

TRATAMIENTO ESTADÍSTICO DE SEÑALES

VARIABLES Y PROCESOS ALEATORIOS - CONCEPTOS

Variable aleatoria X

• fdp. $p(x)$ 

valor medio: $\mu = E[x]$

pot. media: $P = E[x^2]$

varianza: $\sigma^2 = E[(x - E[x])^2]$

$$P = (\mu)^2 + (\sigma^2)$$

$$E[x] = \int x \cdot p(x) dx$$

$$E[f(x)] = \int f(x) \cdot p(x) dx$$

valor medio de $f(x)$

$$E[x|y] = \int x \cdot p(x|y) dx$$

$$E[f(x_1, x_2)] = \int f(x_1, x_2) \cdot p(x_1, x_2) dx_1 dx_2$$

• fdp 2D



correlación $E[x \cdot y^*]$

covarianza $E[(x - E[x]) \cdot (y - E[y])^*]$

indep $\Leftrightarrow p(x, y) = p(x) \cdot p(y)$
 $E[x \cdot y] = E[x] \cdot E[y]$

• V.A. Gaussiana

$x: N(\mu, \sigma^2)$ 

$$p(x) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

• Ruido blanco media nula

$\mu = 0$
 $R(m) = \begin{cases} P & m=0 \\ 0 & \text{resto} \end{cases}$ $S(\Omega) = P$

• V.A. Chi-cuadrado con N grados de lib

Elevar al cuadrado un vector de N gaussianas

$$x = Z \cdot Z^T$$

$$x = \sum_{i=1}^N Z_i^2$$

$$p(x|H_1) = \frac{x^{\frac{(N-1)}{2}}}{2^{\frac{N}{2}} \sigma^N \Gamma(\frac{N}{2})} \exp\left(-\frac{x}{2\sigma^2}\right) u(x)$$

si $N=2 \rightarrow p(x|H_1) = \frac{1}{2\sigma^2} \exp\left(-\frac{x}{2\sigma^2}\right) u(x)$

escalón

• Función de V.A.

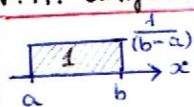
$y = f(x) \rightarrow p_x(x)$
 $x = g(y)$

$p_y(y) = p_x(x) \cdot \frac{dx}{dy}$

Fácil recordar

$dx \cdot p(x) = dy \cdot p(y)$

• V.A. Uniforme



$E(X) = \frac{b+a}{2}$

$\text{var}(X) = \frac{(b-a)^2}{12}$

1

• Función suma de 2 V.A.

$X = Y + Z$

$p_x(x) = p_y(y) * p_z(z)$

convolución

V.A. exponencial

$f_x(x) = a e^{-ax}$ $E[x] = \frac{1}{a}$

$\text{var}(x) = \frac{1}{a^2}$

Cálculos típicos

VECTOR de observaciones: $y = [y_0, y_1, \dots, y_{N-1}]^T$

$$\rightarrow Pr(y) = \prod_{i=0}^{N-1} Pr(y_i)$$

$$\rightarrow p(y|\alpha) = \prod_{i=0}^{N-1} p(y_i|\alpha)$$

ej. típico $y_i = \alpha + n_i \xrightarrow{N(0, \sigma_n^2)} N(\alpha, \sigma_n^2)$

$$p(y|\alpha) = \prod_{i=0}^{N-1} \frac{1}{\sqrt{2\pi\sigma_n^2}} e^{-\frac{(y_i - \alpha)^2}{2\sigma_n^2}} = \frac{1}{(2\pi\sigma_n^2)^{\frac{N}{2}}} \prod_{i=0}^{N-1} e^{-\frac{(y_i - \alpha)^2}{2\sigma_n^2}}$$

$$\ln p(y|\alpha) = K - \sum_{i=0}^{N-1} \frac{(y_i - \alpha)^2}{2\sigma_n^2}$$

VARIANZA (calcular autocorrelación)

ej. $y \rightarrow y_i = \alpha + n_i \xrightarrow{N(0, \sigma_n^2)}$

$$\rightarrow \hat{a}_{ML} = \frac{1}{N} \sum_{i=0}^{N-1} y_i$$

$$\text{var}(\hat{a}_{ML})$$

$$= E[\hat{a}_{ML}^2] - E^2[\hat{a}_{ML}]$$

$$= E\left[\frac{1}{N^2} \sum_i \sum_j y_i y_j\right] - \alpha^2$$

$$= \frac{1}{N^2} \sum_i \sum_j E[y_i y_j] - \alpha^2$$

$$\begin{cases} E[y_i^2] = \sigma_y^2 + \alpha^2 & \text{si } i=j \\ E[y_i y_j] = E[y_i] \cdot E[y_j] = \alpha^2 & \text{si } i \neq j \end{cases}$$

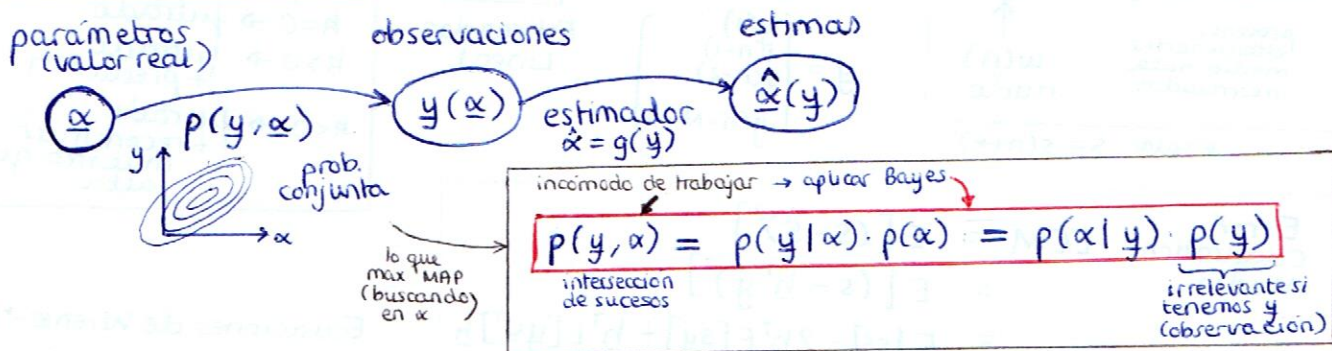
↑ indep

$$= \frac{1}{N^2} \left[\sum_i \frac{(\sigma_n^2 + \alpha^2)}{E[y_i^2]} + \sum_{i \neq j} \alpha^2 \right] - \alpha^2$$

$$= \frac{1}{N^2} [N \cdot (\sigma_n^2 + \alpha^2) + N(N-1)\alpha^2] - \alpha^2 = \frac{\sigma_n^2}{N}$$

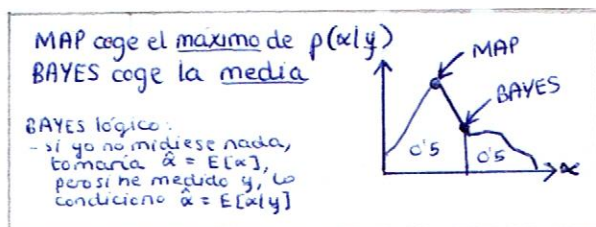
ESTIMACIÓN

Espacios involucrados



Estimadores óptimos

- Maximum a Posteriori **MAP** $\rightarrow \hat{\alpha}_{MAP} = \max_{\{\alpha\}} (p(\alpha|y))$ $\leftarrow$ útil para entenderlo
 $= \max (p(y|\alpha) \cdot p(\alpha))$ $\leftarrow$ útil para calcular
 - Máxima Verosimilitud **ML** $\rightarrow \hat{\alpha}_{ML} = \max (p(y|\alpha))$ $\leftarrow$ no tiene en cuenta la prob del parámetro (es como considerar α determinista, no conocemos su jdp)
 - Bayes **BAYES**
 Función error/coste $E(\alpha, \hat{\alpha}) = \begin{cases} |\alpha - \hat{\alpha}| \\ (\alpha - \hat{\alpha})^2 \\ \vdots \end{cases}$
 $\hat{\alpha}_{Bayes} = \min_{\hat{\alpha}} E(\alpha, \hat{\alpha}) \rightarrow$ minimiza el coste
 $= \min \iint E(\alpha, \hat{\alpha}) \cdot p(\alpha, y) d\alpha dy$ $\leftarrow$ irrelevante para y conocido
 $\hat{\alpha}_{Bayes} = \min_{\hat{\alpha}} \int E(\alpha, \hat{\alpha}) \cdot p(\alpha|y) d\alpha$ $\leftarrow$ Coste medio condicionado a y
- $\rightarrow$ para $E = (\alpha - \hat{\alpha})^2$: $\hat{\alpha}_{Bayes ECM} = E[\alpha|y]$



Demostración
 si $E(\alpha, \hat{\alpha}) = (\alpha - \hat{\alpha})^2$
 entonces $\min \int (\alpha - \hat{\alpha})^2 p(\alpha|y) d\alpha$
 $\frac{d}{d\hat{\alpha}} \left(\int (\alpha - \hat{\alpha})^2 p(\alpha|y) d\alpha \right) = 0$
 $= -2 \int (\alpha - \hat{\alpha}) p(\alpha|y) d\alpha$
 $\int \alpha p(\alpha|y) d\alpha = \int \hat{\alpha} p(\alpha|y) d\alpha$
 $E[\alpha|y] = \hat{\alpha} \cdot \int p(\alpha|y) d\alpha$
 $\int p(\alpha|y) d\alpha = 1$

Propiedades de los estimadores

- Sesgo($\hat{\alpha}$) = $E[\hat{\alpha}] - E[\alpha]$
- Varianza($\hat{\alpha}$) = $E[(\hat{\alpha} - E[\hat{\alpha}])^2]$
 $= E[\hat{\alpha}^2] - E[\hat{\alpha}]^2$

$$E[\epsilon] = \text{sesgo}^2(\hat{\alpha}) + \text{var}(\hat{\alpha})$$

error

- insesgado: $E[\hat{\alpha}] = E[\alpha]$

- consistente: $\lim_{N \rightarrow \infty} \text{var}(\hat{\alpha}) = 0$

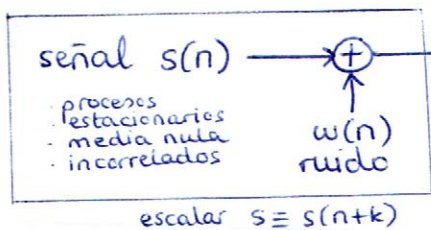
- mínima varianza: $\text{var} \hat{\alpha}_{MV} < \text{var} \hat{\alpha}$ para cualquier $\hat{\alpha}$

- eficiente: $\text{var} \hat{\alpha} = \text{CRB}$

- Cota de Cramer-Rao
 suponiendo α determinista
 $\hat{\alpha}$ insesgado
 Nunca podemos conseguir que $\text{var}(\hat{\alpha}) < \text{CRB}$

$$\text{CRB} = \frac{-1}{E_{\{y\}} \left[\frac{\partial^2}{\partial \alpha^2} \ln(p(y|\alpha)) \right]}$$

Estimación lineal de señales



$$y = \begin{bmatrix} y(n) \\ y(n-1) \\ y(n-2) \\ \vdots \\ y(n-N+1) \end{bmatrix}$$

FIR
 h
Estimador
Lineal

$$\hat{s}(n+k) \equiv \hat{s}$$

es un escalar
muestra
que predécimos

$k=0 \rightarrow$ filtrado
 $k>0 \rightarrow$ filtrado
+ predicción
 $k<0 \rightarrow$ filtrado
+ reconstruir
muestra que
falta

Error
Cuadrático
Medio

$$\begin{aligned} ECM &= E[(s - \hat{s})^2] \\ &= E[(s - h^T y)^2] \\ &= E[s^2] - 2h^T E[sy] + h^T E[yy^T] h \end{aligned}$$

$N \times 1$ $N \times N$
 R_{sy} R_{yy}

Minimizar ECM
con h

$$\frac{\partial ECM}{\partial h} = \begin{bmatrix} \partial ECM / \partial h_1 \\ \partial ECM / \partial h_2 \\ \vdots \end{bmatrix} = -2 R_{sy} + 2 R_{yy} h = 0$$

Ecuaciones de Wiener-Hopf

solución lineal óptima

$$h = R_{yy}^{-1} R_{sy}$$

$$\begin{aligned} R_{yy} &= E[yy^T] \\ R_{sy} &= E[s(n+k)y] \end{aligned}$$

Si todo es gaussiano
coincide con MAP, ML, ...

$$R_{yy} = \text{matriz de autocorrelación de } y = E[y \cdot y^T]$$

$N \times N$

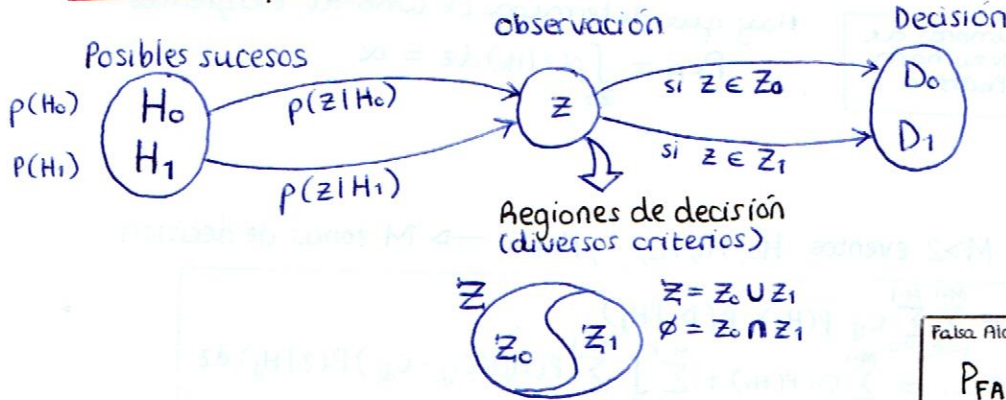
$$\begin{aligned} R_{sy} &= \text{vector} = E[s \cdot y] \\ R_{sy}(i) &= E[s(n+k)y(n-i+1)] \\ &= E[s(n+k)s(n-i+1)] \end{aligned}$$

ruido independiente con s

$$e.g. R_{sy}(0) = E[s(n+k) \cdot s(n)]$$

DETECCIÓN

Test de Hipótesis



Coste

• Func normal $H_0 \rightarrow D_0$	C_{00}
• Detección $H_1 \rightarrow D_1$	C_{11}
• Falsa Alarma $H_0 \rightarrow D_1$	C_{10}
• Miss $H_1 \rightarrow D_0$	C_{01}

Criterios de decisión

MAP

$$p(H_1|z) \stackrel{H_1}{\underset{H_0}{\gtrless}} p(H_0|z) \quad \leftarrow \text{entenderlo Bayes}$$

$$p(z|H_1) \cdot p(H_1) \stackrel{H_1}{\underset{H_0}{\gtrless}} p(z|H_0) \cdot p(H_0) \quad \leftarrow \text{aplicarlo}$$

$$\underbrace{\frac{p(z|H_1)}{p(z|H_0)}}_{\Lambda(z) \text{ razón de verosimilitud}} \stackrel{H_1}{\underset{H_0}{\gtrless}} \underbrace{\frac{p(H_0)}{p(H_1)}}_{\lambda \text{ umbral (no umbral en } Z \text{ sino en } \Lambda(z))}$$

ML

$$p(z|H_1) \stackrel{H_1}{\underset{H_0}{\gtrless}} p(z|H_0)$$

$$\frac{p(z|H_1)}{p(z|H_0)} \stackrel{H_1}{\underset{H_0}{\gtrless}} 1$$

Bayes

Escoge la región Z_0 que minimiza la parte variable del coste medio

$$p(H_1)(C_{01}-C_{11})p(z|H_1) \stackrel{H_1}{\underset{H_0}{\gtrless}} p(H_0)(C_{10}-C_{00})p(z|H_0)$$

MISS DETECT FA

$$\underbrace{\frac{p(z|H_1)}{p(z|H_0)}}_{\Lambda(z)} \stackrel{H_1}{\underset{H_0}{\gtrless}} \underbrace{\frac{C_{10}-C_{00}}{C_{01}-C_{11}} \frac{p(H_0)}{p(H_1)}}_{\lambda}$$

