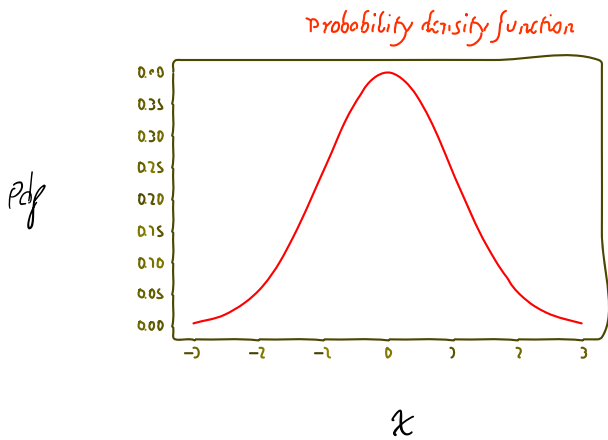


1. Basic Principles

CONTINUOUS RANDOM VARIABLE

Characterized by pdf $f_X(x; \theta)$



- Probability to be inside an interval

$$P(x \leq X \leq x+dx) = f_X(x; \theta) dx$$

- Expected Value

$$E[X] = \int x f_X(x; \theta) dx$$

$$\bullet \text{Var}[X] = E[(X - E[X])^2]$$

$$\bullet \text{cdf: } F(t, \theta) = P(X \leq t; \theta)$$

Population Parameters

Expected value

$$E[X] = \mu \quad \bullet \quad E[M] = E[X]$$

Variance

$$\text{var}[X] = V[X] = \sigma^2 = E[(X - \mu)^2] \quad \bullet \quad V[X] = E[X^2] - (E[X])^2$$

Covariance

$$\text{cov}(X, Y) = \sigma_{xy} = E[(X - \mu_x)(Y - \mu_y)]$$

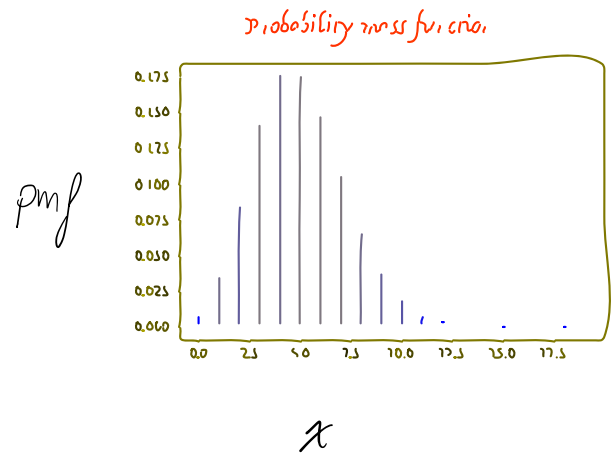
$$\bullet \quad V[M] = \frac{V[X]}{n}$$

Moments

$$E[X^n] = m_n$$

DISCRETE RANDOM VARIABLE

Characterized by pmf $P_X(x, \theta)$



- $P_X(x, \theta)$ = prob. that $X = x$

$$\bullet E[X] = \sum_i x_i P_X(x_i; \theta)$$

Independence

X and Y are independent if

$$f_{XY}(x,y) = f_X(x) f_Y(y)$$

Remark: Dependence $\neq$ Correlation

Correlation is measured by $\text{cor}(X,Y) = E[XY] - E[X]E[Y]$

$\leadsto \text{cov}(X,Y) = 0$ if uncorrelated $\leadsto$ Independence $\Rightarrow$ Uncorrelation
BUT NOT CONTRARY

EMPIRICAL QUANTITIES

Sample mean

$$M = \frac{1}{n} \sum_{i=1}^n X_i$$

Sample median: assuming $X_1 < X_2 < \dots < X_n$

- median = $\frac{X_{n/2} + X_{1+n/2}}{2}$ (n even)
- median = $X_{(n+1)/2}$ (n odd)

