

$$\hat{f}_{MAP}(\underline{x}) = \underset{\underline{d}}{\operatorname{argmax}} \frac{P_{\underline{x}|\underline{D}}(\underline{x}|\underline{d}) P_{\underline{D}}(\underline{d})}{P_{\underline{x}}(\underline{x})}$$

R.19

$$= \underset{\underline{d}}{\operatorname{argmax}} \underbrace{P_{\underline{x}|\underline{D}}(\underline{x}|\underline{d})}_{\text{likelihood}} \underbrace{P_{\underline{D}}(\underline{d})}_{\text{prior}} \quad \text{because } P_{\underline{x}}(\underline{x}) \text{ does not depend on } \underline{d}.$$

or Taking the log (an increasing function does not alter the order relationship)

$$\hat{f}_{MAP}(\underline{x}) = \underset{\underline{d}}{\operatorname{argmax}} \left[\underbrace{\log P_{\underline{x}|\underline{D}}(\underline{x}|\underline{d})}_{\text{log-likelihood}} + \log P_{\underline{D}}(\underline{d}) \right]$$

The log transformation may be useful when the likelihood is in the exponential family of distributions (i.e. Gaussian).

ML ESTIMATION (KV MASSHA VEROSITIALIANZA)

Note that the MAP estimator requires knowledge of both the prior $P_{\underline{D}}(\underline{d})$ and the likelihood $P_{\underline{x}|\underline{D}}(\underline{x}|\underline{d})$.

The prior $P_{\underline{D}}(\underline{d})$ represents our marginal knowledge of the hidden variable. When $P_{\underline{D}}(\underline{d})$ is not available, or it's equivalently assumed to be uniform the estimator is

$$\hat{f}_{ML}(\underline{x}) = \underset{\underline{d}}{\operatorname{argmax}} P_{\underline{x}|\underline{D}}(\underline{x}|\underline{d})$$

MAXIMUM
LIKELIHOOD
ESTIMATOR

or taking the log

R.20

$$\hat{f}^{ML}(\underline{x}) = \underset{\underline{d}}{\operatorname{argmax}} \underbrace{\log P(\underline{x}|\underline{d})}_{\text{log-likelihood}}$$

In solving regression problems (cont. D) ML is very popular because seldom we have knowledge of the prior on D.

The argmax of $P_{\underline{x}|\underline{d}}(\underline{x}|\underline{d})$ could be found solving $\nabla_{\underline{d}} P_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) = \underline{0}$

likelihood equation

or

$$\nabla_{\underline{d}} \log P_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) = \underline{0}$$

This condition must be verified to give a maximum value, that in general could be multiple.

Let us look now at some simple but typical examples of ML estimators

Example 1

$$X = \underset{\substack{\uparrow \\ \text{unknown} \\ \text{constant}}}{d} + W \leftarrow \begin{matrix} \text{Additive} \\ \text{Gaussian} \\ \text{noise} \end{matrix}$$

$$P_W(W) = \mathcal{N}(W; 0, \sigma_W^2)$$

$$P_{X|d}(x|d) = P_W(x-d) = \frac{1}{\sqrt{2\pi\sigma_W^2}} e^{-\frac{(x-d)^2}{2\sigma_W^2}}$$

$$\hat{d}^{ML} = \underset{d}{\operatorname{argmax}} P_{X|d}(x|d) = x$$

Trivial result:
take the observation
as it is!!

EXAMPLE 2

$$\underline{x} = d \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix} + \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_N \end{bmatrix}$$

This is an "unknown" constant plus noise

Assume all w_i to be independent and distributed according to N densities $\{p_{w_i}(w_i), i=1, \dots, N\}$. Therefore the components of $\underline{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_N \end{bmatrix}$ are independent, distributed according to the same densities of w_i but with mean d .

$$P_{\underline{x}|d}(\underline{x}|d) = \prod_{i=1}^N p_{w_i}(x_i - d)$$

Consider now various noise distributions.