Minima prob. de error

supone $C_{00} = C_{11} = 0$ III $\bar{C} = P_e \Rightarrow$ minimizar coste (Bayes) III minimizar P_e (MAP)

$C_{01} = C_{10} = 1$

Minimax

No conocemos $P(H_1)$ así que tomamos el caso peor suponemos $P(H_1) = P_1^*$

Falsa Alarma

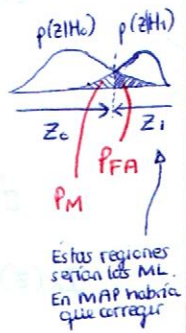
$$P_{FA} = \int_{Z_1} p(z|H_0) dz$$

$$P_M = \int_{Z_0} p(z|H_1) dz$$

no tienen en cuenta $p(H_0)$ y $p(H_1)$

Prob. error:

$$P_e = P_{FA} \cdot P(H_0) + P_M \cdot P(H_1)$$



COSTE MEDIO

$$\begin{aligned} \bar{C} &= C_{00} \cdot P(H_0, D_0) + C_{11} \cdot P(H_1, D_1) + C_{10} \cdot P(H_0, D_1) + C_{01} \cdot P(H_1, D_0) \\ &= C_{00} \cdot P(H_0) \cdot P(D_0|H_0) + C_{11} \cdot P(H_1) \cdot P(D_1|H_1) + C_{10} \cdot P(H_0) \cdot P(D_1|H_0) + C_{01} \cdot P(H_1) \cdot P(D_0|H_1) \\ &= C_{00} \cdot P(H_0) \cdot \int_{Z_0} p(z|H_0) dz + C_{11} \cdot P(H_1) \cdot [1 - \int_{Z_0} p(z|H_1) dz] + C_{10} \cdot P(H_0) \cdot [1 - \int_{Z_0} p(z|H_0) dz] + C_{01} \cdot P(H_1) \cdot \int_{Z_0} p(z|H_1) dz \end{aligned}$$

lógicamente $\int_{Z_1} p(z|H_1) dz = 1 - \int_{Z_0} p(z|H_1) dz$ lo hacemos para integrar siempre en Z_0

$$\bar{C} = C_{11} P(H_1) + C_{10} P(H_0) + \int_{Z_0} p(H_1)(C_{01}-C_{11})p(z|H_1) - p(H_0)(C_{10}-C_{00})p(z|H_0) dz$$

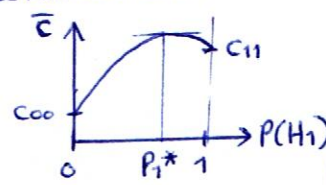
coste fijo variable (según Z_0 y Z_1)

$$\bar{C} = C_{11} P(H_1) + C_{10} P(H_0) + (C_{01}-C_{11})P(H_1) \underbrace{\int_{Z_0} p(z|H_1) dz}_{P_M} - (C_{10}-C_{00})P(H_0) \underbrace{\int_{Z_0} p(z|H_0) dz}_{1-P_{FA}}$$

Y utilizando que $P(H_0) = 1 - P(H_1)$ se llega a

$$\bar{C} = C_{00}(1-P_F) + C_{10} P_F + P(H_1)[(C_{11}-C_{00}) + (C_{01}-C_{11})P_M - (C_{10}-C_{00})P_F]$$

La relación con $P(H_1)$ es:



$$P_1^* \Leftrightarrow \frac{d\bar{C}}{dP(H_1)} = 0 \Leftrightarrow (C_{11}-C_{00}) + (C_{01}-C_{11})P_M - (C_{10}-C_{00})P_{FA} = 0$$

si $C_{00} = C_{11} = 0$
 $C_{10} = C_{01} = 1$

$P_M = P_{FA}$

Criterio de Neymann-Pearson

C_{ij} desconocidas
 $P(H_i)$ desconocidas

se trata de fijar $P_{FA} = \alpha$ (dato) y minimizar P_M (max P_{DET})

Para ello

$$\frac{p(z|H_1)}{p(z|H_0)} \underset{H_0}{\overset{H_1}{>}} \lambda \leftarrow \text{umbral de Neymann Pearson}$$

Hay que determinar el umbral exigiendo:

$$P_{FA} = \int_{Z_1} p(z|H_0) dz = \alpha$$

Hipótesis múltiple

$M > 2$ eventos $H_0, H_1, H_2, \dots, H_{M-1} \rightarrow M$ zonas de decisión

El coste medio:

$$\begin{aligned} \bar{C} &= \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} C_{ij} \cdot P(H_j) \cdot P(D_i|H_j) \\ &= \underbrace{\sum_{i=0}^{M-1} C_{ii} P(H_i)}_{\text{coste fijo}} + \underbrace{\sum_{i=0}^{M-1} \int_{Z_i} \sum_{j=0, j \neq i}^{M-1} P(H_j) (C_{ij} - C_{jj}) P(z|H_j) dz}_{\text{coste variable que minimizamos}} \end{aligned}$$

Escogemos las Z_i que MINIMICEN

$$\begin{aligned} I_0(z) &= \sum_{j \neq 0} P(H_j) (C_{0j} - C_{jj}) P(z|H_j) \rightarrow J_0(z) = \sum_{j \neq 0} P(H_j) (C_{0j} - C_{jj}) \cdot \frac{p(z|H_j)}{p(z|H_0)} \\ I_1(z) &= \sum_{j \neq 1} P(H_j) (C_{1j} - C_{jj}) P(z|H_j) \rightarrow J_1(z) = \sum_{j \neq 1} P(H_j) (C_{1j} - C_{jj}) \cdot \frac{p(z|H_j)}{p(z|H_0)} \\ &\vdots \\ I_i(z) &= \sum_{j \neq i} P(H_j) (C_{ij} - C_{jj}) P(z|H_j) \rightarrow J_i(z) = \sum_{j \neq i} P(H_j) (C_{ij} - C_{jj}) \cdot \frac{p(z|H_j)}{p(z|H_0)} \end{aligned}$$

normalizando por $p(z|H_0)$

$$\Lambda_j(z) \begin{cases} \Lambda_0(z) = 1 \\ \Lambda_1(z) = p(z|H_1)/p(z|H_0) \\ \vdots \\ \Lambda_{M-1}(z) = p(z|H_{M-1})/p(z|H_0) \end{cases}$$

se escoge

$$z \in Z_i \iff J_i(z) \leq J_j(z) \quad \forall j$$

Condición en el dominio $J(z) \rightarrow \Lambda(z)$
 Hay que convertirla en condición en el dominio z

CLASIFICACIÓN

Problema de clasificación



Preprocessing

→ sensing
separate pieces from one another and from background (segmentation)

Feature extraction

Reduce data by measuring certain features
e.g. lightness
length
width

Classifier

Threshold decision
e.g.



→ salmon
→ sea bass

Generalization

- control trade-off between training error and test error

e.g.



would give 0% training-error but bad test-error

→ OVERFITTING

• Features should be:

- incorrelated with each other
- do not add "noisy features"
- too many features is not good (problem of high dimensionality)

Poner demasiadas dimensiones puede hacer que se estime mal la distribución

- usar 'discriminative features'
- saber elegir features es un arte

Tipos

Statistical learning

$$x \rightarrow f(\cdot) \rightarrow y$$

- Supervised learning: - sample set of x and y during training
- $f(\cdot)$ is learned

- Unsupervised learning (clustering)

- only x is known
- try to cluster data into sets

- $Y = \{0, 1\} \rightarrow$ detection
- $Y = \{0, 1, 2, \dots, N\} \rightarrow$ classification
- $Y = \mathbb{R} \rightarrow$ regression

Bayesian Theory

$$P_{X,Y}(x, y) = P_{X|Y}(x|y) \cdot P_Y(y)$$

marginalization (eliminating variables)

$$P_{X_1, X_4}(x_1, x_4) = \iint P_{X_1, X_2, X_3, X_4}(x_1, x_2, x_3, x_4) dx_2 dx_3$$

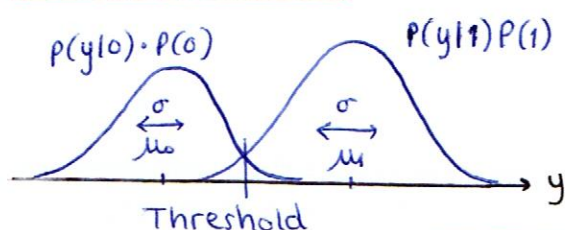
Bayes Rule:

$$P_{Y|X}(y|x) = \frac{P_{X|Y}(x|y) \cdot P_Y(y)}{P_X(x)} \rightarrow \text{no importante para clasificar}$$

MAR Rule classification

$$i^*(x) = \arg \max_i \frac{P(x|y_i) \cdot P(y_i)}{P(x)}$$

Gaussiana en 1D



$$y \geq \frac{\mu_1 + \mu_0}{2} + \frac{1}{\left(\frac{\mu_1 - \mu_0}{\sigma^2}\right)} \cdot \ln \frac{P(0)}{P(1)}$$

$\frac{1}{\left(\frac{\mu_1 - \mu_0}{\sigma^2}\right)}$ weighs the prior
distance between the means in units of variance

- classes very far apart:
"forget the prior"

- classes with same mean
"forget the observation, use prior"

The Gaussian Classifier

Gaussian
N-dim.

$$P_{X|Y}(\vec{x}|i) = \frac{1}{\sqrt{(2\pi)^d |\underline{\Sigma}_i|}} e^{-\frac{1}{2}(\vec{x}-\vec{\mu}_i)^T \underline{\Sigma}_i^{-1}(\vec{x}-\vec{\mu}_i)}$$

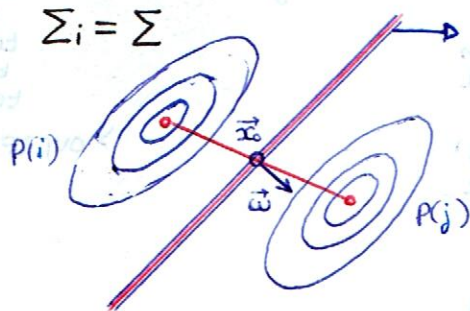
MAP $i^*(\vec{x}) = \arg \max [\ln P_{X|Y}(\vec{x}|i) + \ln P_Y(i)]$

General
Case
 $\vec{\mu}_1, \vec{\mu}_2$
 $\underline{\Sigma}_1, \underline{\Sigma}_2$

$$i^*(\vec{x}) = \arg \min \left[\underbrace{(\vec{x}-\vec{\mu}_i)^T \underline{\Sigma}_i^{-1}(\vec{x}-\vec{\mu}_i)}_{\text{Mahalanobis distance } d_i(\vec{x}, \vec{\mu}_i)} + \underbrace{\ln(2\pi)^d |\underline{\Sigma}_i| - 2 \ln P_Y(i)}_{\propto \text{class prior}} \right]$$


↳ special case: all classes have same covariance

$$\underline{\Sigma}_i = \underline{\Sigma}$$



Hyperplane $\vec{w}^T \cdot (\vec{x} - \vec{x}_0) = 0$

↳ passes through $\vec{x}_0$
↳ normal to $\vec{w}$

changes direction
direction joining the means

$$\vec{w} = \underline{\Sigma}^{-1} \cdot (\vec{\mu}_i - \vec{\mu}_j)$$

$$\vec{x}_0 = \underbrace{\frac{\vec{\mu}_i + \vec{\mu}_j}{2}}_{\text{midpoint}} - \underbrace{\frac{(\vec{\mu}_i - \vec{\mu}_j)}{(\vec{\mu}_i - \vec{\mu}_j)^T \underline{\Sigma}^{-1}(\vec{\mu}_i - \vec{\mu}_j)}}_{\text{weight to prior } \downarrow \text{ Mahalanobis distance between means}} \cdot \underbrace{\ln \frac{P_Y(i)}{P_Y(j)}}_{\text{prior}}$$

$$\underline{\Sigma} = \sigma^2 \underline{I}$$

Mahalanobis distance becomes EUCLIDEAN DISTANCE SQUARED in units of variance

$$d_i(\vec{x}_1, \vec{x}_2) = \frac{\|\vec{x}_1 - \vec{x}_2\|^2}{\sigma^2}$$

El peso al prior pasa a ser:

$$\frac{\vec{\mu}_i - \vec{\mu}_j}{d_i(\vec{\mu}_i, \vec{\mu}_j)} = \frac{\vec{\mu}_i - \vec{\mu}_j}{\left(\frac{\|\vec{\mu}_i - \vec{\mu}_j\|^2}{\sigma^2} \right)}$$

clara analogía al 1D

↳ sub-special case

same variance along all dimensions
(variables incorreladas)

$$\underline{\Sigma} = \sigma^2 \underline{I} = \begin{bmatrix} \sigma^2 & & \\ & \sigma^2 & \\ & & \sigma^2 \end{bmatrix}$$



$$\vec{w} = \frac{\vec{\mu}_i - \vec{\mu}_j}{\sigma^2} \quad \text{along the line joining the means}$$

$$\vec{x}_0 = \frac{\vec{\mu}_i + \vec{\mu}_j}{2} - \frac{(\vec{\mu}_i - \vec{\mu}_j)}{\left(\frac{\|\vec{\mu}_i - \vec{\mu}_j\|^2}{\sigma^2} \right)} \cdot \ln \frac{P_Y(i)}{P_Y(j)}$$

↳ sub-special case
 $P_Y(i) = P_Y(j)$
equiprobables

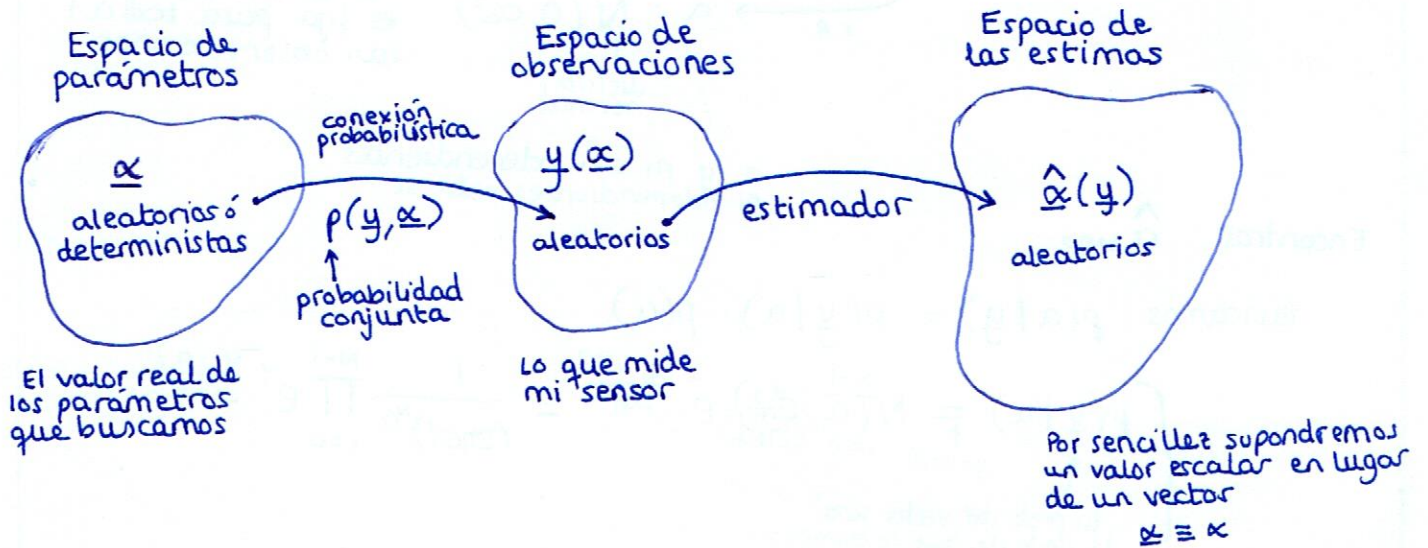
$$\vec{x}_0 = \frac{\vec{\mu}_i + \vec{\mu}_j}{2} \quad \text{midpoint of the means}$$



$$\ln \frac{P_Y(i)}{P_Y(j)} = \ln \frac{1}{2} + \ln \frac{1}{2} = 0$$

ESTIMACIÓN

Espacios involucrados y modelos de probabilidad



Estimadores óptimos

1. Estimador "Maximum a posteriori" (MAP)

Supondremos α aleatorio con $p(\alpha)$ conocida (si no conociésemos $p(\alpha)$ podríamos decir que es determinista desconocido)

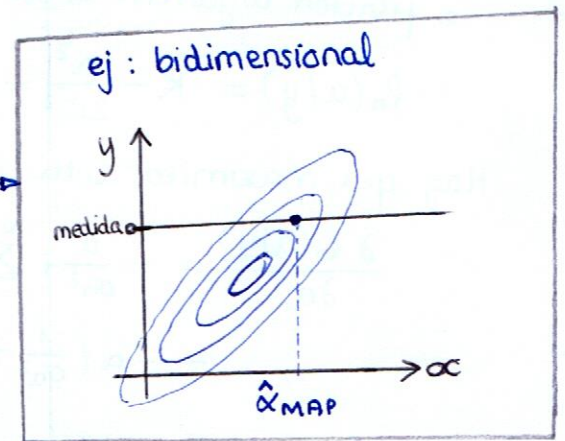
$$\hat{\alpha}_{\text{MAP}} = \hat{\alpha}(y) = \max_{\{\alpha\}} P(y, \alpha)$$

Annotations for the equation:

- Under $\{ \alpha \}$: "he medido en y "
- Under $P(y, \alpha)$: "me muevo en α para buscar el max."