Sample variance

$$S_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - M)^2 = \frac{1}{n} \sum_{i=1}^n X_i^2 - M^2$$

biased

$$S_{n-1}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - M)^2$$

unbiased

Sample covariance matrix

$$Q = \frac{1}{n-1} \sum_{i=1}^n (\vec{X}_i - \vec{M})(\vec{X}_i - \vec{M})^T$$

FUNCTIONS OF RANDOM VARIABLES

Functions of random variables are also random variables

$$Z = g(\underbrace{X_1, \dots, X_n}_{\text{sample}})$$

$\nearrow$ empirical quantiles
 $\downarrow$ function

The mean is

$$E[Z] \approx g(E[X_1], \dots, E[X_n]) + \frac{1}{2} \sum_{i=1}^n \frac{\partial^2 g}{\partial X_i^2} \text{var}[X_i] + \sum_{i,j} \frac{\partial^2 g}{\partial X_i \partial X_j} \text{cov}(X_i, X_j)$$

Variance

$$\text{var}[Z] \approx \sum_{i=1}^n \sum_{j=1}^n \frac{\partial g}{\partial X_i} \frac{\partial g}{\partial X_j} \text{cov}(X_i, X_j)$$

$$\bullet \text{ If } X_i \text{ are iid} \Rightarrow \frac{\partial g}{\partial X_i} = \frac{\partial g}{\partial X_j}$$

$$\Rightarrow \text{cov}(X_i, X_j) = \text{cov}(X_i, X_i) = \text{var}(X_i)$$

Binomial and Multinomial Distributions

Binomial

$$P(K; n, p) = \binom{n}{K} p^K (1-p)^{n-K}$$

$$= \frac{n!}{K!(n-K)!} p^K (1-p)^{n-K}$$

"Describe the probability of having K successes in n independent, identical trials where each trial has only 2 possible outcomes"

Properties

- $E[K] = np$
- $\text{var}[K] = np(1-p)$

if $n=1 \Rightarrow$ Bernoulli

Multinomial

$$P(n_1, \dots, n_m; n, p_1, \dots, p_m) = \frac{n!}{n_1! n_2! \dots n_m!} p_1^{n_1} \dots p_m^{n_m}$$

$\leadsto$ m : number of possible results in a trial
 n_i : number of results of type i ($i \in [1, m]$), $\sum_i n_i = n$
 p_i : probability that result in a trial is of type i

Properties

- $E[n_i] = np_i$
- $\text{var}[n_i] = np_i(1-p_i)$
- $\text{cov}(n_i, n_j) = -np_i p_j$

Example [\[edit\]](#)

Suppose that in a three-way election for a large country, candidate A received 20% of the votes, candidate B received 30% of the votes, and candidate C received 50% of the votes. If six voters are selected randomly, what is the probability that there will be exactly one supporter for candidate A, two supporters for candidate B and three supporters for candidate C in the sample?

Note: Since we're assuming that the voting population is large, it is reasonable and permissible to think of the probabilities as unchanging once a voter is selected for the sample. Technically speaking this is sampling without replacement, so the correct distribution is the [multivariate hypergeometric distribution](#), but the distributions converge as the population grows large in comparison to a fixed sample size^[1].

$$\Pr(A=1, B=2, C=3) = \frac{6!}{1!2!3!} (0.2^1)(0.3^2)(0.5^3) = 0.135$$

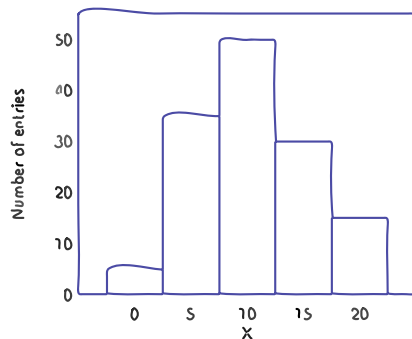
Filling an histogram from a FIXED sample is a MULTINOMIAL PROCESS N FIXED

$$\Rightarrow P(\text{histogram}) = \frac{n!}{\prod_i n_i!} \prod_i p_i^{n_i}$$

$\leadsto n_i$: content of bin i

p_i : probability to be in bin i

$$p_i = \int_{x \in \text{bin}} f_X(x) dx$$



N = # bins

$\leadsto$ Bins are not independent $\Rightarrow$ What if n is not fixed?