2.1 GAUSSIAN NOISE

$$p_{w_i}(w_i) = \mathcal{N}(w_i, 0, \sigma_w^2)$$

$$P_{\underline{x}|d}(\underline{x}|d) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma_w^2}} e^{-\frac{(x_i - d)^2}{2\sigma_w^2}}, \text{ Taking the log}$$

$$\ln P_{\underline{x}|d}(\underline{x}|d) = \sum_{i=1}^N \left(\frac{1}{\sqrt{2\pi\sigma_w^2}} - \frac{(x_i - d)^2}{2\sigma_w^2} \right)$$

constant with respect to d

$$d^{ML} = \underset{d}{\operatorname{argmax}} \ln P_{\underline{x}|d}(\underline{x}|d) = \underset{d}{\operatorname{argmax}} - \sum_{i=1}^N \frac{(x_i - d)^2}{2\sigma_w^2}$$

$$= \underset{d}{\operatorname{argmin}} \sum_{i=1}^N (x_i - d)^2$$

$$\text{Taking the derivative w.r.t. } d \quad \& \sum_{i=1}^N (x_i - d)(-1) = 0$$

$$\sum_{i=1}^N x_i - Nd = 0$$

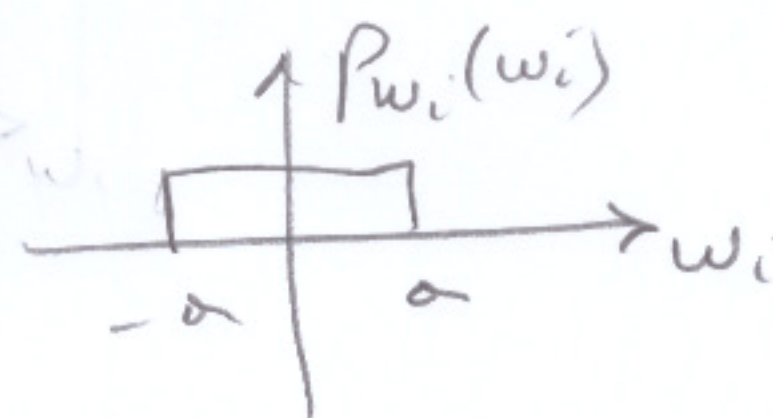
$$d^{ML} = \frac{1}{N} \sum_{i=1}^N x_i$$

Sample mean.

