f_{X,Y}(x,y) dx dy$$

$$\rho_{XY} = \frac{E\{X - \mu_X\}(Y - \mu_Y)\}}{\sigma_X \sigma_Y}$$

3

Estimación de Parámetros

Diagram: A vertical axis labeled 'Espacio del Parámetro' (Parameter Space) has a point 'a'. A horizontal axis labeled 'Espacio de observación' (Observation Space) has a point 'R'. A blue arrow labeled 'Mapeo Probabilístico' (Probabilistic Mapping) points from 'a' to 'R'. A red arrow labeled 'Regla de Estimación' (Estimation Rule) points from 'R' to 'â(R)'.

a y $\hat{a}(R)$ son variables continuas

I. Caso: a es una variable aleatoria

A cada par $[a, \hat{a}(R)]$ se asigna un costo $C(a, \hat{a})$

Error de la estimación $a_e(R) \triangleq \hat{a}(R) - a$

Costo depende del error $C(a_e)$

Three graphs showing different cost functions $C(a_e)$ versus error a_e :

- Quadratic cost: $C(a_e) = a_e^2$ (labeled '(mse) error medio cuadrático')
- Absolute cost: $C(a_e) = |a_e|$ (labeled 'error absoluto')
- Uniform cost: $C(a_e) = \begin{cases} 0 & |a_e| \leq \frac{\Delta}{2} \\ 1 & |a_e| > \frac{\Delta}{2} \end{cases}$ (labeled 'función de costo uniforme')

4

Función de costo debe

- Satisfacer los requerimientos del usuario
- Proporcionar resolver el problema

$P_0(A)$ Densidad de probabilidad a priori (conocida)

Si $P_0(A)$ es desconocida, entonces se sigue un procedimiento MINIMAX

$$\text{Riesgo: } R \triangleq E\{C[a, \hat{a}(R)]\} = \int_{-\infty}^{\infty} dA \int_{-\infty}^{\infty} C[A - \hat{a}(R)] P_{A|R}(A, R) dR$$

Estimación de Bayes minimiza el riesgo:

a) Criterio del error medio cuadrático:

$$R_{ms} = \int_{-\infty}^{\infty} dR \int_{-\infty}^{\infty} dA [A - \hat{a}(R)]^2 P_{A|R}(A, R)$$

$$\text{donde } P_{A|R}(A, R) = P_1(R) P_{A|R}(A|R)$$

$$\therefore R_{ms} = \int_{-\infty}^{\infty} dR P_1(R) \int_{-\infty}^{\infty} dA [A - \hat{a}(R)]^2 P_{A|R}(A|R)$$

Derivando con respecto a $\hat{a}(R)$:
 Para minimizar R_{ms} hay que minimizar esta integral,
 o sea, la varianza a posteriori o varianza condicional

$$\frac{d}{d\hat{a}} \int_{-\infty}^{\infty} dA [A - \hat{a}(R)]^2 P_{A|R}(A|R) = -2 \int_{-\infty}^{\infty} dA A P_{A|R}(A|R) + 2 \hat{a}(R) \int_{-\infty}^{\infty} dA P_{A|R}(A|R) \underset{=1}{dA}$$

Igualando a cero, obtenemos $\hat{a}_{ms}(R)$:

$$\hat{a}_{ms}(R) = \int_{-\infty}^{\infty} dA A P_{A|R}(A|R)$$

media de la densidad a posteriori,
media condicional.