The probability changes as follows:

$$P(n_1, \dots, n_m; n, p_1, \dots, p_m) \rightarrow P(n, n_1, \dots, n_m; \nu, p_1, \dots, p_m)$$

What is $P(n, n_1, \dots, n_m)$?

$$\rightarrow P(n, n_1, \dots, n_m) = \underbrace{P(n_1, \dots, n_m; n)}_{\text{multinomial}} \times P(n)$$

↓
Parameters of the distribution of n

Often n is a Poisson random variable → See next slide

Poisson Distribution

$$P(n; \nu) = \frac{\nu^n}{n!} e^{-\nu}$$

"Probability of a given number of events occurring in a fixed interval of time or space if those events occur with a known constant mean rate"

Properties

- $E[n] = \nu$
- $\text{var}[n] = \nu$

Poisson = limit of binomial when $n \rightarrow \infty$ and $p \rightarrow 0$

Histogram Summary

n Fixed: histogram $\sim$ multinomial

$n \sim$ Poisson: histogram $\sim$ product of poisson in each bin

MULTIDIMENSIONAL SAMPLES

2D Case where N measurements are performed of M different variables

$$\vec{x}_i = \{\vec{x}_1, \vec{x}_2, \dots, \vec{x}_N\} \quad \text{with} \quad \begin{cases} \vec{x}_1 : x_1^{(1)}, x_1^{(2)}, \dots, x_1^{(M)} \\ \vdots \\ \vec{x}_N : x_N^{(1)}, x_N^{(2)}, \dots, x_N^{(M)} \end{cases}$$

Def

$$\text{cov}(x, y) = \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})(y_i - \bar{y}) = \overline{xy} - \bar{x} \bar{y}$$

Correlation Factor

$$f_{xy} = \frac{\text{cov}(x, y)}{\sigma_x \sigma_y} \quad \text{with} \quad -1 \leq f_{xy} \leq 1$$

- $f = 0 \Rightarrow x$ and y are uncorrelated ($\neq$ independent)
- $f_{xy} = \pm 1 \Rightarrow x$ and y are fully correlated

Covariance Matrix

$$C = \begin{pmatrix} \text{cov}(x_1, x_1) & \dots & \text{cov}(x_1, x_N) \\ \vdots & \text{cov}(x_i, x_j) & \vdots \\ \text{cov}(x_N, x_1) & \dots & \text{cov}(x_N, x_N) \end{pmatrix} = \begin{pmatrix} \sigma_1^2 & \dots & f_{1N} \sigma_1 \sigma_N \\ \vdots & f_{ij} \sigma_i \sigma_j & \vdots \\ f_{N1} \sigma_N \sigma_1 & \dots & \sigma_N^2 \end{pmatrix}$$

Correlation Matrix

$$f = \begin{pmatrix} 1 & \dots & f_{1N} \\ \vdots & 1 & \vdots \\ f_{N1} & \dots & 1 \end{pmatrix}$$