2.2 UNIFORM NOISE

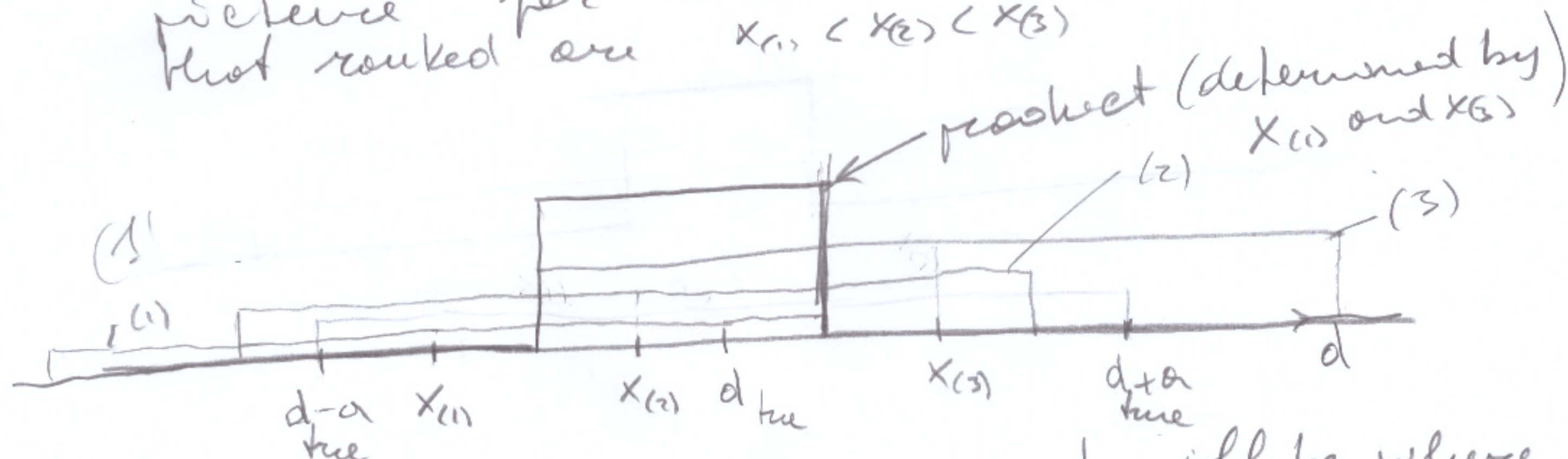
R, 22

$$P_{w_i}(w_i) = U(-a, a) = \frac{1}{2a} \pi\left(\frac{w_i}{2a}\right)$$



$$P_{x|d}(x|d) = \prod_{i=1}^N \pi\left(\frac{x_i - d}{2a}\right)$$

Let us visualize the product with a picture for $N=3$. Assume we get x_1, x_2 and x_3 that ranked are $x_{(1)} < x_{(2)} < x_{(3)}$



In general the non-zero segment will be where all the functions are superimposed. But this is determined by the min $x_{(1)}$ and the max $x_{(n)}$

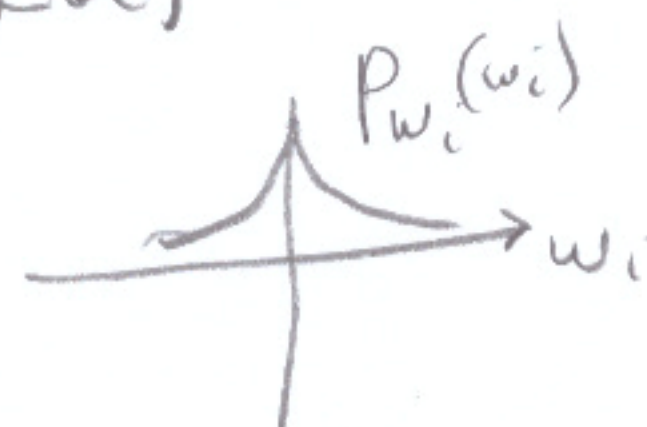
$$\hat{d} = \frac{x_{(1)} + x_{(n)}}{2}$$

MIDRANGE OR MIN-MAX ESTIMATOR

2.3 LAPLACIAN NOISE (Double exponential)

R.23

$$p_{w_i}(w_i) = \frac{\lambda}{2} e^{-\lambda |w_i|}$$



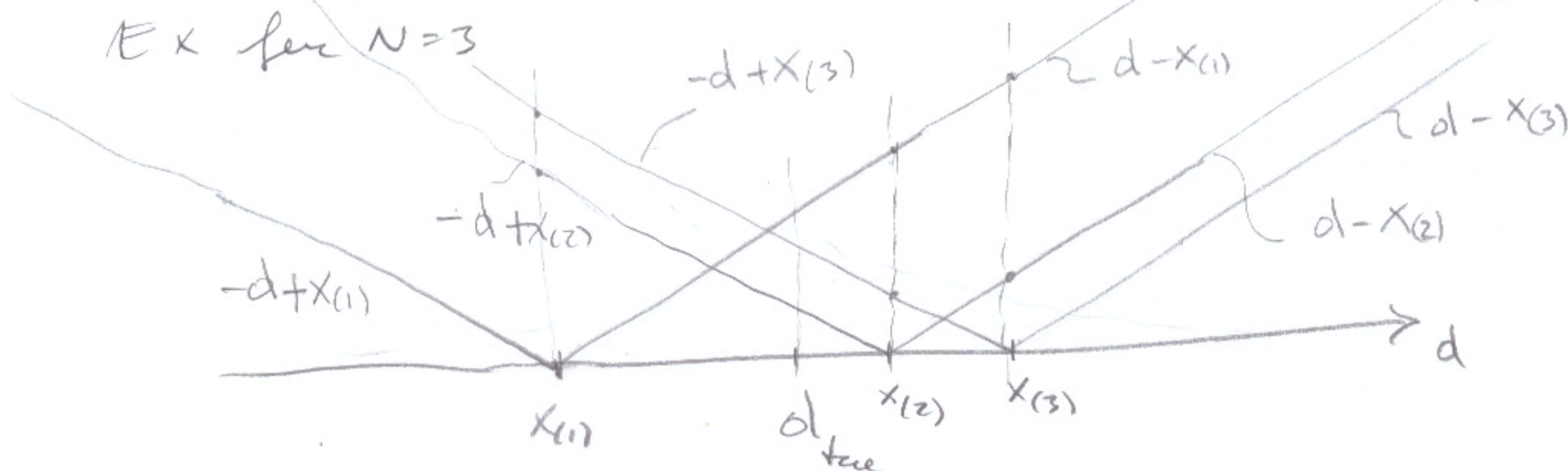
$$P_{\underline{x}|d}(\underline{x}|d) = \prod_{i=1}^N p_{w_i}(x_i - d) = \prod_{i=1}^N \frac{\lambda}{2} e^{-\lambda |x_i - d|}$$