5

b) Criterio del error absoluto:

$$R_{abs} = \int_{-\infty}^{\infty} dR P_1(R) \int_{-\infty}^{\infty} dA [A - \hat{a}(R)] P_{A|R}(A|R)$$

Para minimizar la integral interior:

$$I(R) = \int_{-\infty}^{\infty} dA [\hat{a}(R) - A] P_{A|R}(A|R) + \int_{-\infty}^{\infty} dA [A - \hat{a}(R)] P_{A|R}(A|R)$$

Derivando e igualando a cero:

$$\int_{-\infty}^{\infty} dA P_{A|R}(A|R) = \int_{-\infty}^{\infty} dA P_{A|R}(A|R)$$

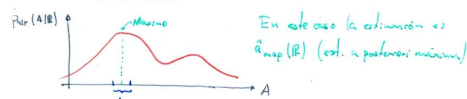
Mediana de la densidad a posteriori

c) Función de costo uniforme

$$R_{unf} = \int_{-\infty}^{\infty} dR P_1(R) \left[1 - \int_{\hat{a}_{unf}(R) - \Delta/2}^{\hat{a}_{unf}(R) + \Delta/2} P_{A|R}(A|R) dA \right]$$

Para minimizar R_{unf} hay que maximizar esta integral

Si $\Delta \rightarrow 0$, entonces:



$\hat{a}_{map}$ es el punto donde $P_{A|R}(A|R)$ es máxima
 Si el máximo es interior a A y la $P_{A|R}(A|R)$ es continua, entonces una
 condición necesaria para encontrar $\hat{a}_{map}(R)$ se puede obtener así:

$$\left. \frac{\partial \ln P_{A|R}(A|R)}{\partial A} \right|_{A=\hat{a}(R)} = 0 \quad \text{Ecuación MAP}$$

6

En el caso de la ecuación MAP, hay que notar que la solución es un máximo absoluto

separando las contribuciones del vector R y de la información a priori de $P_{\text{air}}(A|R)$

$$P_{\text{air}}(A|R) = \frac{P_{\text{ria}}(R|A)P_0(A)}{P_r(R)}$$

$$\ln P_{\text{air}}(A|R) = \ln P_{\text{ria}}(R|A) + \ln P_0(A) - \ln P_r(R)$$

independiente de A

$$L(A) \triangleq \underbrace{\ln P_{\text{ria}}(R|A)}_{\text{dependencia de } R \text{ de } A} + \underbrace{\ln P_0(A)}_{\text{conocimiento a priori}}$$

Ecuación MAP:

$$\left. \frac{\partial L(A)}{\partial A} \right|_{A=\hat{A}(R)} = \left. \frac{\partial \ln P_{\text{ria}}(R|A)}{\partial A} \right|_{A=\hat{A}(R)} + \left. \frac{\partial \ln P_0(A)}{\partial A} \right|_{A=\hat{A}(R)} = 0$$

Ejemplo $r_i = a + n_i \quad i=1, 2, \dots, N$

a es una var. aleat. Gaussiana $N(0, \sigma_a^2)$ y n_i son variables aleat. Gaussianas independientes $N(0, \sigma_n^2)$

$$P_{\text{ria}}(R|A) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma_n} e^{-\frac{(r_i - A)^2}{2\sigma_n^2}}$$

$$P_0(A) = \frac{1}{\sqrt{2\pi}\sigma_a} e^{-\frac{A^2}{2\sigma_a^2}}$$



7

4) Encontrar $\hat{a}_{\text{ms}}(R)$

Necesitamos $P_{\text{air}}(A|R) = \frac{P_{\text{ria}}(R|A)P_0(A)}{P_r(R)} = \frac{P_{\text{ria}}(R|A)}{P_r(R)} \cdot P_0(A)$ de que hace que $\int_{-\infty}^{\infty} P_{\text{air}}(A|R) dA = 1$

$$P_{\text{air}}(A|R) = \frac{\left(\frac{1}{\sqrt{2\pi}\sigma_n} \right)^N \frac{1}{\sqrt{2\pi}\sigma_a}}{P_r(R)} e^{-\frac{1}{2} \left[\frac{\sum_{i=1}^N (R_i - A)^2}{\sigma_n^2} + \frac{A^2}{\sigma_a^2} \right]}$$