Es engorroso trabajar con $P(y, \alpha)$
Pero aplicando Teorema de Bayes

$$P(y, \alpha) = \frac{P(y|\alpha) \cdot P(\alpha)}{P(y)} \leftarrow \text{no depende de } \alpha$$



$$\hat{\alpha}_{\text{MAP}} = \max_{\{\alpha\}} \underbrace{P(y|\alpha) \cdot P(\alpha)}_{= P(\alpha|y)}$$

← en la práctica ésta es la función fácil de conseguir

← es equivalente a maximizar $P(\alpha|y)$ lo cual le da nombre al método

A la función que maximizamos ($p(y, \alpha) \rightarrow p(y|\alpha) \cdot p(\alpha) \rightarrow p(\alpha|y)$) se le llama función de verosimilitud (likelihood)

$$l(y/\alpha)$$

$$\ln[l(y/\alpha)] = L(y/\alpha) \quad \text{función logarítmica de verosimilitud}$$

ejemplo 1:

suponemos $y = [y_0, y_1, \dots, y_{N-1}]^T$

donde $y_i = a + n_i$

$n_i \sim N(0, \sigma_n^2)$

$a \sim N(0, \sigma_a^2)$

distrib. normal

'a' es una V.A. pero es fija para todas las observaciones

a y n_i son independientes
ni independientes entre si

Encontrar $\hat{a}_{MAP}$

Buscamos $p(a|y) = p(y|a) \cdot p(a)$

$$p(y|a) = \prod_{i=0}^{N-1} \frac{1}{\sqrt{2\pi\sigma_n^2}} e^{-\frac{(y_i-a)^2}{2\sigma_n^2}} = \frac{1}{(2\pi\sigma_n^2)^{N/2}} \prod_{i=0}^{N-1} e^{-\frac{(y_i-a)^2}{2\sigma_n^2}}$$

la prob. del vector será la prob. de que se cumplan todas sus elementos (que como son independientes es el producto de gaussianas)

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

La función logarítmica de verosimilitud será

$$\ln(a|y) = K - \frac{a^2}{2\sigma_a^2} - \sum_{i=0}^{N-1} \frac{(y_i-a)^2}{2\sigma_n^2}$$

Hay que maximizar esta función con respecto a 'a'

$$\frac{\partial \ln(a|y)}{\partial a} = -\frac{a}{\sigma_n^2} + \sum_{i=0}^{N-1} \frac{(y_i-a)}{\sigma_n^2}$$

$$= -a \left(\frac{1}{\sigma_a^2} + \frac{N}{\sigma_n^2} \right) + \sum_{i=0}^{N-1} \frac{y_i}{\sigma_n^2} = 0$$

despejamos $\hat{a}_{MAP}$ que cumple esta ecuación

$$\Rightarrow \hat{a}_{MAP} = \frac{1}{\frac{\sigma_n^2}{\sigma_a^2} + N} \sum_{i=0}^{N-1} y_i$$

Esta fórmula es muy lógica sobretodo si la vemos en casos extremos

$$\text{si } \frac{\sigma_n^2}{\sigma_a^2} \ll N \rightarrow \hat{a}_{MAP} \approx \frac{1}{N} \sum_{i=0}^{N-1} y_i$$

la media de las observaciones

$$\text{si } \frac{\sigma_n^2}{\sigma_a^2} \gg N \rightarrow \hat{a}_{MAP} \approx \frac{\sigma_a^2}{\sigma_n^2} \sum_{i=0}^{N-1} y_i$$

esta ya no es tan intuitivo!

sabiendo que $\sigma_a \rightarrow 0$ el valor estimado lo reducimos acercándolo a 0

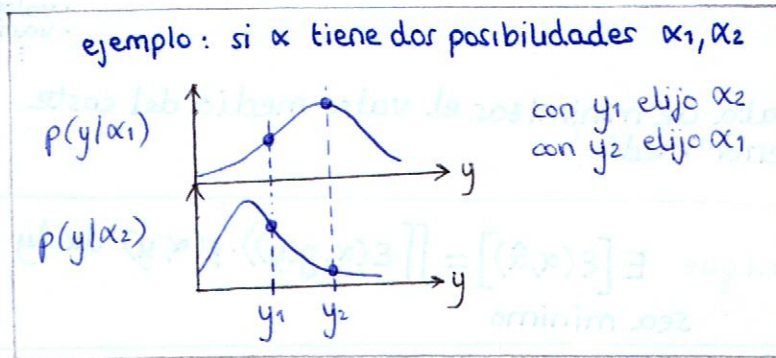
2. Estimador de máxima verosimilitud (ML)

Suponemos α determinista

$$\hat{\alpha}_{ML} = \max_{\{\alpha\}} p(y|\alpha)$$

Equivale a MAP
en el caso de
 $p(\alpha) = K$ uniforme

Es decir, no introducimos información
a priori acerca del parámetro



ejemplo 2: igual que 1 pero ' a ' determinista

$$\ln\{p(y|a)\} = K - \sum_{i=0}^{N-1} \frac{(y_i - a)^2}{2\sigma_n^2} \quad \text{función logarítmica de verosimilitud}$$

maximizando respecto a ' a '

$$\frac{\partial \ln p(y|a)}{\partial a} = - \sum_{i=0}^{N-1} \frac{(y_i - a)}{\sigma_n^2} = 0 \rightarrow$$

$$\hat{a}_{ML} = \frac{1}{N} \sum_{i=0}^{N-1} y_i$$

coincide con el MAP cuando no
hay info de ' a ' a priori
 $\sigma_a^2 \rightarrow \infty$

En resumen:

$$\hat{\alpha}_{ML} = \max_{\alpha} p(y|\alpha) \quad \text{escoge el parámetro que da la máxima probabilidad de lo que he medido}$$

$$\hat{\alpha}_{MAP} = \max_{\alpha} p(\alpha|y) = \max_{\alpha} p(y|\alpha) \cdot p(\alpha) \quad \text{escoge el parámetro que con máxima probabilidad ha generado lo que he medido}$$

← no estoy seguro

3. Estimador de Bayes

Cada pareja $\alpha, \hat{\alpha}$ tiene un coste (típicamente grande si he errado mucho)
a través de una función de coste

$$C(\alpha, \hat{\alpha})$$

típicamente

$$C(\alpha, \hat{\alpha}) = \begin{cases} (\alpha - \hat{\alpha})^2 \\ |\alpha - \hat{\alpha}| \\ \vdots \end{cases} = \mathcal{E}(\alpha, \hat{\alpha})$$

↓
será una
V.A.
· valor medio
· varianza

El estimador de Bayes trata de minimizar el valor medio del coste
(i.e. minimiza el 'error' medio)

$$\hat{\alpha}_{\text{BAYES}} = g(y) \text{ tal que } E[\mathcal{E}(\alpha, \hat{\alpha})] = \iint \mathcal{E}(\alpha, g(y)) \cdot p(\alpha, y) \, d\alpha \, dy$$

sea mínimo

Recuerda

$$E\{X\} = \int x \cdot p(x) \, dx$$

$$E\{f(x)\} = \int f(x) \cdot p(x) \, dx$$

$$E\{f(x_1, x_2)\} = \int f(x_1, x_2) \cdot p(x_1, x_2) \, dx_1 \, dx_2$$

se puede expresar la integral:

$$E[\mathcal{E}(\alpha, \hat{\alpha})] = \iint \mathcal{E}(\alpha, g(y)) \cdot \underbrace{p(\alpha, y)}_{p(\alpha|y) \cdot p(y)} \, d\alpha \, dy$$

$$= \int p(y) \left[\int \mathcal{E}(\alpha, g(y)) p(\alpha|y) \, d\alpha \right] dy$$

↑
siempre
es positivo

esto dará lugar a
un valor que depende de y

↓
puedo limitarme a minimizar esto
(llamado:

Coste medio condicionado a y)

si lo minimizo para todo y estaré minimizando
todo el coste medio

particular:
si escogemos $E(\alpha, g(y)) = (\alpha - g(y))^2 = (\alpha - \hat{\alpha})^2$
i.e. error cuadrático como función de coste

- $$\frac{\partial}{\partial g} \int (\alpha - g(y))^2 p(\alpha | y) d\alpha = -2 \int (\alpha - g(y)) \cdot p(\alpha | y) d\alpha = 0$$

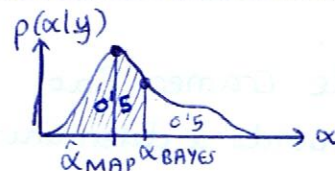
La integral de cualquier densidad de probabilidad (aunque sea condicionada, etc...) debe ser 1

- En general es un procesado no lineal en y
- si todo es gaussiano, sí es un operador lineal
(es decir $x = A \cdot y$)
 $\quad\quad\quad \uparrow$
 $\quad\quad\quad$ matriz

si yo no midiere nada, tomaría $\hat{\alpha} = E[\alpha]$

pero si he medido y , tomo el valor medio **CONDICIONADO** a y

Es decir:
MAP coge el MÁXIMO de $p(x|y)$
BAYES coge la MEDIA de $x|y$



minimiz. error cuadrático medio

$$\hat{a}_{\text{BAYES}} \stackrel{\downarrow}{=} E[a|y] = \int_{-\infty}^{\infty} a \cdot p(a|y) da$$

recuerda:

$$p(a|y) = \frac{1}{(2\pi\sigma_n^2)^{N/2}} \left[\prod_{i=0}^{N-1} e^{-\frac{(y_i - a)^2}{2\sigma_n^2}} \right] \frac{1}{\sqrt{2\pi\sigma_a^2}} e^{-\frac{a^2}{2\sigma_a^2}}$$

Propiedades de los estimadores

- Insesgado: $E[\hat{\alpha}] = E[\alpha]$

i.e. si hago ∞ estimaciones, la media será el valor correcto

$$\text{sesgo}(\alpha) = E[\hat{\alpha}] - E[\alpha]$$

- Consistente:

$\lim_{N \rightarrow \infty} \Pr(\hat{\alpha} - \alpha < \epsilon) = 0$ i.e. puedo conseguir que el valor estimado sea tan cercano al real como se desee sin más que repetir la estimación N veces

- Mínima varianza: $\text{var} \hat{\alpha}_{MV} \leq \text{var} \hat{\alpha}$, para cualquier $\hat{\alpha}$

recuerda, se cumple:

$$\text{var} \hat{\alpha} = E[\hat{\alpha}^2] - E^2[\hat{\alpha}]$$

$$E[E^2] = \text{var} \hat{\alpha} + \text{sesgo}^2(\hat{\alpha})$$

↑
error

Normalmente preocupa más la varianza que el sesgo (a veces el sesgo se puede compensar)

- Cota de Cramer-Rao

suponiendo α determinista, $\hat{\alpha}$ insesgado

$$\text{var}(\hat{\alpha}) \geq \text{CRB}$$

nunca podemos conseguir que la varianza baje por debajo de la cota

La distancia a la cota indica lo mejorable que es el estimador

$$\text{CRB} = \frac{1}{E \left[\frac{\partial^2}{\partial \alpha^2} \ln(p(y|\alpha)) \right]}$$

- Estimador eficiente

$$\text{var}(\hat{\alpha}) = \text{CRB}$$

ejemplo 4: ML

Hipótesis de ejemplos 1, 2, 3

$$\hat{a}_{ML} = \frac{1}{N} \sum_{i=0}^{N-1} y_i$$

$$E[\hat{a}_{ML}] = \frac{1}{N} \sum_{i=0}^{N-1} E[y_i] = a \rightarrow \text{estimador insesgado}$$

$$\text{var}[\hat{a}_{ML}] = E[\hat{a}_{ML}^2] - E^2[\hat{a}_{ML}]$$

$$= E\left[\frac{1}{N^2} \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} y_i y_j\right] - a^2$$

$$= \frac{1}{N^2} \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} E[y_i y_j] - a^2$$

$$= \begin{cases} E[y_i^2] = \sigma_n^2 + a^2 & \text{si } i=j \\ E[y_i y_j] = E[y_i] \cdot E[y_j] = a^2 & \text{si } i \neq j \end{cases}$$

↑ independientes

!

$$= \frac{1}{N^2} \left[\sum_{i=0}^{N-1} E[y_i^2] + \sum_{i=0}^{N-1} \sum_{j \neq i}^{N-1} E[y_i] \cdot E[y_j] \right] - a^2$$

$$= \frac{1}{N^2} \left[N \cdot (\sigma_n^2 + a^2) + (N \times N - 1) \cdot a^2 \right] - a^2$$

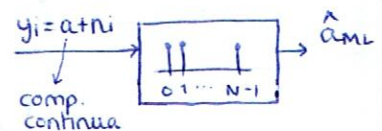
$$\text{var} \hat{a}_{ML} = \frac{\sigma_n^2}{N}$$

Estamos reduciendo el ruido

$$\begin{aligned} y_i &= a + n_i \rightarrow \sigma_n^2 \\ \hat{a}_{ML} &= a + n_i' \rightarrow \frac{\sigma_n^2}{N} \end{aligned}$$

$$\text{var} \hat{a}_{ML} \xrightarrow{N \rightarrow \infty} 0 \rightarrow \text{estimador consistente}$$

si te fijas bien, $\hat{a}_{ML}$ en este caso no es mas que un filtro FIR que elimina ruido



Calculemos la Cota de Kramer-Rao

recuerda que ya teniamos

$$\frac{\partial \ln p(y|a)}{\partial a} = \frac{\partial}{\partial a} \left\{ k - \sum_{i=0}^{N-1} \frac{(y_i - a)^2}{2\sigma_n^2} \right\}$$

$$= \frac{1}{2\sigma_n^2} (-2) \sum_{i=0}^{N-1} (y_i - a)(-1) = \frac{1}{\sigma_n^2} \sum_{i=0}^{N-1} (y_i - a)$$

$$\frac{\partial^2 \ln(p(y|a))}{\partial^2 a} = \frac{1}{\sigma_n^2} \sum_{i=0}^{N-1} (-1) = \frac{-N}{\sigma_n^2}$$

$$\text{CRB} = - \frac{1}{E\left[\frac{\partial^2 \ln p(y|a)}{\partial^2 a}\right]} = \frac{\sigma_n^2}{N}$$

$$\text{var} \hat{a}_{ML} = \text{CRB} \rightarrow \text{estimador eficiente}$$

Si el ruido no fuese gaussiano, habria otras formas mejores de estimar el valor medio

Ejemplo 5.