$$\ln P_{\underline{x}|d}(\underline{x}|d) = \sum_{i=1}^N \left(\log \frac{\lambda}{2} - \lambda |x_i - d| \right)$$

$$d^{ML} = \underset{d}{\operatorname{argmax}} -\lambda \sum_{i=1}^N |x_i - d| = \underset{d}{\operatorname{argmin}} \sum_{i=1}^N |x_i - d|$$

MEAN ABSOLUTE DEVIATION

Ex for $N=3$



Analyzing the four intervals

- (a) $d < x_{(1)}$ $-d + x_{(1)} - d + x_{(2)} - d + x_{(3)} = -3d + x_{(1)} + x_{(2)} + x_{(3)}$
- (b) $x_{(1)} < d < x_{(2)}$ $d - x_{(1)} - d + x_{(2)} - d + x_{(3)} = -d - x_{(1)} + x_{(2)} + x_{(3)}$
- (c) $x_{(2)} < d < x_{(3)}$ $d - x_{(1)} + d - x_{(2)} - d + x_{(3)} = -d - x_{(1)} - x_{(2)} + x_{(3)}$
- (d) $d > x_{(3)}$ $d - x_{(1)} + d - x_{(2)} + d - x_{(3)} = 3d - x_{(1)} - x_{(2)} - x_{(3)}$

- (a) $d_{\min} = x_{(1)}$
- (b) $d_{\min} = x_{(2)}$
- (c) $d_{\min} = x_{(2)}$
- (d) $d_{\min} = x_{(3)}$

$$d^{ML} = x_{(2)}$$

Sample
Median
estimator

The solution could be found using the gradient

$$\nabla_{\underline{d}} \sum_{i=1}^N (x_i - g_i(\underline{d}))^2 = \underline{0}$$

$$-2 \sum_{i=1}^N (x_i - g_i(\underline{d})) \nabla_{\underline{d}} g_i(\underline{d}) = \underline{0}$$

When this equation cannot be solved, we could use Newton method using an iteration

$$\underline{d}(k) = \underline{d}(k-1) - \mu \nabla_{\underline{d}(k-1)}$$

$$\underline{d}(k) = \underline{d}(k-1) + \mu \sum_{i=1}^N (x_i - g_i(\underline{d}(k-1))) \nabla_{\underline{d}(k-1)} g_i(\underline{d}(k-1))$$

Given $x_1, x_2, \dots, x_N \xrightarrow{\text{Newton search}} \underline{d}^{MC}$