Rearreglando:

$$P_{\text{air}}(A|R) = k(R) e^{-\frac{1}{2\sigma_n^2} \left[A - \frac{\sigma_n^2}{\sigma_n^2 + \sigma_a^2/N} \left(\frac{1}{N} \sum_{i=1}^N R_i \right) \right]^2}$$

f. de densidad gaussiana.

$$\text{donde } \sigma_f^2 \triangleq \left(\frac{1}{\sigma_n^2} + \frac{N}{\sigma_a^2} \right)^{-1} = \frac{\sigma_n^2 \sigma_a^2}{N\sigma_n^2 + \sigma_a^2}$$

Varianza a posteriori

$\hat{a}_{\text{ms}}(R)$ es la media condicional:

$$\hat{a}_{\text{ms}}(R) = \frac{\sigma_n^2}{\sigma_n^2 + \sigma_a^2/N} \left(\frac{1}{N} \sum_{i=1}^N R_i \right)$$

$\hat{R}$ nuevo m.s.e. es la varianza a posteriori

Observaciones:

1) $\{R\} = \sum_{i=1}^N R_i \Rightarrow$ Estadística suficiente

a) Si $\sigma_n^2 \ll \frac{\sigma_a^2}{N} \Rightarrow$ conocimiento a priori es mayor que los datos observados y el valor estimado es parecido al (valor a priori) (0)

b) Si $\sigma_n^2 \gg \frac{\sigma_a^2}{N} \Rightarrow$ conocimiento a priori es de poco valor, y el valor estimado son los datos recibidos

c) En el límite: $\lim_{\sigma_n^2 \rightarrow 0} \hat{a}_{\text{ms}}(R) = \frac{1}{N} \sum_{i=1}^N R_i$ Promedio aritmético de R_i

8

b) $\hat{a}_{\text{map}} : P_{\text{air}}(A|\mathbb{R})$ es Gaussiana, Por tanto, el valor máximo ocurre en la media condicional. Entonces:

$$\hat{a}_{\text{map}}(\mathbb{R}) = \hat{a}_{\text{ms}}(\mathbb{R})$$

En funciones Gaussianas, la mediana condicional también ocurre en la media condicional. Por lo tanto:

$$\hat{a}_{\text{abs}}(\mathbb{R}) = \hat{a}_{\text{ms}}(\mathbb{R})$$

Cuando la función de costo es simétrica, convexa hacia arriba y no decreciente, y la función de densidad a posteriori $P_{\text{air}}(A|\mathbb{R})$ es simétrica con respecto a su media condicional, entonces, el valor estimado de $\hat{a}$ que minimiza cualquier función de costo, dentro de la clase mencionada, es igual a $\hat{a}_{\text{ms}}$.

9

Teorema de Bayes para Clasificación

- Sea "x" un vector de características y "Ci" una clase:

$$P(C_i|x) = \frac{P(x|C_i)P(C_i)}{P(x)}$$



10

- Si tenemos dos clases, C_1 y C_2 , y queremos saber a cuál clase pertenece x , simplemente usamos la anterior fórmula y vemos cuál de las dos clases tiene mayor probabilidad.

$$P(C_1|x) = \frac{P(x|C_1)P(C_1)}{P(x)}$$

$$P(C_2|x) = \frac{P(x|C_2)P(C_2)}{P(x)}$$

Se observa que como el denominador es el mismo, entonces el denominador no importa y típicamente solo se calcula el numerador y se ve cuál es más grande.