Demstrar que cualquier estimador de Bayes es insesgado
(lo demostramos para el caso del error cuadrático como función de coste)

$$E[\hat{\alpha}_{\text{BAYES}}] = E[E[\alpha|y]]$$

valor medio de la estima para varias observaciones (i.e. la aleatoriedad es en y)

valor medio de α en la función $p(\alpha|y)$ (i.e. aleatoriedad en α)

$$= E \left[\int \alpha \cdot p(\alpha|y) d\alpha \right] = \int \alpha \cdot E[p(\alpha|y)] \cdot d\alpha$$

$$= \int \alpha \cdot \left[\int p(\alpha|y) \cdot p(y) dy \right] d\alpha$$

$p(y|\alpha) \cdot p(\alpha)$ Bayes

$$= \int \alpha \cdot \left[\int p(y|\alpha) \cdot p(\alpha) dy \right] d\alpha$$

$$= \int \alpha \cdot p(\alpha) \left[\int p(y|\alpha) dy \right] d\alpha$$

1

$$= \int \alpha \cdot p(\alpha) d\alpha = E[\alpha]$$

recuerda:

valor medio de x

$$E[x] = \int x p(x) dx$$

valor medio de $f(x)$

$$E[f(x)] = \int f(x) p(x) dx$$

Estimación lineal de señales

Vamos a trabajar con → Procesos estocásticos (¡señales!)
→ Estimadores lineales

$$y(n) = \underset{\substack{\downarrow \\ \text{señal}}}{s(n)} + \underset{\substack{\downarrow \\ \text{ruido}}}{w(n)}$$

- $\{\tilde{s}(n)\}$ proceso estacionario en sentido amplio
 $E[\tilde{s}(n+m) \cdot s(n)] = R_{ss}(m)$
 y con media nula $E[\tilde{s}(n)] = 0$
 (si no lo fuera, se le resta y ya está)

- $\{\tilde{w}(n)\}$ proceso estacionario en sentido amplio
 $E[\tilde{w}(n+m) \cdot w(n)] = R_{ww}(m)$
 y con media nula $E[\tilde{w}(n)] = 0$

- la señal y el ruido son incorrelados
 $E[\tilde{w}(n+m) \cdot \tilde{s}(n)] = R_{ws}(m) = 0$

Problema: Estimar $s(n+k) \equiv s$
a partir de $y = [y(n), y(n-1), \dots, y(n-N+1)]$

si $k=0 \rightarrow$ filtrado $\rightarrow$ eliminar el ruido

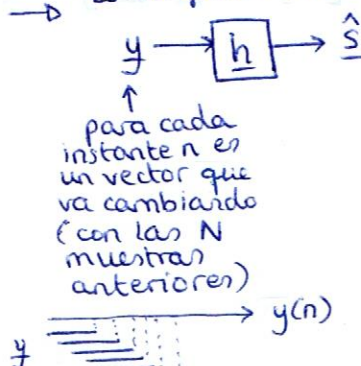
$k>0 \rightarrow$ filtrado + predicción

$k<0 \rightarrow$ filtrado + interpolación $\rightarrow$ intento reconstruir muestras que me faltan en el pasado

y quiero que sea un procesado LINEAL

$$\hat{s}(n+k) = \underline{h}^T \underline{y} \rightarrow \text{Es un filtro FIR}$$

↑
hay que encontrar esta $\underline{h}$



Error cuadrático
medio

$$\begin{aligned} ECM &= E[(s - \underline{h}^T \underline{y})^2] \\ &= E[s^2 - 2s\underline{h}^T \underline{y} + \underline{h}^T \underline{y} \underline{y}^T \underline{h}] \end{aligned}$$

NOTA:

$\underline{h}^T \underline{y}$ no es más que un número

$\underline{h}^T \underline{y} \underline{y}^T \underline{h}$ no es más que ese número al cuadrado. He cambiado los traspuestos porque me conviene y el producto escalar es conmutativo

$$= E[s^2] - 2\underline{h}^T E[s\underline{y}] + \underline{h}^T E[\underline{y} \underline{y}^T] \underline{h}$$

un vector
de valores
medios
N x 1

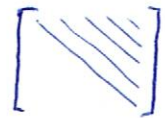
matriz
N x N

indaguemos en ambos

$$E[\underline{y} \underline{y}^T] = \underline{R}_{yy} \longrightarrow \underline{R}_{yy}(i, j) = E[y(i) \cdot y(j)] = E[y(n-i+1) y(n-j+1)] = R_{yy}(j-i)$$

la comp.
(i,j) de la
matriz

la matriz de autocorrelación
puede construirse conociendo la
función de autocorrelación
del proceso estocástico $y(n)$



$$E[s \cdot \underline{y}] = \underline{R}_{sy} \longrightarrow R_{sy}(i) = E[s(n+k) y(n-i+1)]$$

podemos quitar la parte
del ruido de la y , ya
que es incorrelado
con s

$$\begin{aligned} &= E[s(n+k) \cdot s(n-i+1)] \\ &= R_{ss}(k+i-1) \end{aligned}$$

a la correlación
sólo le importa
la diferencia
de índices

$$ECM = E[s^2] - 2\underline{h}^T \underline{R}_{sy} + \underline{h}^T \underline{R}_{yy} \underline{h}$$

Nos interesa minimizar el ECM, por tanto derivamos

(con Reglas de derivación con matrices y vectores)
(se resuelve muy fácilmente)

$$\frac{\partial \text{ECM}}{\partial \underline{h}} = \begin{bmatrix} \frac{\partial \text{ECM}}{\partial h(n)} \\ \vdots \\ \frac{\partial \text{ECM}}{\partial h(n)} \end{bmatrix} = -2 \underline{R}_{sy} + 2 \underline{R}_{yy} \underline{h} = 0$$

Iguando $\frac{\partial \text{ECM}}{\partial \underline{h}}$ a cero se obtiene un sistema de ecuaciones lineales que se resuelve con:

$$\underline{h} = \underline{R}_{yy}^{-1} \underline{R}_{sy} \quad \text{Ecuaciones de Wiener-Hopf}$$

Esta es la solución lineal óptima

Además si todo es gaussiano resulta que coincide esta solución con las soluciones MAP, ML, ... $p(x/y)$

Ejemplo: Predecir el valor de un proceso estacionario (sin ruido)
(por ej. predecir la bolsa)

Ejemplo 6: construir el predictor lineal óptimo (ECM mínimo) de $s(n+1)$ a partir de $s(n)$

$$w(n) = 0 \rightarrow R_{ww}(m) = 0 \quad s \equiv s(n+1) \quad \uparrow k=1$$

$$\underline{y} = [s(n)]^T = s(n) \quad \uparrow N=1$$

La matriz de autocorrelación

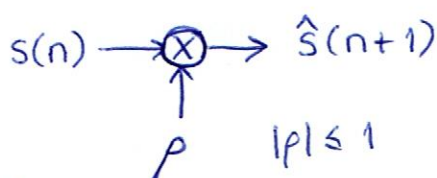
$$\underline{R}_{yy} = \underline{R}_{ss} = R_{ss}(0) \quad \left. \begin{array}{l} \text{matriz} \\ \text{tamaño } 1 \end{array} \right\} \text{escalar}$$

vector de correlación señal

$$\underline{R}_{sy} = E[s(n+1)s(n)] = R_{ss}(1)$$

$$\underline{h} = h = \frac{R_{ss}(1)}{R_{ss}(0)} = \rho$$

es decir



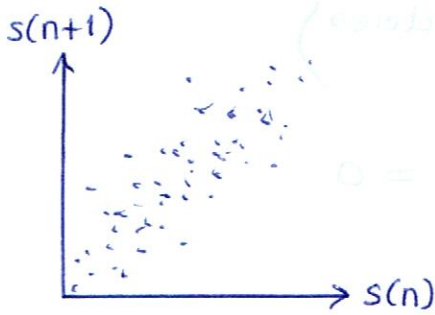
si $\rho = 1 \rightarrow$ mucha correlación entre n y $n+1$
(tomo el valor anterior)

si $\rho = 0 \rightarrow$ no tengo ni idea así que digo
o que es la media

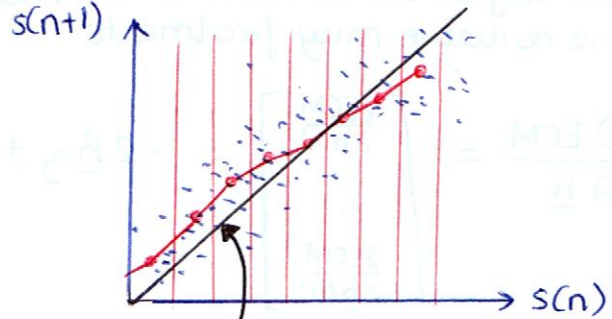
nota de campo:
algunos que trabajan
en esto dicen que una
matriz es mínimo
1000x1000. Cualquier
matriz menor que
eso lo consideran
un escalar... 😊

si emparejo muestras contiguas.

Puedo calcular la media alrededor de ventanas de $s(n)$



Forma sencilla experimental de hacer un predictor con $N=1$



será la media de los puntos si el proceso es gaussiano

$$\hat{s}(n) = \frac{1}{N} \sum_{k=1}^N s(n-k)$$

Problema 1

Tengo un conjunto de números binarios

$$x_0, x_1, \dots, x_{N-1}$$

independientes

MUY IMPORTANTE

$$\text{Prob} \{ x_i = 1 \} = \theta \quad \forall i$$

$$\text{Hallar } \hat{\theta}_{ML} \rightarrow \hat{\theta}_{ML} = \max_{\theta} p(x|\theta)$$

calculemos:

$$p(x|\theta) = \prod_{i=0}^{N-1} p(x_i|\theta)$$

$$p(x_i|\theta) = \begin{cases} 1-\theta & x_i=0 \\ \theta & x_i=1 \end{cases} = \theta^{x_i} (1-\theta)^{1-x_i}$$

compactando (i)

Por tanto:

$$p(x|\theta) = \prod_{i=0}^{N-1} \theta^{x_i} (1-\theta)^{1-x_i}$$

$$= \theta^k (1-\theta)^{N-k}$$

siendo k el número de unos $k = \sum_{i=0}^{N-1} x_i$

ln (i) muy aconsejable aplicar el ln para simplificar

$$\ln p(x|\theta) = k \ln \theta + (N-k) \ln(1-\theta)$$

Maximizándolo:

$$\begin{aligned} \frac{\partial \ln p(x|\theta)}{\partial \theta} &= \frac{k}{\theta} - (N-k) \frac{1}{1-\theta} \\ &= \frac{k(1-\theta) - (N-k)\theta}{\theta(1-\theta)} = 0 \end{aligned}$$

$$\Rightarrow k(1-\theta) - (N-k)\theta = 0$$

$$\Rightarrow \boxed{\hat{\theta}_{ML} = \frac{k}{N}}$$

Esto es cierto y sencillo gracias a que los x_i son incorrelados
(ej: intención de voto: PP o PSOE?
... si siempre preguntas al mismo.....)

Problema 3

$x_0, \dots, x_{M-1}$

independientes

$$p(x_i/\theta) = \theta e^{-\theta x_i} \quad E[x_i] = \frac{1}{\theta}$$

$(x_i \geq 0, \theta \geq 0)$

- calcular $\hat{\theta}_{ML}$
- calcular $\hat{\theta}_{MAP}$ suponiendo $p(\theta) = \lambda e^{-\lambda \theta} \quad \begin{matrix} \theta \geq 0 \\ \lambda \geq 0 \end{matrix}$

• ¿ $\hat{\theta}_{ML}$? $\hat{\theta}_{ML} = \max_{\theta} p(\underline{x}|\theta)$

$$p(\underline{x}|\theta) = \prod_{i=0}^{M-1} p(x_i/\theta) = \prod_{i=0}^{M-1} \theta e^{-\theta x_i} = \theta^M e^{-k\theta} \quad \text{siendo } k = \sum_i x_i$$

Hay que buscar el máximo.
Puede ser más sencillo tomando la función logarítmica

$$\ln p(\underline{x}|\theta) = M \ln \theta - k\theta$$

$$\frac{\partial \ln p(\underline{x}|\theta)}{\partial \theta} = \frac{M}{\theta} - k = 0$$

$$\Rightarrow \hat{\theta}_{ML} = \frac{M}{k} = \frac{1}{\frac{1}{M} \sum_{i=0}^{M-1} x_i}$$

la inversa de la media

lo cual es lógico ya que teníamos $E[x_i] = \frac{1}{\theta}$

• ¿ $\hat{\theta}_{MAP}$? $\hat{\theta}_{MAP} = \max_{\theta} p(\theta|\underline{x})$

$$= \max_{\theta} p(\underline{x}|\theta) \cdot p(\theta)$$

$$= \max_{\theta} \left[\prod_{i=0}^{M-1} p(x_i/\theta) \right] \cdot (\lambda e^{-\lambda \theta}) = \max_{\theta} [\theta^M e^{-k\theta}] \cdot (\lambda e^{-\lambda \theta})$$

tomando logaritmo neperiano para maximizar más fácil

$$\ln [p(\underline{x}|\theta) \cdot p(\theta)] = M \ln \theta - k\theta + \ln \lambda - \lambda \theta$$

$$\frac{\partial \ln [p(\underline{x}|\theta) \cdot p(\theta)]}{\partial \theta} = \frac{M}{\theta} - k - \lambda$$

$$\Rightarrow \hat{\theta}_{MAP} = \frac{M}{k + \lambda}$$

se ha corregido el valor anterior a partir del conocimiento del parámetro
si $\lambda \rightarrow 0$ se tiende a una varianza muy grande en la V.A. θ
 $p(\theta) = \lambda e^{-\lambda \theta} \rightarrow E[\theta] = \frac{1}{\lambda} \rightarrow \text{var}[\theta] = \frac{1}{\lambda^2}$
una varianza muy grande no aporta información

Problema 4

$x_0, \dots, x_{M-1}$ independientes

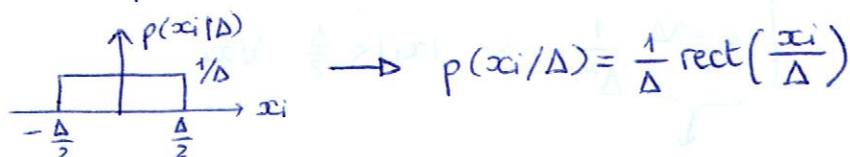
x_i uniforme entre $-\frac{\Delta}{2}, \frac{\Delta}{2}$

• Hallar $\hat{\Delta}_{ML}$

• Hallar $\hat{\Delta}_{MAP}$ sabiendo $p(\Delta) = \lambda e^{-\lambda \Delta}$

• ¿ $\hat{\Delta}_{ML}$?

Vamos a poner la distribución de cada x_i



$$\hat{\Delta}_{ML} = \max_{\Delta} p(\underline{x} | \Delta)$$

$$= \max_{\Delta} \prod_{i=0}^{M-1} p(x_i | \Delta)$$

$$= \max_{\Delta} \frac{1}{\Delta^M} \prod \text{rect}\left(\frac{x_i}{\Delta}\right)$$

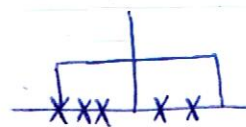
En lugar de derivar vamos a ver si encontramos el máximo razonando

$$p(\underline{x} | \Delta) = \begin{cases} 0 & \text{si } |x_i| > \frac{\Delta}{2} \text{ en algún } x_i \text{ (basta con uno!)} \\ \frac{1}{\Delta^M} & \text{si } |x_i| < \frac{\Delta}{2} \forall x_i \leftarrow \text{i.e. todo el vector} \end{cases}$$

me interesaría Δ muy pequeño para tener valor muy alto pero no lo suficiente como para anular algún x_i !!!