The error ellipse

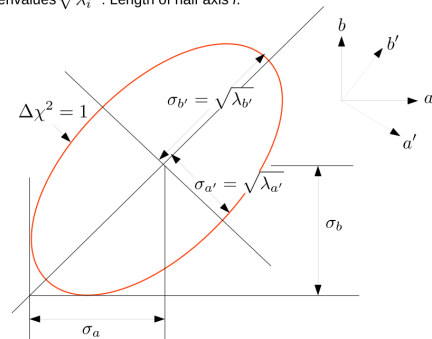
From the covariance matrix

we can find the eigenvalues λ_i

$\Rightarrow$ They are the variance of x_i
 $= \sigma_i^2$

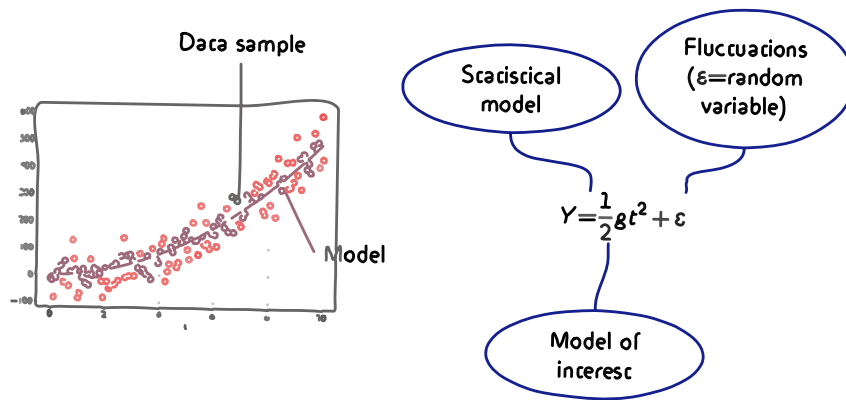
$$\theta? \rightarrow \theta = \frac{1}{2} \tan^{-1} \left(\frac{2 f_{12} \sigma_1 \sigma_2}{\sigma_1^2 - \sigma_2^2} \right)$$

Eigenvalues $\sqrt{\lambda_i}$: Length of half axis i .



2. STATISTICAL MODEL

It's the mathematical expression describing your observation



Statistical model For iid samples

Suppose you have a sample $X = (X_1, \dots, X_n)$

The statistical model can be written as

$$f_{X_1 \dots X_n}(x_1, \dots, x_n; \theta) = f_{X_1}(x_1; \theta) \times f_{X_2}(x_2; \theta) \times \dots \times f_{X_n}(x_n; \theta)$$

For IID variables

$$f_{X_1 \dots X_n}(x_1, \dots, x_n; \theta) = f_X(x_1; \theta) \times f_X(x_2; \theta) \times \dots \times f_X(x_n; \theta)$$

$$\Rightarrow f_{X_1 \dots X_n}(x_1, \dots, x_n; \theta) = \prod_{i=1}^n f_X(x_i; \theta)$$

~> IF X_i 's are normally distributed

$$\begin{aligned} \text{stat. mod} &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x_i - \mu)^2}{2\sigma^2}\right) \\ &= \left(\frac{1}{\sqrt{2\pi}\sigma}\right)^n \exp\left(-\sum_{i=1}^n \frac{(x_i - \mu)^2}{2\sigma^2}\right) \end{aligned}$$