11

- Entonces el Clasificador solo necesitaría calcular esto para clasificar a un vector "x": $P(C_i|x) \propto P(x|C_i)P(C_i)$

- Y, por tanto, en el entrenamiento, solamente se necesita obtener:

$$P(x|C_i) \quad P(C_i)$$



12

Bayes Ingenuo (*Naive Bayes*)

Imaginando que "x" es un vector de 3 componentes,
tendríamos la siguiente probabilidad conjunta:

$$P(x|C_i) = P(x_1, x_2, x_3|C_i)$$



13

- Y recordando la expansión de la probabilidad conjunta en probabilidades condicionales, tendríamos:

$$\begin{aligned} P(x_1, x_2, x_3|C_i) \\ = P(x_1|C_i) P(x_2|x_1, C_i) P(x_3|x_2, x_1, C_i) \end{aligned}$$

14

- Intentar calcular las probabilidades para todo x_1 , x_2 y x_3 , suele resultar muy complicado, por lo que **se asume que x_1 , x_2 y x_3 son independientes.**
- Asumiendo que son independientes la anterior fórmula se simplifica:

$$P(x_1, x_2, x_3 | C_i) = P(x_1 | C_i) P(x_2 | x_1, C_i) P(x_3 | x_2, x_1, C_i) \\ = P(x_1 | C_i) P(x_2 | C_i) P(x_3 | C_i)$$

15

Bayes Ingenuo



- A este clasificador se le llama **Bayes Ingenuo**, ya que resulta ingenuo pensar que los elementos del vector de características son independientes.
- Sin embargo se simplifica mucho la fórmula, haciendo que sea uno de los clasificadores más rápidos para entrenar.

16

- Generalizando a un vector x de tamaño n , simplemente tendríamos:

$$P(x|C_i) = \prod_{j=1}^n P(x_j|C_i)$$

- Y entonces en el entrenamiento tendríamos que calcular

$$P(x_j|C_i)$$

para toda j para todo i .

17

Naive Bayes

Resumiendo:

$$p(C_j | x_1, x_2, \dots, x_d) \propto p(x_1, x_2, \dots, x_d | C_j) p(C_j)$$

$$p(X | C_j) \propto \prod_{k=1}^d p(x_k | C_j)$$

$$p(C_j | X) \propto p(C_j) \prod_{k=1}^d p(x_k | C_j)$$



18



Ejemplo: Clasificador de SPAM

Training Set: Correos preprocesados

Tenemos 8 correos, 3 son spam y 5 no son spam. Tras quitar palabras duplicadas y palabras que no interesan, quedan así:

Spam	No Spam
oferta es secreto	practica deportes hoy
click link secreto	fue practica deportes
link deportes secreto	evento deportes secreto
	deportes es hoy
	deportes cuesta dinero

19

Entrenamiento (Contar)

Palabra	Spam	No Spam
Oferta	1	0
Es	1	1
Secreto	3	1
Click	1	0
Link	2	0
Deportes	1	5
Practica	0	2
Hoy	0	2
Fue	0	1
Evento	0	1
Cuesta	0	1
Dinero	0	1
TOTAL:	9	15

20

Entrenamiento (Calcular Probabilidades)

Palabra	P(Palabra Spam)	P(Palabra No Spam)
Oferta	1 / 9	0
Es	1 / 9	1 / 15
Secreto	3 / 9	1 / 15
Click	1 / 9	0
Link	2 / 9	0
Deportes	1 / 9	5 / 15
Practica	0	2 / 15
Hoy	0	2 / 15
Fue	0	1 / 15
Evento	0	1 / 15
Cuesta	0	1 / 15
Dinero	0	1 / 15
TOTAL:	1	1

21

Entrenamiento (Calcular Probabilidades)

- Así ya obtuvimos todas las $P(x_j | C_i)$
- Luego para obtener $P(C_i)$ recordamos que teníamos solo 8 correos, 3 eran de spam y 5 no eran de spam, entonces:

$$P(\text{Spam}) = 3/8$$

$$P(\text{No Spam}) = 5/8$$

22

Ejemplo de Predicción

- Tenemos un mensaje M, lo preprocesamos y terminamos con "secreto es deportes"
- Queremos clasificarlo, así que buscaremos calcular:

$$P(\text{Spam}|\text{secreto es deportes}) \propto P(\text{secreto es deportes}|\text{Spam}) P(\text{Spam})$$

$$P(\text{No Spam}|\text{secreto es deportes}) \propto P(\text{secreto es deportes}|\text{No Spam}) P(\text{No Spam})$$