Cogeremos el mínimo posible siempre y cuando $\Delta/2$ esté por encima de todos los x_i

$$\Rightarrow \hat{\Delta}_{ML} = \max_i 2|x_i|$$



Nota: y si no supiéramos que la f.d.p. es simétrica, cogeríamos entre el mínimo y el máximo

lógicamente coger un Δ menor sería estúpido (tratamos de estimar la f.d.p. de x_i viendo sus muestras)

$$\begin{aligned} \hat{\Delta}_{\text{MAP}} &= \max_{\Delta} p(\underline{x}|\Delta) p(\Delta) \\ &= \max_{\Delta} \left[\frac{1}{\Delta^M} \prod_{i=0}^{M-1} \text{rect}\left(\frac{x_i}{\Delta}\right) \right] \cdot \lambda e^{-\lambda \Delta} \end{aligned}$$

$$= \max_{\Delta} \begin{cases} 0 & \text{si } |x_i| \geq \frac{\Delta}{2} \text{ en algún } x_i \\ \lambda e^{-\lambda \Delta} \cdot \frac{1}{\Delta^M} & \text{si } |x_i| < \frac{\Delta}{2} \forall x_i \end{cases}$$

este conocimiento no nos ha aportado nada ya que nos sigue pidiendo que busquemos el valor más pequeño posible de Δ (siempre y cuando entren dentro todas las observaciones)

$$\Rightarrow \hat{\Delta}_{\text{MAP}} = \max_i 2|x_i|$$

Problema 2

$x_0, \dots, x_{M-1}$ independientes

$$p(x_i | \theta) = e^{-\theta} \cdot \frac{1}{x_i!} \theta^{x_i}$$

Distribución de Poisson

$$x_i = 0, 1, 2, 3, \dots$$

- Hallar $\hat{\theta}_{ML}$
- Indicar si $\hat{\theta}_{ML}$ es eficiente

$$p(\underline{x} | \theta) = \prod_{i=0}^{M-1} p(x_i | \theta) = \frac{e^{-M\theta}}{\prod_{i=0}^{M-1} x_i!} \theta^k \quad k = \sum_{i=0}^{M-1} x_i$$

$$\ln p(\underline{x} | \theta) = -M\theta - \ln \left(\prod_{i=0}^{M-1} x_i! \right) + k \ln \theta$$

Hay que hallar el máximo

$$\frac{\partial \ln p(\underline{x} | \theta)}{\partial \theta} = -M + \frac{k}{\theta}$$

$$\Rightarrow \boxed{\hat{\theta}_{ML} = \frac{k}{M} = \frac{1}{M} \sum x_i}$$

la estimación de θ es la media muestral.

Tiene sentido ya que el valor medio de una distrib de Poisson

$$p(x_i) = e^{-\theta} \cdot \frac{1}{x_i!} \theta^{x_i} \rightarrow E[x_i] = \theta$$

- Para hallar si es eficiente, hay que hallar la cota de Kramer-Raus, la cual solo existe si el estimador es insesgado

$$\bullet \text{ ¿es insesgado? } E[\hat{\theta}_{ML}] = E\left[\frac{1}{M} \sum_{i=0}^{M-1} x_i\right] = \frac{1}{M} \sum_{i=0}^{M-1} E[x_i] = \frac{1}{M} \cdot M\theta = \theta$$

- ¿varianza del estimador?

$$\text{var}[\hat{\theta}_{ML}] = E[(\hat{\theta}_{ML} - \theta)^2] = E[\hat{\theta}_{ML}^2] - 2E[\hat{\theta}_{ML} \cdot \theta] + E[\theta^2]$$

ya que no tenemos info de θ , para nosotros es una constante

$$= E[\hat{\theta}_{ML}^2] - 2\theta \underbrace{E[\hat{\theta}_{ML}]}_{\theta} + \theta^2$$

$$= E[\hat{\theta}_{ML}^2] - \theta^2$$

Podríamos haber llegado usando

siempre

$$\boxed{E[x^2] = E^2[x] + \text{var}[x]}$$

↓
falta hallar este término
(valor cuadrático medio)

$$E[\hat{\theta}_{ML}^2] = \frac{1}{M^2} E\left[\left(\sum x_i\right)^2\right] \quad \text{se puede hacer con sumatorios e integrales}$$

$$= \frac{1}{M^2} E\left[\sum_{n=0}^{M-1} \sum_{m=0}^{M-1} x_n \cdot x_m\right]$$

$$= \frac{1}{M^2} \sum_n \sum_m E[x_n \cdot x_m]$$

$$\begin{cases} E[x_n^2] \text{ si } m=n \leftarrow M \text{ casos} \\ E[x_n x_m] \text{ si } m \neq n \leftarrow M^2 - M \text{ casos} \end{cases}$$

al ser independientes,
 $E[x_i x_j] = E[x_i] \cdot E[x_j]$

$$= \frac{1}{M^2} \left(\sum_{n=0}^{M-1} E[x_n^2] + \sum_{n=0}^{M-1} \sum_{m \neq n} E[x_n] \cdot E[x_m] \right)$$

$$\begin{aligned} \sum E[x_i] &= M \cdot E[x_i] \\ \sum \sum E[x_i] E[x_j] &= \sum E^2[x_i] \\ &= (M^2 - M) E^2[x_i] \end{aligned}$$

$$= \frac{1}{M^2} \left(M \cdot E[x_n^2] + (M^2 - M) E^2[x_n] \right)$$

Sabiendo que, para x_i de Poisson se tiene

$$E[x_i] = \theta$$

$$E[x_i^2] = \theta$$

$$= \frac{1}{M^2} \cdot M \cdot \theta + \frac{1}{M^2} M(M-1) \cdot \theta^2 = \frac{\theta}{M} + \frac{M-1}{M} \theta^2$$

Por tanto ya sabemos:

$$E[\hat{\theta}_{ML}] = \frac{\theta}{M} + \frac{M-1}{M} \theta^2$$

$$\text{var}[\hat{\theta}_{ML}] = \frac{\theta}{M} + \frac{M-1}{M} \theta^2 - \theta^2$$

Por tanto:

$$CCR = \frac{-1}{E\left[\frac{\partial^2 p(x/\theta)}{\partial \theta^2}\right]} = \frac{1}{E\left[\frac{\partial(-M + \frac{k}{\theta})}{\partial \theta}\right]} = \frac{1}{E\left[\frac{k}{\theta^2}\right]}$$

① no olvidar aplicar $E[\cdot]$

$$= \frac{1}{E\left[\frac{\sum x_i}{\theta^2}\right]} = \frac{\theta^2}{M\theta} = \frac{\theta}{M}$$

Por lo tanto el estimador es asintóticamente eficiente (cuando $M \rightarrow \infty$)

DETECCIÓN

TEORÍA DE LA DETECCIÓN: BIBLIOGRAFÍA

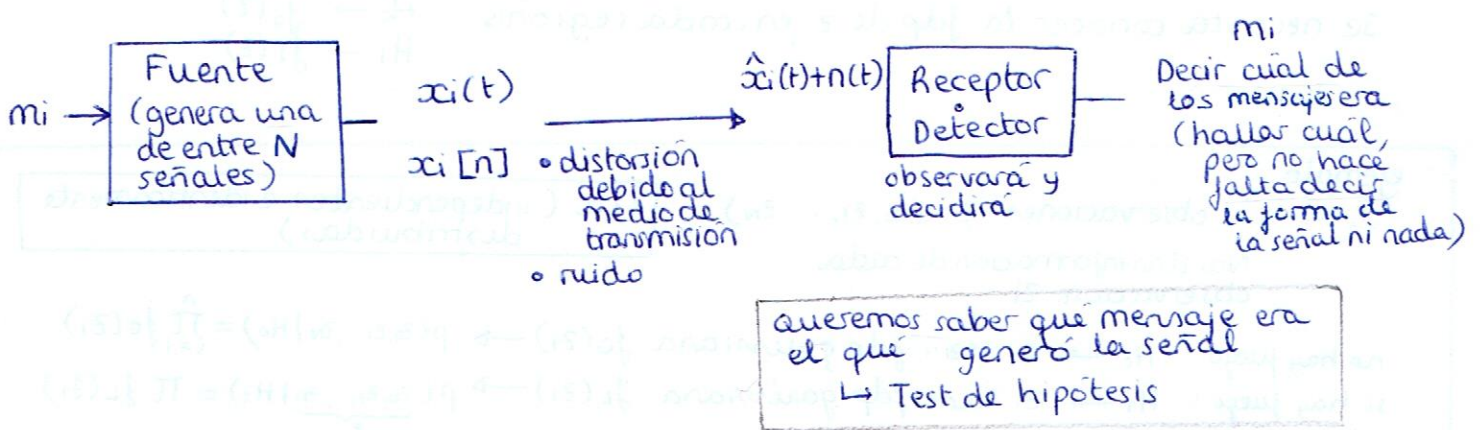
- [1] M.D. Srinath, P.K. Rajasekaran, ... "Introduction to statistical signal processing with applications" (Capítulo 3) Ed. Prentice Hall
- [2] Hippewstiel, Ralph Dieter: "Detection theory: applications and digital signal processing"

Índice

1. Introducción
2. Test de Hipótesis
3. Criterios de decisión
 - 3.1. Criterio de mínima probabilidad de error
 - 3.2. Criterio minimax
 - 3.3. Test de Neyman Pearson
 - 3.4. Detección Máxima a posteriori (MAP)
 - 3.5. Criterio de máxima verosimilitud (ML)
4. Hipótesis Múltiples
5. Test de múltiples hipótesis
6. Test de hipótesis compuestas

1. Introducción

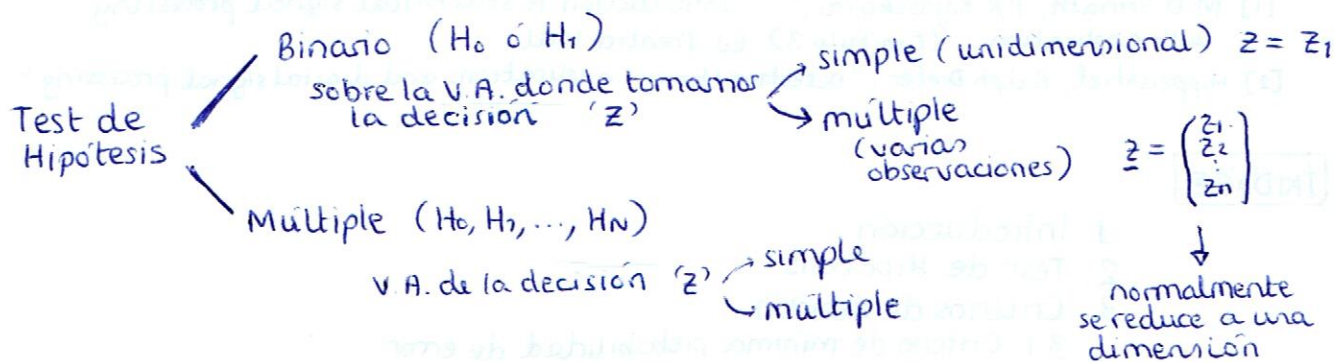
- deseamos elegir entre varias situaciones de manera automática
- seleccionar la opción correcta



La observación es una V.A.
(ya que la distorsión y el ruido no lo conocemos)

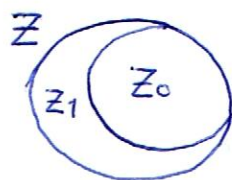
2. Test de Hipótesis

la detección y la clasificación son problemas que resolvemos aplicando Test de Hipótesis



Z : región de los valores posibles de z

- En Test de Hipótesis binario se divide en dos zonas Z_0 y Z_1



$$Z = Z_0 \cup Z_1$$

$$Z_0 \cap Z_1 = \emptyset$$

regiones de decisión

De forma que decidimos la hipótesis según a qué región pertenezca z

$$z \in Z_0 \rightarrow H_0$$

$$z \in Z_1 \rightarrow H_1$$

Se necesita conocer la fdp de z para cada hipótesis

$$\begin{aligned} H_0 &\rightarrow f_0(z) \\ H_1 &\rightarrow f_1(z) \end{aligned}$$

ejemplo:

vector de observaciones $\underline{z} = (z_0, z_1, \dots, z_n)$

Nos dan información de cada observación z_i

i.i.d (independientes e idénticamente distribuidas)

no hay fuego: $H_0 \rightarrow z_i$ tiene fdp gaussiana $f_G(z_i) \rightarrow p(z_0, z_1, \dots, z_n | H_0) = \prod_{i=1}^n f_G(z_i)$

sí hay fuego: $H_1 \rightarrow z_i$ tiene fdp laplaciana $f_L(z_i) \rightarrow p(\underbrace{z_0, z_1, \dots, z_n}_{\underline{z}} | H_1) = \prod f_L(z_i)$

¿con qué criterio escojo las regiones?

suele reducirse el problema a escoger un umbral

ejemplo: Test de Hipótesis Compuesto

$z_1, z_2, \dots, z_N$ iid gaussianas con media θ y varianza 1

$$H_1 \rightarrow \theta = 0 \rightarrow p(\underline{z} | H_0) = \prod_{i=1}^N N(0, 1) = \frac{1}{(\sqrt{2\pi})^N} e^{-\sum_{i=1}^N \frac{z_i^2}{2}}$$

$H_0 \rightarrow \theta \neq 0$ (Test de Hipótesis compuesto)

¿cuál será la fdp en ese caso?

↓
// Yo no sé cuál será la media de los z_i en ese caso!!

$$p(\underline{z} | H_1) = \prod_{i=1}^N N(\theta, 1)$$

↑
todas las z_i tendrán la misma θ , pero yo no sé cuál es.
Previamente tendré que estimar θ (teoría de la estimación)

recuerda

$$N(m, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(z-m)^2}{2\sigma^2}}$$

↑
comprueba

3. Criterios de decisión

3.1 criterio MAP (el que nos dice el sentido común)

si $p(H_0 | z) > p(H_1 | z) \rightarrow H_0$

si $p(H_1 | z) > p(H_0 | z) \rightarrow H_1$

$$\Leftrightarrow \frac{p(H_1 | z)}{p(H_0 | z)} \underset{H_0}{\overset{H_1}{\geq}} 1$$

criterio MAP

↑
¿cómo podemos conocer esto?

Por Bayes: $p(H_i | z) = \frac{p(z | H_i) \cdot p(H_i)}{p(z)}$

Por tanto

$$\frac{p(z | H_1)}{p(z | H_0)} \underset{H_0}{\overset{H_1}{\geq}} \frac{p(H_0)}{p(H_1)}$$

$$\Lambda(z) \geq \lambda_{\text{umbral}}$$

razón de verosimilitud

ejemplo:

Canal de comunicaciones:

Fuente: T segundos $\begin{matrix} \swarrow 1 \\ \searrow 0 \end{matrix}$

Canal: Añade ruido $v(t) \rightarrow$ gaussiana media nula
varianza uno

Recibido: $r(t) = \begin{cases} 1 + v(t) \\ 0 \\ v(t) \end{cases}$ en T segundos

Con una sola observación decidir H_0 o H_1

$$H_0: z = v$$

$$H_1: z = 1 + v$$

Podemos deducir

$$\bullet p(z | H_0) = p(v) = \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}}$$

$\bullet p(z | H_1)$: será también gaussiana pero con la media desplazada

$$E[z] = E[1+v] = E[1] + E[v] = 1$$

$$\text{var}[z] = E[(z-1)^2] = E[v^2] = 1$$

$$p(z | H_1) = \frac{1}{\sqrt{2\pi}} e^{-\frac{(z-1)^2}{2}}$$

Por lo tanto podemos calcular

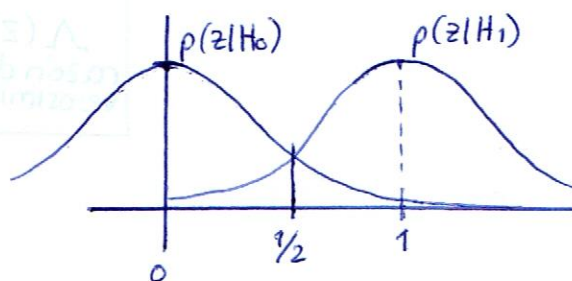
$$\Lambda(z) = \frac{p(z | H_1)}{p(z | H_0)} = \frac{e^{-\frac{(z-1)^2}{2}}}{e^{-\frac{z^2}{2}}} = e^{z - \frac{1}{2}}$$

$$\lambda = \frac{p(H_0)}{p(H_1)}$$

$$e^{z - \frac{1}{2}} \underset{\substack{H_1 \\ H_0}}{\geq} \frac{p(H_0)}{p(H_1)}$$

↓ tomando ln

$$z \geq \frac{1}{2} + \ln\left(\frac{p(H_0)}{p(H_1)}\right)$$



El principal problema es obtener la fdp. bajo cada hipótesis

- Criterio de Bayes

Vamos a tener en cuenta el coste que supone equivocarse, y vamos a minimizar el coste.