Example 4 GENERAL GAUSSIAN LIKELIHOOD R.26

$$p(\underline{x}|\underline{d}) = \mathcal{N}(\underline{x}; \underline{\mu}_{\underline{x}|\underline{d}}, \underline{\Sigma}_{\underline{x}|\underline{d}})$$

$$= \frac{1}{(2\pi)^{N/2} |\underline{\Sigma}_{\underline{x}|\underline{d}}|^{1/2}} e^{-\frac{1}{2}(\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}})^T \underline{\Sigma}_{\underline{x}|\underline{d}}^{-1} (\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}})}$$

$$\underline{d}^{ML} = \underset{\underline{d}}{\operatorname{argmax}} p(\underline{x}|\underline{d}) = \underset{\underline{d}}{\operatorname{argmax}} \ln p(\underline{x}|\underline{d})$$

$$= \underset{\underline{d}}{\operatorname{argmax}} \left[-\ln(2\pi)^{N/2} |\underline{\Sigma}_{\underline{x}|\underline{d}}|^{1/2} - \frac{1}{2}(\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}})^T \underline{\Sigma}_{\underline{x}|\underline{d}}^{-1} (\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}}) \right]$$

If the conditional covariance $\underline{\Sigma}_{\underline{x}|\underline{d}}$ does not depend on $\underline{d}$ and is $\underline{\Sigma}_x$

$$\underline{d}^{ML} = \underset{\underline{d}}{\operatorname{argmin}} (\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}})^T \underline{\Sigma}_x^{-1} (\underline{x} - \underline{\mu}_{\underline{x}|\underline{d}})$$

etc.

UNBIASED ESTIMATORS

An estimator of $\underline{D}$ is unbiased, if its mean matches the mean of the desired variable.

$$E[f(x)] = E[D]$$

When the variable is a constant, but unknown, $f(x)$ is unbiased if $E[f(x)] = \underline{d}$ (the estimator's mean matches the true value)

MSE LOWER BOUND FOR UNBIASED ESTIMATORS

There exist a quite striking result that predicts a lower bound for ^{the mean squared error} any unbiased estimator. This may be quite useful because, even if we do not have found the estimator, we may be able to predict a theoretical limit for the MSE. Consider first the scalar case.

THEOREM (CRAMER-RAO LOWER BOUND):

If $f(x)$ is an unbiased estimate of d .

$$E[(f(x) - d)^2] \geq \frac{1}{E\left[\left(\frac{\partial \ln p_{x|d}(x|d)}{\partial d}\right)^2\right]}$$

or equivalently

$$E[(f(x) - d)^2] \geq - \frac{1}{E\left[\frac{\partial^2 \ln p_{x|d}(x|d)}{\partial d^2}\right]}$$

with the following conditions satisfied

CRLB.2⁴²

(c) $\frac{\partial p_{\underline{x}|d}(\underline{x}|d)}{\partial d}$ and $\frac{\partial^2 p_{\underline{x}|d}(\underline{x}|d)}{\partial d^2}$

exist and are absolutely integrable.

Proof: Since $f(\underline{x})$ is unbiased, we have

$$E[f(\underline{x}) - d] = \int_{\underline{x}} (f(\underline{x}) - d) p_{\underline{x}|d}(\underline{x}|d) d\underline{x} = 0$$

Differentiating with respect to d , we have

$$\frac{\partial}{\partial d} \int_{\underline{x}} (f(\underline{x}) - d) p_{\underline{x}|d}(\underline{x}|d) d\underline{x} = 0$$

Integrability is ensured by the hypothesis of the theorem, therefore we can write

$$\int_{\underline{x}} \frac{\partial}{\partial d} [(f(\underline{x}) - d) p_{\underline{x}|d}(\underline{x}|d)] d\underline{x} = 0$$

$$-\int_{\underline{x}} p_{\underline{x}|d}(\underline{x}|d) d\underline{x} + \int_{\underline{x}} (f(\underline{x}) - d) \frac{\partial}{\partial d} p_{\underline{x}|d}(\underline{x}|d) d\underline{x} = 0$$

Now observe that

$$\frac{\partial \ln p_{\underline{x}|d}(\underline{x}|d)}{\partial d} = \frac{1}{p_{\underline{x}|d}(\underline{x}|d)} \cdot \frac{\partial p_{\underline{x}|d}(\underline{x}|d)}{\partial d}$$

Extracting $\frac{\partial p_{\underline{x}|d}}{\partial d}$ and substituting it into the above equation we get:

$$\int_X \frac{\partial \ln p_{x|d}(x|d)}{\partial d} p_{x|d}(x|d) (f(x) - d) dx = 1$$