Likelihoods

$$\mathcal{L}(\theta, \vec{x}) = \prod_{i=1}^N \mathcal{L}_i(\theta, \vec{x})$$

where $\mathcal{L}_i = f_i$

Extended Likelihood Function

If the size of the sample is not constant but Follows Poisson Distribution

$$n \sim \text{Pois}(\tau)$$

In such cases the likelihood function has to be extended

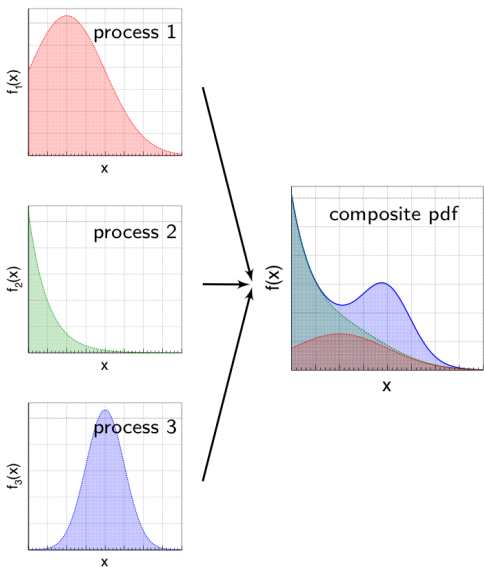
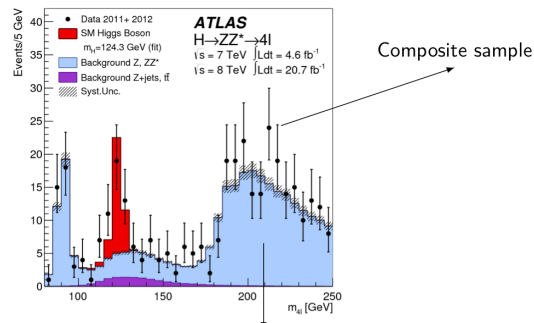
$$\mathcal{L}_{\text{ext}}(\theta, \vec{v}; x, n) = \frac{v^n}{n!} e^{-v} \times \mathcal{L}(\theta, \vec{x})$$

Composite Samples

In realistic cases, samples are often composite

→ Composite Sample = sample in which events can come from different origins

Example: $\text{sample} = (\underline{E_1}, \underline{E_2}, \dots, \underline{E_n})$
source bkg source



☐ **Composite pdf:**

$$f(\mathbf{x}; \{\mu_p\}, \theta) = \frac{\sum_{p=1}^P \mu_p f_p(\mathbf{x}; \theta)}{\sum_{p=1}^P \mu_p}$$

where:

- p : process index
- μ_p : expected number of elements in sample from process p
- $f_p(x; \theta)$: pdf for process p

□ **Remark:** the μ_p 's are often unknown

- Determining them can be one of the objective

General extended unblinded $\mathcal{L}$

General extended unbinned likelihood:

General extended unbinned likelihood:

$$\mathcal{L}_{\text{ext}}(\{\mu_p\}, \theta; \{x_i\}, n) = \underbrace{\frac{(\sum \mu_p)^n}{n!}}_{\text{extended term}} e^{-\sum \mu_p} \times \prod_{i=1}^n \underbrace{\frac{\sum_{p=1}^P \mu_p f_p(x_i; \theta)}{\sum_{p=1}^P \mu_p}}_{\text{composite pdf}}$$

Assuming the x_i 's are iid

Assuming the x_i 's are iid

Treatment of nuisance parameters

Suppose your model has one parameter of interest and other parameter you do not care about

Two cases must be distinguished:

- ① **Nuisance parameters perfectly known:** formalism presented previously works perfectly fine
- ② **Nuisance parameters not perfectly known:** formalism presented previously should be extended

→ New term in likelihood

$$\mathcal{L}_{\text{full}} = \mathcal{L}(\theta_1, \theta_2, \dots, \theta_N) \times \underbrace{g(\theta_1, \theta_2, \dots, \theta_N)}_{\text{New const. term}}$$

Detailed Example

Let's consider a Poissonian process with background (Measurement of radioactivity)

$$\text{Statist. Model: } \mathcal{L}(s, b) = \frac{(s+b)^n}{n!} e^{-(s+b)}$$

- n : number of measured events (observable)
- s : number of signal events (parameter of interest)
- b : number of background events (nuisance parameter)

You can get a value for $b \Rightarrow$ But + error : $b \pm \sigma_b$

Full likelihood:

• Incomplete Stat. Model: $\mathcal{L}(s, b) = \frac{(s+b)^n}{n!} e^{-(s+b)}$

• b determined in an auxiliary poissonian measurement $\mathcal{L}_{\text{aux}}(b) = \frac{(\alpha b)^{n_b}}{n_b!} e^{-(\alpha b)}$

• Full likelihood $\mathcal{L}_{\text{full}} = \mathcal{L}(s, b) \times \mathcal{L}_{\text{aux}}(b) = \frac{(s+b)^n}{n!} e^{-(s+b)} \times \frac{(\alpha b)^{n_b}}{n_b!} e^{-(\alpha b)}$