Llamamos D_i a la decisión tomada $i=0,1$
 D_0 ó D_1 según creamos que la hipótesis es H_0 ó H_1

Pueden pasar 4 cosas:

1. H_0 hipótesis verdadera y decidimos $D_0 \rightarrow p(H_0, D_0)$
2. H_1 hipótesis verdadera y decidimos $D_1 \rightarrow p(H_1, D_1)$
3. H_0 hipótesis verdadera y decidimos $D_1 \rightarrow p(H_0, D_1) \rightarrow$ falsa alarma $\nearrow C_{10}$
4. H_1 hipótesis verdadera y decidimos $D_0 \rightarrow p(H_1, D_0) \rightarrow$ pérdida (miss) $\nwarrow C_{01}$

prob. conjunta
intersección

↓
 cada decisión tendrá un coste según el contexto.

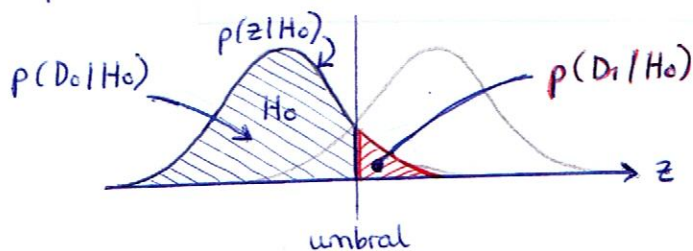
C_{ij} : coste de ser la hipótesis " j " verdadera y nosotros haber decidido " i "

Vamos a intentar minimizar el coste medio $\bar{C}$

$$\bar{C} = C_{00} \cdot p(H_0, D_0) + C_{11} \cdot p(H_1, D_1) + C_{10} \cdot p(H_0, D_1) + C_{01} \cdot p(H_1, D_0)$$

Por Bayes:
$$p(D_i | H_j) = \frac{p(D_i, H_j)}{p(H_j)}$$

$$\bar{C} = C_{00} \cdot p(H_0) \cdot p(D_0 | H_0) + C_{11} \cdot p(H_1) \cdot p(D_1 | H_1) + C_{10} \cdot p(H_0) \cdot p(D_1 | H_0) + C_{01} \cdot p(H_1) \cdot p(D_0 | H_1)$$



Es fácil ver que:

$$p(D_i | H_j) = \int_{Z_i} p(z | H_j) dz$$

Además, puesto que $Z = Z_0 \cup Z_1$ y $Z_0 \cap Z_1 = \emptyset$

$$\int_Z p(z | H_j) dz = 1 = \int_{Z_0} p(z | H_j) dz + \int_{Z_1} p(z | H_j) dz$$

$$\Rightarrow \int_{Z_1} p(z | H_j) = 1 - \int_{Z_0} p(z | H_j) dz$$

Por lo tanto se deduce que podemos expresarlo todo en función de la integral en una única zona:

$$\bar{C} = C_{00} \cdot p(H_0) \cdot \int_{Z_0} p(z|H_0) dz$$

$$+ C_{11} p(H_1) \cdot \left[1 - \int_{Z_0} p(z|H_1) dz \right]$$

$$+ C_{10} \cdot p(H_0) \cdot \left[1 - \int_{Z_0} p(z|H_0) dz \right]$$

$$+ C_{01} p(H_1) \cdot \int_{Z_0} p(z|H_1) dz$$

el coste de equivocarnos es mayor que el de acertar

$$\bar{C} = \underbrace{C_{11} p(H_1) + C_{10} p(H_0)}_{\text{cantidad fija (viene de los unos)}} + \underbrace{\int_{Z_0} p(H_1) \cdot (C_{01} - C_{11}) p(z|H_1) - p(H_0) \cdot (C_{10} - C_{00}) p(z|H_0) dz}_{\text{minimizar esto (puede ser negativo)}}$$

Nuestro criterio será:

$$p(H_1) \cdot (C_{01} - C_{11}) p(z|H_1) \underset{H_0}{\overset{H_1}{\geq}} p(H_0) \cdot (C_{10} - C_{00}) p(z|H_0)$$

que queda como:

$$\frac{p(z|H_1)}{p(z|H_0)} \underset{H_0}{\overset{H_1}{\geq}} \frac{C_{10} - C_{00}}{C_{01} - C_{11}} \cdot \frac{p(H_0)}{p(H_1)}$$

Criterio de Bayes

$$\Lambda(z) \geq \lambda$$



ejercicio: una fuente emite señales con 2 posibles varianzas



señal fdp
gausiana
media $\mu=0$

H_1 : varianza σ_1^2

$\sigma_1^2 > \sigma_0^2$

H_0 : varianza σ_0^2

Hay que hallar:

$$p(z|H_0) = N(0; \sigma_0^2) = \frac{1}{\sqrt{2\pi}\sigma_0} \exp\left(-\frac{z^2}{2\sigma_0^2}\right)$$

$$p(z|H_1) = N(0; \sigma_1^2) = \frac{1}{\sqrt{2\pi}\sigma_1} \exp\left(-\frac{z^2}{2\sigma_1^2}\right)$$

Por tanto:

$$\Lambda(z) = \frac{p(z|H_1)}{p(z|H_0)} = \frac{\sigma_0}{\sigma_1} \exp\left(-\frac{1}{2}z^2\left(\frac{1}{\sigma_1^2} - \frac{1}{\sigma_0^2}\right)\right)$$

Aplicando la regla de decisión (con el $\ln$) se obtiene

$$\ln \Lambda(z) \underset{H_0}{\overset{H_1}{\gtrless}} \ln \lambda$$

$$\lambda = \frac{C_{10} - C_{00}}{C_{01} - C_{11}} \cdot \frac{P(H_0)}{P(H_1)}$$

$$\ln\left(\frac{\sigma_0}{\sigma_1}\right) - \frac{1}{2}z^2\left(\frac{1}{\sigma_1^2} - \frac{1}{\sigma_0^2}\right) \underset{H_0}{\overset{H_1}{\gtrless}} \ln \lambda$$

$$\frac{1}{2}z^2\left(\frac{\sigma_1^2 - \sigma_0^2}{\sigma_1^2 \cdot \sigma_0^2}\right) \underset{H_0}{\overset{H_1}{\gtrless}} \ln \lambda - \ln\left(\frac{\sigma_0}{\sigma_1}\right)$$

$$z^2 \underset{H_0}{\overset{H_1}{\gtrless}} \underbrace{\frac{2\sigma_1^2\sigma_0^2}{\sigma_1^2\sigma_0^2} \ln\left(\frac{\sigma_1}{\sigma_0}\lambda\right)}_{\text{umbral } \gamma}$$

si z era
gausiana,
esta sera
chi-cuadrado

$$\boxed{z^2 \underset{H_0}{\overset{H_1}{\gtrless}} \gamma}$$

¿Cómo es de bueno el sistema?

P_F Probabilidad de
falsa alarma

$$P_I = Pr(D_1|H_0) = \int_{Z_1} p(z|H_0) dz = \int_{\gamma}^{\infty} p(\Lambda(z)|H_0) dz$$

P_M Probabilidad de
pérdida (miss)

$$P_{II} = Pr(D_0|H_1) = \int_{Z_0} p(z|H_1) dz = \int_{-\infty}^{\gamma} p(\Lambda(z)|H_1) dz$$

Probabilidad
de error

$$P_e = P_I \cdot p(H_0) + P_{II} \cdot p(H_1)$$

⚠
 P_F, P_M no tienen en
cuenta $pr(H_0)$ frente
a $p(H_1)$.
 P_e sí

Probabilidad
de detección

$$P_D = Pr(D_1|H_1)$$

Sabemos que el criterio de Bayes es óptimo, pero se requiere conocer ciertos datos.

Hay otros criterios (ya no óptimos) que necesitan menos datos

• Criterio de mínima probabilidad de error

supone $C_{00} = C_{11} = 0$
 $C_{10} = C_{01} = 1$ (observador ideal)

$$\hookrightarrow \lambda = \frac{P(H_0)}{P(H_1)} \rightarrow \text{Equivale al MAP}$$

Requiere conocer las prob. a priori

El coste medio

$$\begin{aligned} \bar{C} &= C_{00}(1-P_F) + C_{10}P_F + p(H_1) \cdot [(C_{11}-C_{00}) + (C_{01}-C_{11})P_M - (C_{10}-C_{00})P_F] \\ &= P_F + p(H_1) \cdot [P_M - P_F] \\ &= P_F(1-p(H_1)) + p(H_1)P_M \\ &= P_F p(H_0) + P_M p(H_1) \\ &= P_e \end{aligned}$$

El coste medio es igual a la prob. de error

ejemplo:

$$\begin{aligned} H_1 &: m + n \\ H_0 &: n \end{aligned}$$

↑ ruido $N(0, \sigma^2)$

Determinar la regla de decisión aplicando el criterio de mínima probabilidad de error, sabiendo $p(H_0) = p(H_1) = 0.5$

$$\Lambda(z) = \frac{p(z|H_1)}{p(z|H_0)} = \frac{N(m, \sigma^2)}{N(0, \sigma^2)} = \frac{\frac{1}{\sqrt{2\pi}\sigma} \exp(-\frac{(z-m)^2}{2\sigma^2})}{\frac{1}{\sqrt{2\pi}\sigma} \exp(-\frac{z^2}{2\sigma^2})} = \exp\left(-\frac{(z-m)^2}{2\sigma^2} + \frac{z^2}{2\sigma^2}\right)$$

Por tanto el criterio queda $\Lambda(z) \underset{H_0}{\overset{H_1}{\geq}} \frac{p(H_0)}{p(H_1)} = 1$

$$\ln \Lambda(z) \underset{H_0}{\overset{H_1}{\geq}} 0$$

$$-\frac{z^2}{2\sigma^2} + \frac{2mz - m^2 + z^2}{2\sigma^2} \underset{H_0}{\overset{H_1}{\geq}} 0$$

$$\boxed{z \underset{H_0}{\overset{H_1}{\geq}} \frac{m}{2}}$$

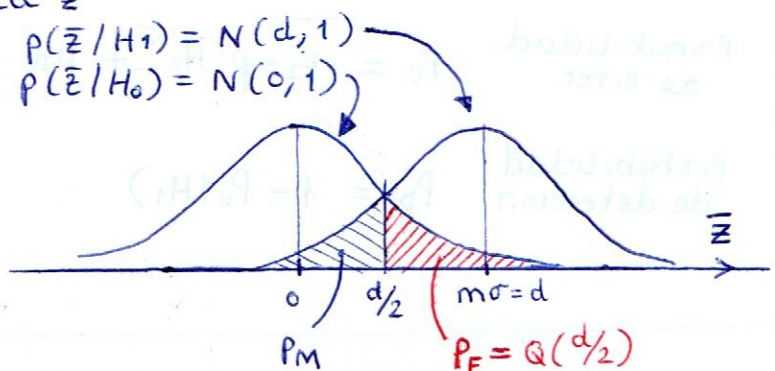
Típicamente se normaliza z

$$\boxed{\bar{z} = z/\sigma}$$

para que

$$\boxed{\bar{z} \underset{H_0}{\overset{H_1}{\geq}} \frac{m}{2\sigma} = \frac{d}{2}}$$

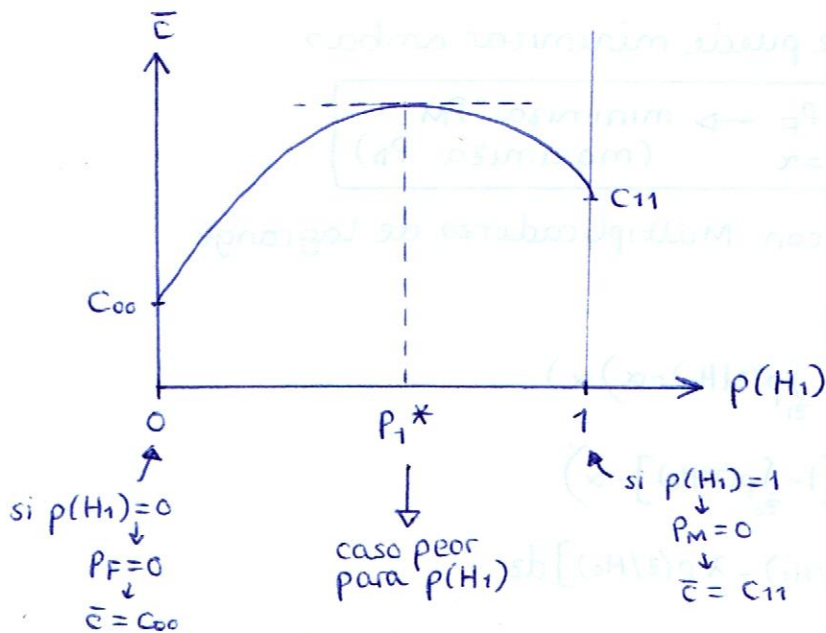
$$d = \frac{m}{\sigma}$$



• Criterio minimax

Se aplica cuando no conocemos las probabilidades a priori

Recuerda: $\bar{c} = C_{00}(1-P_F) + C_{10}P_F + p(H_1)[(C_{11}-C_{00}) + (C_{01}-C_{11})P_M - (C_{10}-C_{00})P_F]$



El criterio minimax supone que $p(H_1) = p_1^*$

p_1^* se obtiene maximizando $\bar{c}$

$$\frac{d\bar{c}}{dp(H_1)} = 0$$

$$= (C_{11}-C_{00}) + (C_{01}-C_{11})P_M - (C_{10}-C_{00})P_F$$

si además no conocemos los costes

y suponemos $C_{00} = C_{11} = 0$
 $C_{01} = C_{10} = 1$

$$\frac{d\bar{c}}{dp(H_1)} = P_M - P_F = 0$$

$$\Rightarrow \boxed{P_M = P_F}$$

• Criterio de Neyman-Pearson

- C_{ij} desconocidos
- $p(H_i)$ desconocidas

¿Puedo minimizar a la vez $P_F = \alpha$ y $P_M = \beta$?

Neumann vio que no se puede minimizar ambas.

se le ocurrió:

$\begin{array}{l} \text{fijar} \\ P_F = \alpha \end{array} \rightarrow \begin{array}{l} \text{minimizar } P_M \\ (\text{maximizar } P_D) \end{array}$

Pudo hacerlo con Multiplicadores de Lagrange

$$J = P_M + \lambda(P_F - \alpha)$$

$$J = \int_{z_0} p(z|H_1) dz + \lambda \left(\int_{z_1} p(z|H_0) dz - \alpha \right)$$

$$= \int_{z_0} p(z|H_1) dz + \lambda \left(\left[1 - \int_{z_0} p(z|H_0) dz \right] - \alpha \right)$$

$$= \lambda(1-\alpha) \int_{z_0} [p(z|H_1) - \lambda p(z|H_0)] dz$$

Puedo escoger el λ que minimice J condicionada a la decisión sin más que hacer:

$$p(z|H_1) \underset{H_0}{\overset{H_1}{\gtrless}} p(z|H_0) \lambda \Rightarrow \frac{p(z|H_1)}{p(z|H_0)} \underset{H_0}{\overset{H_1}{\gtrless}} \lambda \leftarrow \text{el umbral de Neyman-Pearson}$$

Es decir, hay que resolver la ecuación:

$$\int_{\lambda}^{\infty} p(z|H_0) dz = \alpha$$

$\uparrow$
 umbral que debemos determinar

$\uparrow$
 prob. de falsa alarma (dato)
 (exijimos un valor)

Hipótesis múltiple

- Supongamos $M > 2$ eventos $H_0, H_1, \dots, H_{M-1}$
- Asumimos:
 - a) $p(H_0), p(H_1), \dots, p(H_{M-1})$ conocidos ($\sum p_i = 1$)
 - b) C_{ij} es el coste de D_i/H_j
 - c) $p(z|H_i)$

• El objetivo de decisión es minimizar el coste medio $\bar{c}$

$$\bar{c} = \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} C_{ij} \cdot P(H_j) \cdot P(D_i/H_j)$$

Sabiendo que $P(D_i/H_j) = \int_{Z_i} p(z|H_j) dz$

siendo Z_i la región de decisión de H_i
 $Z = \bigcup_{i=0}^{M-1} Z_i$
 $Z_i \cap Z_j = \emptyset \quad \forall i \neq j$

$$\bar{c} = \sum_{i=0}^{M-1} C_{ii} P(H_i) \int_{Z_i} p(z|H_i) dz + \sum_{\substack{j=0 \\ j \neq i}}^{M-1} C_{ij} P(H_j) \int_{Z_j} p(z|H_j) dz$$

y puesto que: $\int_{Z_i} p(z|H_i) dz = 1 - \sum_{\substack{j=0 \\ j \neq i}}^{M-1} \int_{Z_j} p(z|H_j) dz$

$$\bar{c} = \underbrace{\sum_{i=0}^{M-1} C_{ii} P(H_i)}_{\text{coste fijo (no depende de } Z_i)} + \underbrace{\sum_{i=0}^{M-1} \int_{Z_i} \sum_{\substack{j=0 \\ j \neq i}}^{M-1} P(H_j) (C_{ij} - C_{jj}) P(z|H_j) dz}_{\text{coste variable}}$$

Escogeremos las regiones Z_i que minimicen

tomando referencia $p(z|H_0)$
 $I_i(z) = \sum_{\substack{j=0 \\ j \neq i}}^{M-1} P(H_j) (C_{ij} - C_{jj}) P(z|H_j)$ (sin referencia)

minimizar:
 $J_i(z) = \frac{I_i(z)}{p(z|H_0)} = \sum_{\substack{j=0 \\ j \neq i}}^{M-1} P(H_j) \cdot (C_{ij} - C_{jj}) \left(\frac{p(z|H_j)}{p(z|H_0)} \right) \Lambda_j(z)$

 siendo $\Lambda_0(z) = 1$ (referencia!)