We can rewrite it as

$$\int_X \left[\frac{\partial \ln p_{x|d}(x|d)}{\partial d} \sqrt{p_{x|d}(x|d)} \right] \left[\sqrt{p_{x|d}(x|d)} (f(x) - d) \right] dx = 1$$

Using Schwarz's inequality

$$\left| \int f(x) g(x) dx \right|^2 \leq \int f^2(x) dx \cdot \int g^2(x) dx$$

$$1 \leq \int_X \left[\frac{\partial \ln p_{x|d}(x|d)}{\partial d} \right]^2 p_{x|d}(x|d) dx \cdot \int_X (f(x) - d)^2 p_{x|d}(x|d) dx$$

Therefore

$$E[(f(x) - d)^2] \leq \frac{1}{E\left[\left(\frac{\partial \ln p_{x|d}(x|d)}{\partial d}\right)^2\right]}$$

The equality in Schwarz inequality holds iff

$f(x) = K g(x)$ for some K . In our case

the equality holds iff.

$$K (f(x) - d) = \frac{\partial \ln p_{x|d}(x|d)}{\partial d} \text{ for some } K.$$

To prove from (6), observe that

CRLB.4

15

$$\int_x p_{x|d}(x|d) dx = 1$$

Differentiating with respect to d , we have

$$\int_x \frac{\partial p_{x|d}(x|d)}{\partial d} dx = 0$$

or

$$\int_x \frac{\partial \ln p_{x|d}(x|d)}{\partial d} p_{x|d}(x|d) dx = 0$$

Differentiating again with respect to d , we have

$$\int_x \frac{\partial^2 \ln p_{x|d}(x|d)}{\partial d^2} p_{x|d}(x|d) dx + \int_x \left(\frac{\partial \ln p_{x|d}(x|d)}{\partial d} \right)^2 p_{x|d}(x|d) dx = 0,$$

which can be written as

$$E \left[\left(\frac{\partial \ln p_{x|d}(x|d)}{\partial d} \right)^2 \right] = -E \left[\frac{\partial^2 \ln p_{x|d}(x|d)}{\partial d^2} \right].$$

The above expression immediately gives us (6) $\square$

Note that only if the condition

$$K(d)(f(x) - d) = \frac{\partial \ln p_{x|d}(x|d)}{\partial d} \quad (*)$$

is satisfied, the lower bound is achieved. Note that the constant K can depend on d and that in general this may be a rather restrictive condition. Nevertheless if the condition holds and $\frac{\partial \ln p_{x|d}(x|d)}{\partial d} = 0$,

which corresponds to a unique solution for the CRLB.5
ML estimator, the ML estimator reaches the lower
bound. It is rare for general cases that estimator
reach the lower bound.

EXAMPLE case 1

Take example 2 with W gaussian. Namely

$$\underline{X} = d \begin{bmatrix} 1 \\ 1 \end{bmatrix} + \underline{W}$$

the ML equation reads

$$\frac{\partial}{\partial d} \ln p_{\underline{x}|d}(\underline{x}|d) = \frac{1}{\sigma_w^2} \left(\sum_{i=1}^N x_i - d \right) = 0$$

or

$$\frac{\partial}{\partial d} \ln p_{\underline{x}|d}(\underline{x}|d) = \frac{N}{\sigma_w^2} \left(\frac{1}{N} \sum_{i=1}^N x_i - d \right) = 0$$

the first part of the equation satisfies the
proportionality condition^(*), therefore the
bound is achieved. the ^{ML} estimator, (sample mean)

$$\hat{d}_{ML} = \frac{1}{N} \sum_{i=1}^N x_i,$$

gives

$$E[(f(\underline{x}) - d)^2] = \frac{\sigma_w^2}{N}$$

since

$$\frac{\partial^2}{\partial d^2} \ln p_{\underline{x}|d}(\underline{x}|d) = -\frac{N}{\sigma_w^2}$$

and because it is unbiased:

$$E\left[\frac{1}{N} \sum_{i=1}^N x_i\right] = \frac{1}{N} \sum_{i=1}^N E[d + w_i] = d$$