3 Inference Generalities

Are what allow us to gain, from a SAMPLE and a STATISTICAL MODEL, some knowledge about the parameter of interest

→ Knowledge acquisition process is called inference

Two big schools of thought:

FREQUENTIST

BAYESIAN

Different interpretation of probability

INTERPRETATION 1: Probability = Frequency of occurrence

$$P(A) = \lim_{n \rightarrow \infty} \frac{n(A)}{n}$$

where $\begin{cases} A: \text{some event} \\ n(A): \text{number of times event } A \text{ occurs} \\ n: \text{total number of events} \end{cases}$

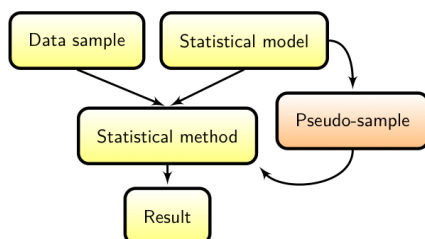
Example: dice roll

INTERPRETATION 2: Probability = Degree of belief on occurrence

Often used for unique events (for which frequentist calculation impossible)
→ Probability tomorrow will rain

Frequentist

- Use frequentist interpretation of probability
- Model parameters are fixed, the only thing that can have a probability associated to it is the data
$$P(\mathbf{x}; \theta)$$
- Inference based on notion of **pseudo-sample** or **pseudo-dataset**



Bayesian

- Use probabilities as degree of belief
- Model parameters have a probability
$$P(\theta; \mathbf{x})$$
 - They are not random variable as the $\mathbf{x}$'s (better qualified as **uncertain variables**)
 - Even if notation is the same, this probability doesn't have the same frequentist meaning
- Inference of parameter of interest based on this probability
- How is this probability computed ?

$$P(\theta; \mathbf{x}) = \frac{P(\mathbf{x}; \theta)P(\theta)}{P(\mathbf{x})} \text{ (Bayes thm)}$$

More on bayesian approach

In bayesian approach, inference proceeds as follows:

1. Write statistical model
2. Find out what is best prior
3. Compute posterior
4. Conclude from posterior

$$\begin{array}{c} \text{Posterior distribution} \quad \text{Stat. model (likelihood)} \quad \text{Prior distribution} \\ \downarrow \quad \downarrow \quad \downarrow \\ P(\theta; x) = \frac{P(x; \theta)P(\theta)}{P(x)} \quad (P: \text{pmf or pdf}) \\ \uparrow \\ \text{Normalization constant} \end{array}$$

Bayes Theorem

For 2 arbitrary events A and B

$$P(A \cap B) = P(A|B)P(B)$$

$$P(B \cap A) = P(B|A)P(A)$$

As $P(A \cap B) = P(B \cap A)$, one has $P(A|B)P(B) = P(B|A)P(A)$

$$\Rightarrow P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

Simple explanation:

Scenario: You want to know the probability of having a rare disease.

- Initial belief (prior probability): 1% of the population has the disease ($P(A) = 0.01$).
- Medical test accuracy: The test is 95% accurate in detecting the disease ($P(B|A) = 0.95$) but has a 10% false positive rate ($P(B|\neg A) = 0.10$).

Using Bayes' Theorem:

1. Calculate the overall probability of getting a positive test result ($P(B)$) using the law of total probability.

$$P(B) = P(B|A) \cdot P(A) + P(B|\neg A) \cdot P(\neg A)$$

2. Apply Bayes' Theorem to find the updated probability of having the disease given a positive test result ($P(A|B)$).

Result: Even with a positive test result, the updated probability of having the disease is around 8.76%. This demonstrates how Bayes' Theorem adjusts your beliefs based on new evidence while considering the accuracy and potential uncertainties in the test.