Decidimos

$$z \rightarrow H_0 \text{ si } J_0 < J_1, J_0 < J_2, J_0 < J_3, \dots, J_0 < J_{M-1}$$

$$z \rightarrow H_1 \text{ si } J_1 < J_0, J_1 < J_3, \dots, J_1 < J_{M-1}$$

$$z \rightarrow H_{M-1} \text{ si } J_{M-1} < J_0, J_{M-1} < J_1, \dots, J_{M-1} < J_{M-2}$$

Ejemplo: Detectar entre 3 señales en fondo de ruido $n \sim N(0, \sigma_n^2)$

se decidirá con
un vector de
N muestras

$$\underline{z} = \begin{bmatrix} z_1 \\ z_2 \\ \vdots \\ z_N \end{bmatrix}$$

$$\begin{aligned} S_0 &= \text{DC } m_0 \\ S_1 &= \text{DC } m_1 \\ S_2 &= \text{DC } m_2 \end{aligned}$$

$$\begin{aligned} P(H_i) &= \frac{1}{3} \\ C_{ii} &= 1 \end{aligned}$$

Los costes

$$\begin{aligned} C_{ii} &= 0 \quad (\text{acertar}) \\ C_{ij} &= 1 \quad (\text{equivocarse}) \end{aligned}$$

Hay que calcular:

$$J_i(\underline{z}) = \sum_{j=0, j \neq i}^2 \frac{1}{3} \Lambda_j(\underline{z})$$

$$\begin{cases} J_0(\underline{z}) = \frac{1}{3} \Lambda_1(\underline{z}) + \frac{1}{3} \Lambda_2(\underline{z}) \\ J_1(\underline{z}) = \frac{1}{3} \underbrace{\Lambda_0(\underline{z})}_1 + \frac{1}{3} \Lambda_2(\underline{z}) \\ J_2(\underline{z}) = \frac{1}{3} \underbrace{\Lambda_0(\underline{z})}_1 + \frac{1}{3} \Lambda_1(\underline{z}) \end{cases}$$

Decidiremos H_0 cuando $J_0(\underline{z}) < J_1(\underline{z})$ y $J_0(\underline{z}) < J_2(\underline{z})$

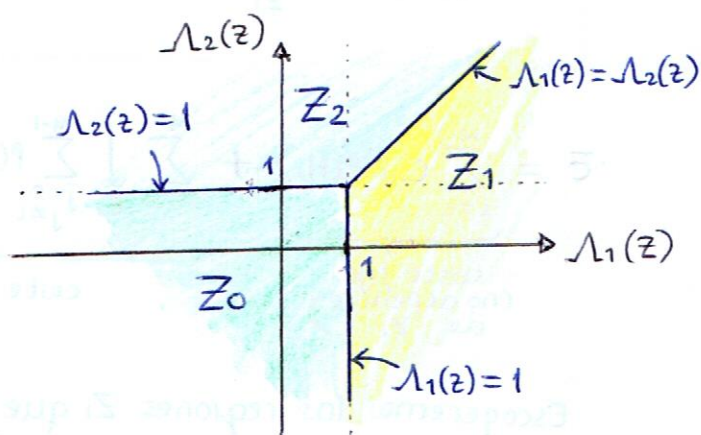
$$\begin{aligned} \bullet \quad \frac{1}{3} \Lambda_1(\underline{z}) + \frac{1}{3} \Lambda_2(\underline{z}) &< \frac{1}{3} + \frac{1}{3} \Lambda_2(\underline{z}) \rightarrow \Lambda_1(\underline{z}) < 1 \\ \bullet \quad \frac{1}{3} \Lambda_1(\underline{z}) + \frac{1}{3} \Lambda_2(\underline{z}) &< \frac{1}{3} + \frac{1}{3} \Lambda_1(\underline{z}) \rightarrow \Lambda_2(\underline{z}) < 1 \end{aligned}$$

$$Z_0 : \begin{aligned} \Lambda_1(\underline{z}) &< 1 \\ \Lambda_2(\underline{z}) &< 1 \end{aligned}$$

Procediendo igual

$$Z_1 : \begin{aligned} \Lambda_1(\underline{z}) &> 1 \\ \Lambda_2(\underline{z}) &< \Lambda_1(\underline{z}) \end{aligned}$$

$$Z_2 : \begin{aligned} \Lambda_2(\underline{z}) &> 1 \\ \Lambda_2(\underline{z}) &> \Lambda_1(\underline{z}) \end{aligned}$$



Habría que expresar estas regiones
como condición en $\underline{z}$ (no como
condición en $\Lambda(\underline{z})$)

$$\Lambda_1(\underline{z}) = \frac{p(\underline{z} | H_1)}{p(\underline{z} | H_0)}$$

$$\Lambda_2(\underline{z}) = \frac{p(\underline{z} | H_2)}{p(\underline{z} | H_0)}$$

$$\text{con } p(\underline{z} | H_i) = \underbrace{\prod_{j=1}^N p(z_j | H_i)}_{N(0, \sigma_n^2)} = \prod_{j=1}^N \frac{1}{(\sqrt{2\pi}\sigma_n)^2} \cdot \exp\left(-\frac{(z_j - m_i)^2}{2\sigma_n^2}\right)$$

Por tanto

$$p(\underline{z} | H_i) = \frac{1}{(2\pi\sigma^2)^{N/2}} \exp\left(-\sum_{j=1}^N \frac{(z_j - m_i)^2}{2\sigma^2}\right)$$

multiplicación
de exponenciales
es el sumatorio

$$\begin{aligned}\Lambda_1(\underline{z}) &= \frac{p(\underline{z} | H_1)}{p(\underline{z} | H_0)} = \exp\left[-\sum_{i=1}^N \frac{(z_i - m_1)}{2\sigma^2} + \sum_{i=1}^N \frac{(z_i - m_0)}{2\sigma^2}\right] \\ &= \exp\left[-\frac{1}{2\sigma^2} \sum_{i=1}^N (z_i^2 - 2m_1 z_i + m_1^2 - z_i^2 + 2m_0 z_i - m_0^2)\right]\end{aligned}$$

entonces $\Lambda_1(\underline{z}) = 1$

$\ln \Lambda_1(\underline{z}) = 0$

$-\frac{1}{2\sigma^2} \sum_{i=1}^N [2(m_0 - m_1)z_i + m_1^2 - m_0^2] = 0$

$2(m_0 - m_1) \sum z_i = (m_0^2 - m_1^2)N$

$\underbrace{\frac{1}{N} \sum_{i=1}^N z_i}_{\text{la media } \bar{z}} = \frac{m_0 + m_1}{2}$

igualmente

$\Lambda_2(\underline{z}) = 1 \rightarrow \frac{1}{N} \sum_{i=1}^N z_i = \frac{m_0 + m_2}{2}$

$\Lambda_2(\underline{z}) = \Lambda_1(\underline{z}) \rightarrow \frac{1}{N} \sum_{i=1}^N z_i = \frac{m_1 + m_2}{2}$

Por lo tanto podemos escribir las regiones como:

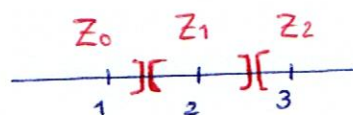
$Z_0 \rightarrow \left[\begin{matrix} \Lambda_1(\underline{z}) < 1 \\ \Lambda_2(\underline{z}) < 1 \end{matrix} \right] \rightarrow \left[\begin{matrix} \bar{z} < \frac{m_0 + m_1}{2} \\ \bar{z} < \frac{m_0 + m_2}{2} \end{matrix} \right]$

$Z_1 \rightarrow \left[\begin{matrix} \Lambda_1(\underline{z}) > 1 \\ \Lambda_2(\underline{z}) < \Lambda_1(\underline{z}) \end{matrix} \right] \rightarrow \left[\begin{matrix} \bar{z} > \frac{m_0 + m_1}{2} \\ \bar{z} < \frac{m_1 + m_2}{2} \end{matrix} \right]$

$Z_2 \rightarrow \left[\begin{matrix} \Lambda_2(\underline{z}) > 1 \\ \Lambda_2(\underline{z}) > \Lambda_1(\underline{z}) \end{matrix} \right] \rightarrow \left[\begin{matrix} \bar{z} > \frac{m_0 + m_1}{2} \\ \bar{z} > \frac{m_1 + m_2}{2} \end{matrix} \right]$

ejemplo: $m_0 = 1$
 $m_1 = 2$
 $m_2 = 3$

$H_0 \rightarrow \frac{\bar{z} < 3/2}{\bar{z} < 2} \rightarrow \bar{z} < 3/2$
 $H_1 \rightarrow \frac{\bar{z} > 3/2}{\bar{z} < 5/2} \rightarrow \frac{3}{2} < \bar{z} < \frac{5}{2}$
 $H_2 \rightarrow \frac{\bar{z} > 5/2}{\bar{z} > 2} \rightarrow \bar{z} > \frac{5}{2}$



Anexo: Distribución chi-cuadrado (producto de gaussianas)

Supongamos

$$\ell = \underline{z}^T \cdot \underline{z} = \sum_{i=1}^N z_i^2 \rightarrow \text{Distribución chi-cuadrado con } N \text{ grados de libertad}$$

$\uparrow$
 z_i gaussianas

$$p(\ell / H_i) = \begin{cases} \frac{\ell^{\frac{(N-1)}{2}}}{(2^{\frac{N}{2}} \sigma_i^N \Gamma(\frac{N}{2}))} \exp\left(-\frac{\ell}{2\sigma_i^2}\right) & \ell \geq 0 \\ 0 & \text{resto} \end{cases}$$

al calcular $\mathcal{L}_i(z)$
se va todo menos esto

se supone
sabido

Es muy útil cuando hay que detectar CAMBIOS en la
varianza de la observación

Problema 1. $H_0 \rightarrow$ transmitida continua DC de 0V
 $H_1 \rightarrow$ transmitida continua DC de 0'75V } equiprobables $C_{00} = C_{11} = 0$
 $C_{01} = C_{10} = 1$
 · Se suma ruido uniforme de varianza 0'5 y media nula
 · sólo una muestra en detección

a) Deduzca el detector óptimo y qué criterio está usando
 Obtenga el umbral adecuado

Dados los costes $C_{00} = C_{11} = 0$, $C_{01} = C_{10} = 1$ se deduce que el coste mínimo (criterio de Bayes) coincide con la mínima probabilidad de error (criterio MAP). Además, por ser las hipótesis equiprobables este es a su vez equivalente al criterio ML.

Por tanto: $\Lambda(z) \underset{H_0}{\overset{H_1}{>}} \lambda = 1$

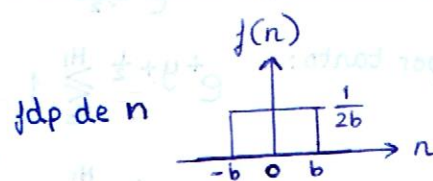
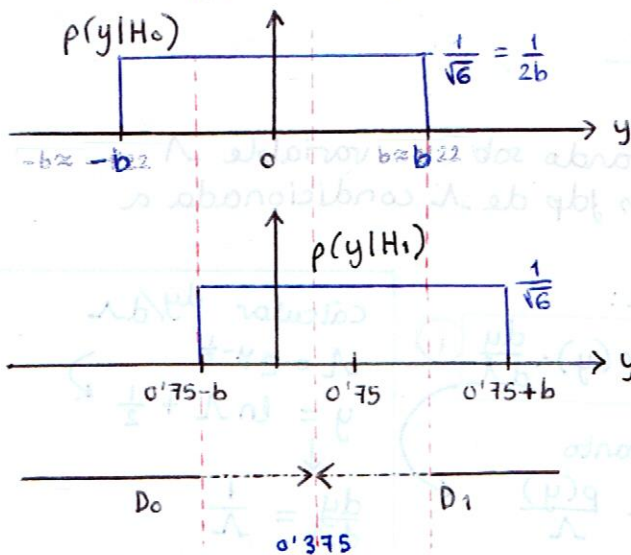
$$\frac{p(z|H_1)}{p(z|H_0)} \underset{H_0}{\overset{H_1}{>}} 1$$

Se tiene:

$$H_0: y = n$$

$$H_1: y = 0'75 + n$$

Parece claro que el valor DC sólo desplazará las jdp de n , por lo tanto:



sabiendo que para una V.A. uniforme se tiene $\sigma_n^2 = \frac{(b-a)^2}{12} = \frac{(2b)^2}{12}$

podemos obtener b :

$$\sigma_n^2 = \frac{b^2}{3} = 0'5 \rightarrow b = \sqrt{1'5} = \sqrt{\frac{3}{2}}$$

Escogeremos H_0 si $y < 0'75 - b$

Escogeremos H_1 si $y > b$

¿qué hacemos en la zona intermedia donde $\Lambda(z) = 1$?