THEOREM Consider any unbiased estimate $\hat{d}$ of d .

then
$$E[(f_i(x) - d_i)^2] \geq J^{ii} \quad \forall i = 1, \dots, M$$

where J^{ii} the i th diagonal element of the $L \times L$ square matrix $\underline{J} = \underline{F}^{-1}$

where

$$[F]_{ij} = E \left[\frac{\partial \ln p_{x|d}(x|d)}{\partial d_i} \cdot \frac{\partial \ln p_{x|d}(x|d)}{\partial d_j} \right]_{i,j=1, \dots, M}$$

which by the property proved above (*) it can be written as

$$[F]_{ij} = -E \left[\frac{\partial^2 \ln p_{x|d}(x|d)}{\partial d_i \partial d_j} \right]_{i,j}$$

The matrix $\underline{F}$ is called FISHER INFORMATION MATRIX. In fact it measures the sensitivity of our observations (in probability) with respect to variation in the observed estimated quantity.

In matrix notation, the Fisher's information matrix can also be expressed as

(*) To prove this equivalence we can just differentiate twice, once with respect to d_i and also with respect to d_j the equation $\int_x p_{x|d}(x|d) dx = 1$, and by noting that $\frac{\partial \ln p_{x|d}(x|d)}{\partial d_i} = \frac{1}{p_{x|d}(x|d)} \frac{\partial p_{x|d}(x|d)}{\partial d_i}$

$$\underline{F} = E \left[\left(\nabla_{\underline{d}} \ln p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) \right) \left(\nabla_{\underline{d}} \ln p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) \right)^T \right]$$

$$= - E \left[\left(\nabla_{\underline{d}} \nabla_{\underline{d}}^T \right) \ln p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) \right]$$

ERLB.748

The bound equality holds only if

$$(f_i(\underline{x}) - d_i) = \sum_{j=1}^L k_{ij}(\underline{d}) \frac{\partial \ln p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d})}{\partial d_j},$$

for all $i=1, \dots, M$, and $\forall \underline{x} \in \mathcal{X}$.

PROOF. The proof follows the same guidelines of that for the one-dimensional case: Since the estimator is assumed to be unbiased, we have

$$\int_{\mathcal{X}} (f_i(\underline{x}) - d_i) p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) d\underline{x} = 0, \quad \forall i=1, \dots, M,$$

or

$$\int_{\mathcal{X}} f_i(\underline{x}) p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) d\underline{x} = d_i, \quad \forall i=1, \dots, M$$

Differentiating both sides with respect to d_j

$$\int_{\mathcal{X}} f_i(\underline{x}) \frac{\partial p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d})}{\partial d_j} d\underline{x} = \int_{\mathcal{X}} f_i(\underline{x}) \frac{\partial \ln p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d})}{\partial d_j} p_{\underline{x}|\underline{d}}(\underline{x}|\underline{d}) d\underline{x}$$

$$= \delta_{ij}, \quad \forall j=1, \dots, M.$$

δ_{ij} is the Kronecker delta: $\delta_{ij} \triangleq \begin{cases} 1 & i=j \\ 0 & \text{else.} \end{cases}$

let us prove the result for $i=1$. The same CRLB.8
 procedure would follow for any $i=1, \dots, M$
 let us define the $(M+1)$ -dimensional vector

$$\underline{\alpha} = \begin{bmatrix} f_1(x) - d_1 \\ \frac{\partial \ln p_{x|d}(x|d)}{\partial d_1} \\ \vdots \\ \frac{\partial \ln p_{x|d}(x|d)}{\partial d_L} \end{bmatrix}$$

The correlation matrix of $\underline{\alpha}$ is ;

$$E[\underline{\alpha}\underline{\alpha}^T] = \begin{bmatrix} \sigma_{\epsilon_1}^2 & 1 & 0 & \dots & 0 \\ 1 & & & & \\ 0 & & \underline{F} & & \\ 0 & & & & \end{bmatrix}$$

The ones come from the above expression, and $\sigma_{\epsilon_1}^2$ is just the mean square error on the first variable. The matrix is non negative definite since it is a covariance matrix. Evaluating the determinant using the cofactor expansion we get

$$\sigma_{\epsilon_1}^2 |\underline{F}| - \text{cofactor } F_{11} \geq 0.$$

Therefore if $\underline{F}$ is non singular

$$\sigma_{\epsilon_1}^2 \geq \frac{\text{cofactor } F_{11}}{|F|} = J_{11}$$

CRLB.9⁵⁰

where J_{11} is the first element of F^{-1} .