23

- Como estamos usando Bayes Ingenuo, calculamos las probabilidades de cada palabra de manera independiente.

$$P(\text{secreto es deportes} | \text{Spam}) P(\text{Spam})$$

$$= P(\text{secreto}|\text{Spam}) P(\text{es}|\text{Spam}) P(\text{deportes}|\text{spam}) P(\text{Spam})$$

$$= (3/9)(1/9)(1/9)(3/8)$$

24

- Haciendo los cálculos para ambas clases, se obtiene:

$$P(\text{secreto es deportes} | \text{Spam}) P(\text{Spam}) = (3/9)(1/9)(1/9)(3/8) = 0.00154$$

$$P(\text{secreto es deportes} | \text{No Spam}) P(\text{No Spam}) = (1/15)(1/15)(5/15)(5/8) = 0.00092$$

Por lo tanto se clasifica como Spam

25

Problema: Datos no vistos

- Si tenemos "Oferta practica secreto" sucedería lo siguiente:

$$P(\text{oferta practica secreto} | \text{Spam}) P(\text{Spam}) = (1/9)(0)(3/9)(3/8) = 0$$

$$P(\text{oferta practica secreto} | \text{No Spam}) P(\text{No Spam}) = (0)(2/15)(1/15)(5/8) = 0$$

¡No podría clasificarlo!

26

Solución: "Smoothing"

Palabra	P(Palabra Spam)	P(Palabra No Spam)
Oferta	$(1 + 1) / (9 + 12)$	$(0 + 1) / (15 + 12)$
Es	$(1 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
Secreto	$(3 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
Click	$(1 + 1) / (9 + 12)$	$(0 + 1) / (15 + 12)$
Link	$(2 + 1) / (9 + 12)$	$(0 + 1) / (15 + 12)$
Deportes	$(1 + 1) / (9 + 12)$	$(5 + 1) / (15 + 12)$
Practica	$(0 + 1) / (9 + 12)$	$(2 + 1) / (15 + 12)$
Hoy	$(0 + 1) / (9 + 12)$	$(2 + 1) / (15 + 12)$
Fue	$(0 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
Evento	$(0 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
Cuesta	$(0 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
Dinero	$(0 + 1) / (9 + 12)$	$(1 + 1) / (15 + 12)$
TOTAL:	1	1

27

En SCIKIT LEARN



- En scikit learn hay 3 tipos de clasificadores de Bayes Ingenuo:

Clasificador	Tipo de Features en X
Gaussiano	Continuo
Multinomial	Discreto (entero no negativo)
Bernoulli	Binario

Así que para poder usarlo es necesario que todas las características del vector x sean del mismo tipo.

- Scikit hace el smoothing adecuado de manera automática para evitar probabilidades iguales a 0 que impidan clasificar a un vector

28

- Para usarlo, simplemente se utiliza este código:

```
from sklearn.naive_bayes import GaussianNB
gnb = GaussianNB()
gnb.fit(X, y)
y_pred = gnb.predict(X)
```

29

Ventajas y desventajas

Ventajas	Desventajas
Es muy rápido de entrenar	Es necesario aplicar Smoothing para evitar probabilidades iguales a 0
Es fácil de implementar	Aunque sirve bien para clasificar, no es un buen estimador de probabilidades
Funciona bien con muchas dimensiones del vector de característica	Puede no generalizar bien a partir del training set
Como se asume independencia, el cálculo de la probabilidad de cada dimensión se puede hacer de manera independiente	Si las características están correlacionadas, puede tener resultados sesgados

30

Referencias

- <https://blog.sicara.com/naive-bayes-classifier-sklearn-python-example-tips-42d100429e44>
- <https://www.kdnuggets.com/2017/03/email-spam-filtering-an-implementation-with-python-and-scikit-learn.html>
- http://scikit-learn.org/stable/modules/naive_bayes.html