Podemos tomar simplemente la mitad como umbral

$$y_{umb} = \frac{0'75 - b + b}{2} = 0'375$$

Por tanto: $D_0 \equiv y < 0'375$
 $D_1 \equiv y > 0'375$

b) Probabilidad de error

$$P_e = P(H_0) \cdot P_{FA} + P(H_1) \cdot P_M$$

$$= 0'347$$

$$P(H_0) = P(H_1) = 0'5$$

$$P_{FA} = \int_{D_1} p(y|H_0) dy = \int_{0'375}^b \frac{1}{2b} dy = 0'347$$

$$P_M = \int_{D_0} p(y|H_1) dy = \int_{0'75-b}^{0'375} \frac{1}{2b} dy = 0'347$$

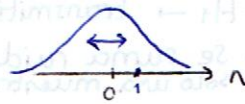
Problema 2 :

· Sólo una muestra en detección

$$H_0 : y = n$$

$$H_1 : y = 1 + n$$

distribución Normal/Gaussiana
 $n = N(0, 1)$
 ↑ media ↑ varianza



a) Detector ML y umbral

Parece claro que: $p(y|H_0) = \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} = N(0, 1)$

$$p(y|H_1) = \frac{1}{\sqrt{2\pi}} e^{-\frac{(y-1)^2}{2}} = N(1, 1)$$

$$ML \rightarrow \Lambda(y) = \frac{p(y|H_1)}{p(y|H_0)} \underset{H_0}{\underset{H_1}{\gtrless}} \lambda = 1$$

$$\Lambda(y) = \frac{e^{-\frac{(y-1)^2}{2}}}{e^{-\frac{y^2}{2}}} = e^{+\frac{y^2}{2} - \frac{(y-1)^2}{2}} = e^{\frac{-y^2 + y^2 + 2y - 1}{2}} = e^{+y - \frac{1}{2}}$$

por tanto: $e^{+y - \frac{1}{2}} \underset{H_0}{\underset{H_1}{\gtrless}} 1$

$$y - \frac{1}{2} \underset{H_0}{\underset{H_1}{\gtrless}} 0 \rightarrow \boxed{y \underset{H_0}{\underset{H_1}{\gtrless}} \frac{1}{2}}$$

calcular $P(D_1|H_0) = \int_{0.5}^{\infty} p(y|H_0) dy = Q(0.5) = 0.3085$



b) Volver a calcular $P(D_1|H_0)$ pero integrando sobre la variable Λ (una exponencial). Deberá obtener las fdp de Λ condicionada a cada hipótesis.

$$\Lambda(y) = e^{y - \frac{1}{2}}$$

$$\Lambda \underset{H_0}{\underset{H_1}{\gtrless}} 1$$

Sabemos que:

$$p(\Lambda) = p(y) \cdot \frac{dy}{d\Lambda}$$

por tanto

$$p(\Lambda) = \frac{p(y)}{\Lambda}$$

calcular $dy/d\Lambda$

$$\Lambda = e^{y - \frac{1}{2}}$$

$$y = \ln \Lambda + \frac{1}{2}$$

$$\downarrow$$

$$\frac{dy}{d\Lambda} = \frac{1}{\Lambda}$$

entonces:

$$p(\Lambda|H_0) = \frac{p(y|H_0)}{\Lambda} = \frac{1}{\Lambda \sqrt{2\pi}} e^{-\frac{y^2}{2}} = \frac{1}{\Lambda \sqrt{2\pi}} e^{-\frac{(\ln \Lambda + \frac{1}{2})^2}{2}}$$

expresar 'y' en función de Λ

$$p(\Lambda|H_1) = \frac{p(y|H_1)}{\Lambda} = \dots = \frac{1}{\Lambda \sqrt{2\pi}} e^{-\frac{(\ln \Lambda - \frac{1}{2})^2}{2}}$$

$$P(D_1|H_0) = \int_{D_1} p(\Lambda|H_0) d\Lambda = \int_1^{\infty} \frac{1}{\Lambda \sqrt{2\pi}} e^{-\frac{(\ln \Lambda + \frac{1}{2})^2}{2}} d\Lambda = Q(\ln 1 + \frac{1}{2}) = Q(0.5)$$

cambio de variable

$$\mu = \ln \Lambda + \frac{1}{2}$$

$$d\mu = \frac{1}{\Lambda} d\Lambda$$

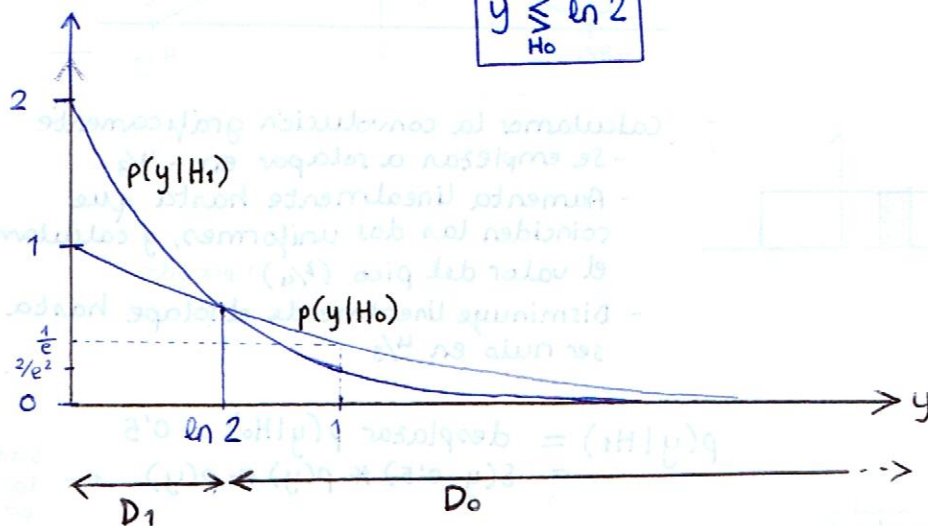
Problema 3. $H_0: f(y) = e^{-y} \cdot u(y)$ $u(y)$ es el escalón
 $H_1: f(y) = 2e^{-2y} u(y)$

a) Detector y umbral, representar gráficamente

$$\Lambda(z) = \frac{p(y|H_1)}{p(y|H_0)} = \frac{2e^{-2y}}{e^{-y}} = 2e^{-y} \stackrel{H_1}{\geq} \stackrel{H_0}{\leq} \lambda = \frac{(C_{10} - C_{00})P(H_0)}{(C_{01} - C_{11})P(H_1)} = 1$$

$$\ln 2 - y \stackrel{H_1}{\geq} \stackrel{H_0}{\leq} 0$$

$$y \stackrel{H_1}{\geq} \stackrel{H_0}{\leq} \ln 2$$



b) P_{FA}

$$P_D = Pr(D_0|H_0) = \int_{D_0} p(y|H_0) dy = \int_{\ln 2}^{\infty} e^{-y} dy = -[e^{-y}]_{\ln 2}^{\infty} = 1/2$$

$$P_{FA} = Pr(D_1|H_0) = \int_{D_1} p(y|H_0) dy = \int_0^{\ln 2} e^{-y} dy = -[e^{-y}]_0^{\ln 2} = 1 - 1/2 = 1/2$$

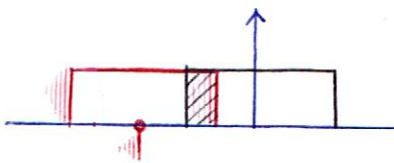
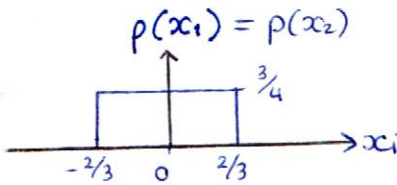
Problema 4.

Sean x_1 y x_2 V.A. indep. uniformemente distribuidas $x_i \sim U(-\frac{2}{3}, \frac{2}{3})$

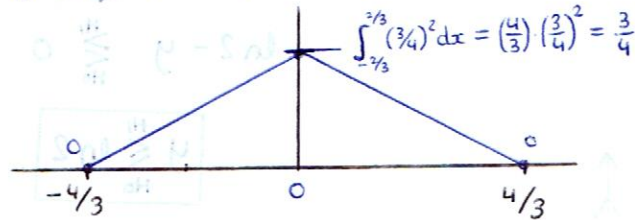
$$H_0 : y = x_1 + x_2$$

$$H_1 : y = 0.5 + x_1 + x_2$$

a) Detector óptimo y el umbral



$$p(y|H_0) = p(x_1) * p(x_2)$$



Calcular la convolución gráficamente

- Se empiezan a solapar en $-4/3$
- Aumenta linealmente hasta que coinciden las dos uniformes, y calculamos el valor del pico $\int_{-\infty}^{\infty} p(x_1) \cdot p(x_2) dx$
- Disminuye linealmente el solape hasta ser nulo en $4/3$

$$p(y|H_1) = \text{desplazar } p(y|H_0) \text{ a } 0.5$$

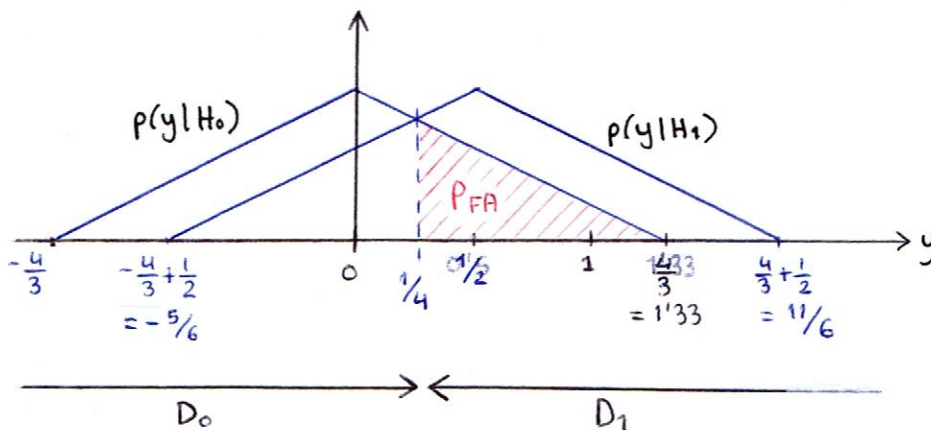
$$= \delta(y - 0.5) * p(y) * p(y)$$

Suma de 3 V.A.
la 1ª de ellas
es una constante
no aleatoria
(su $p(y)$ puede
verse una delta
en el valor
de la cte)

Aplicando:

$$\mathcal{L}(z) = \frac{p(y|H_1)}{p(y|H_0)} \stackrel{H_1}{\geq} \lambda = \frac{(C_{10} - C_{00}) P(H_0)}{(C_{01} - C_{11}) P(H_1)} = 1$$

$$ML \leftrightarrow p(y|H_1) \stackrel{H_1}{\geq} p(y|H_0)$$



Problema 5.

$$H_0 : f_x(x) = \frac{1}{2} e^{-|x|}$$

$$H_1 : f_x(x) = N(0, 4) = \frac{1}{\sqrt{2\pi} \cdot 2} e^{-\frac{x^2}{8}}$$

Determinar la LRT y el umbral

$$\Lambda(z) \stackrel{H_1}{\geq} \lambda$$

sin información de los costes ni de las $p(H_0)$, $p(H_1)$ asumiremos $\lambda = 1$

$$\Lambda(x) = \frac{p(x|H_1)}{p(x|H_0)} = \frac{\frac{1}{\sqrt{2\pi} \cdot 2} e^{-\frac{x^2}{8}}}{\frac{1}{2} e^{-|x|}} = \frac{1}{\sqrt{2\pi}} \cdot e^{-\frac{x^2}{8} + |x|} \stackrel{H_1}{\geq} 1$$

$$\ln \Lambda(x) = -\ln \sqrt{2\pi} - \frac{x^2}{8} + |x| \stackrel{H_1}{\geq} 0$$

$$\boxed{|x| - \frac{x^2}{8} \stackrel{H_1}{\geq} \ln \sqrt{2\pi}} = \frac{1}{2} \ln 2\pi$$

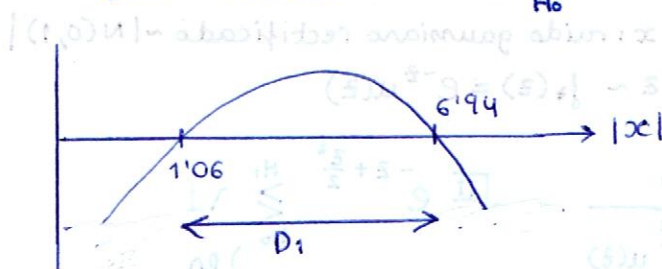
Para obtener los umbrales habrá que resolver la ecuación cuadrática

$$-|x|^2 + 8|x| - 4 \ln 2\pi \geq 0$$

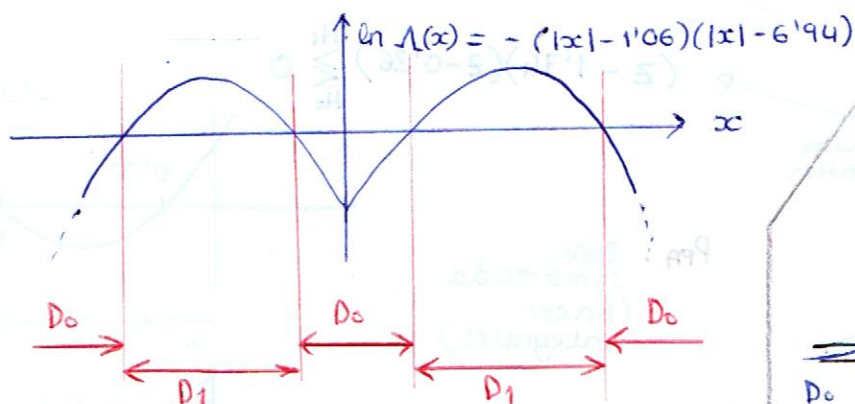
$$[a = -1 \quad b = 8 \quad c = -4 \ln 2\pi] \quad x = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

$$\text{soluciones} \quad x = \begin{cases} 1'06 \\ 6'94 \end{cases}$$

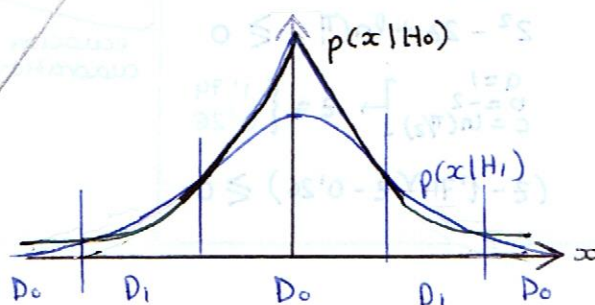
$$-(|x| - 1'06)(|x| - 6'94) \stackrel{H_1}{\geq} 0$$



Nota: Cuidado. Las regiones en x serán más!



Gráficamente tiene sentido:

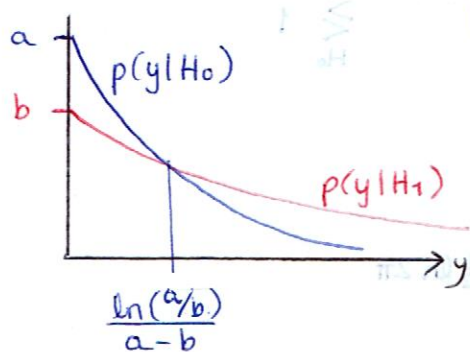


Problema 6 : $H_0 : p(y) = a e^{-ay}$ $a > b > 0$
 $H_1 : p(y) = b e^{-by}$

en el problema original habia que deducir esta constante $\Rightarrow \int_0^{\infty} k e^{-ay} dy = \left[-\frac{k}{a} e^{-ay} \right]_0^{\infty} = \frac{k}{a} = 1$

a) Criterio ML

$$\Lambda(y) = \frac{p(y|H_1)}{p(y|H_0)} = \frac{b e^{-by}}{a e^{-ay}} = \frac{b}{a} e^{+(a-b)y} \underset{H_0}{\overset{H_1}{\geq}} \lambda = 1$$



$$\ln\left(\frac{b}{a}\right) + (a-b)y \underset{H_0}{\overset{H_1}{\geq}} 0$$

$y \underset{H_0}{\overset{H_1}{\geq}} \frac{\ln(a/b)}{a-b}$

b) Probabilidad de falsa alarma

$$P_{FA} = Pr(D_1, H_0) = \int_{D_1} p(y|H_0) dy = \int_{\frac{\ln(a/b)}{a-b}}^{\infty} e^{-ay} dy = \left[-e^{-ay} \right]_{\frac{\ln(a/b)}{a-b}}^{\infty} = e^{-\frac{a}{a-b} \ln(a/b)} = e^{-\frac{a}{a-b} \ln(a/b)} = e^{\ln\left[\left(\frac{a}{b}\right)^{-\frac{a}{a-b}}\right]} = \left(\frac{b}{a}\right)^{\frac{a}{a-b}}$$

Problema 7. $H_0 : y = x$ x : ruido gaussiano rectificado $\sim |N(0,1)|$
 $H_1 : y = z$ $z \sim f_z(z) = e^{-z} u(z)$

$$\Lambda(z) = \frac{p(y|H_1)}{p(y|H_0)} = \frac{e^{-z} u(z)}{\frac{2}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} u(z)} = \sqrt{\frac{\pi}{2}} e^{-z + \frac{z^2}{2}} \underset{H_0}{\overset{H_1}{\geq}} 1$$

gaussiana rectificada

$$\frac{1}{2} \ln\left(\frac{\pi}{2}\right) - z + \frac{z^2}{2} \geq 0$$

Hay que resolver
 $z^2 - 2z + \ln(\pi/2) \geq 0$
 $\left. \begin{matrix} a=1 \\ b=-2 \\ c=\ln(\pi/2) \end{matrix} \right\} \rightarrow z = \begin{cases} 1.74 \\ 0.26 \end{cases}$
 $(z - 1.74)(z - 0.26) \geq 0$

ecuacion cuadratica

P_{FA} : zona sombreada (hacer integrales)