The case when F is singular has to be worked out for every specific problem.

The order for the determinant to be equal to zero, the term $f_1(x) - d_1$ must be expressible as a linear combination of the other terms in x . This is the condition for equality.

Exer

Consider the unknown random variable

$$X = \frac{1}{2} \ln \frac{W}{1+W}$$

known that W has a uniform distribution of W with probability density function

$f_W(w) = \frac{1}{2} \exp(-w) \quad w \geq 0$

RANDOM PARAMETER $\underline{D}$

CRLB.10

Clearly if some a-priori information is given about $\underline{D}$, it should be included in the bound. The condition on an estimator to be unbiased if $\underline{D}$ is a random L -dimensional variable is

$$E[\underline{f}(\underline{x})] = E[\underline{D}]$$

or

$$E[\underline{f}(\underline{x}) - \underline{D}] = \int_{\underline{x} \times \underline{D}} (\underline{f}(\underline{x}) - \underline{d}) p_{\underline{x}|\underline{D}}(\underline{x}|\underline{d}) p_{\underline{D}}(\underline{d}) d\underline{x} d\underline{d}$$

or

$$\int_{\underline{x} \times \underline{D}} \underline{f}(\underline{x}) p_{\underline{x}|\underline{D}}(\underline{x}|\underline{d}) p_{\underline{D}}(\underline{d}) d\underline{x} d\underline{d} = \int_{\underline{D}} \underline{d} p_{\underline{D}}(\underline{d}) d\underline{d}.$$

Rewriting the equation component-by-component

$$\int_{\underline{x} \times \underline{D}} f_i(\underline{x}) p_{\underline{x}|\underline{D}}(\underline{x}|\underline{d}) p_{\underline{D}}(\underline{d}) d\underline{x} d\underline{d} = \int_{\underline{D}} d_i p_{\underline{D}}(\underline{d}) d\underline{d} \quad \forall i=1, \dots, L.$$

Differentiating both sides with respect to d_j , with the usual assumptions that allow to take the derivatives inside the integration signs, we have

$$\int_{\underline{x} \times \underline{D}} f_i(\underline{x}) \left[\frac{\partial p_{\underline{x}|\underline{D}}(\underline{x}|\underline{d})}{\partial d_j} p_{\underline{D}}(\underline{d}) + p_{\underline{x}|\underline{D}}(\underline{x}|\underline{d}) \frac{\partial p_{\underline{D}}(\underline{d})}{\partial d_j} \right] d\underline{x} d\underline{d} =$$

$$= \int_D \left(\delta_{ij} p_{\underline{D}}(\underline{d}) + \underline{d} \frac{\partial p_{\underline{D}}(\underline{d})}{\partial d_j} \right) d \underline{d}$$

...

carrying out all the derivations, in a way similar to the procedure used for the case of $\underline{D}$ nonrandom we get that the same result holds, but the Fisher's information matrix has to be re-defined as

$$\underline{F} = \underline{F}_1 + \underline{F}_2,$$

$\underline{F}_1$ has the usual definition given above, and

$$[\underline{F}_2]_{ij} = E \left[\frac{\partial \ln p_{\underline{D}}(\underline{d})}{\partial d_i} \frac{\partial \ln p_{\underline{D}}(\underline{d})}{\partial d_j} \right]$$

or

$$[\underline{F}_2]_{ij} = -E \left[\frac{\partial^2 \ln p_{\underline{D}}(\underline{d})}{\partial d_i \partial d_j} \right]$$

$$\forall i, j = 1, \dots, M.$$

(see Van Trees for more details)