4. Parameter Estimation

Frequentist Approach

→ Let X be gaussian with $\sigma=1$ and mean μ

Measurement providing a sample of 4 values: (X_1, X_2, X_3, X_4)

Problem: How we determine μ from these X_i ?

Possible solution: Take the sample mean $M = \frac{1}{4} \sum_{i=1}^4 X_i$

Doing this we hope that we get a value that is close to the true value

$$\mu \simeq M$$

Questions

1. What does it mean when we say that an estimate is "close" to the true value $\hat{\theta} \simeq \theta$?

2. How are estimators determined in general?

Answer 1

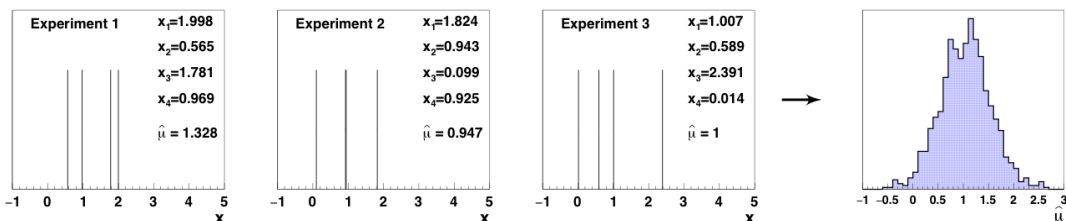
It means nothing → θ is not known

No sense random variable $\simeq$ constant

In order to determine whether an estimator is good or not, need to get back to Fundamental frequentist interpretation:

What would one get if we repeat the same measurement many times?

→ We would get a distribution of $\hat{\theta}$ characterized by a MEAN and a VARIANCE.



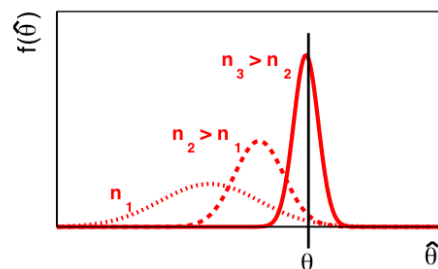
Frequentist Properties of estimators

Three main properties: CONSISTENCY, BIAS, EFFICIENCY

CONSISTENCY

An estimator is consistent when it converges in probability towards the true value as $n \rightarrow \infty$

$$P(|\hat{\theta} - \theta| > \varepsilon) \xrightarrow{n \rightarrow \infty} 0 \quad \forall \varepsilon > 0$$

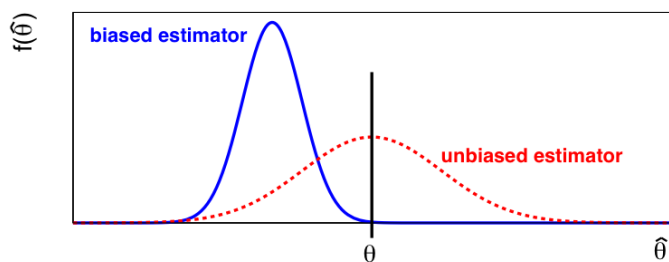


BIAS

An estimator is unbiased when its expectation value is equal to the true value (for all n)

$$E[\hat{\theta}] = \theta$$

The bias of an estimator is defined as $b = E[\hat{\theta}] - \theta$



EFFICIENCY

An estimator is efficient when its variance is equal to the RCF bound

RCF bound:

$$\text{var}[\hat{\theta}] \geq \frac{\left(1 + \frac{\partial b}{\partial \theta}\right)^2}{E\left[\left(\frac{\partial \ln \mathcal{L}}{\partial \theta}\right)^2\right]}$$

RCF bound

sometimes $\geq \frac{1 + \frac{\partial b}{\partial \theta}}{E\left[\frac{\partial^2 \ln \mathcal{L}}{\partial \theta^2}\right]}$

Ideal estimators are consistent, unbiased and efficiency

~> Not always possible $\Rightarrow$ Should at least be consistent

How estimators are determined in general?

1 Maximum Likelihood

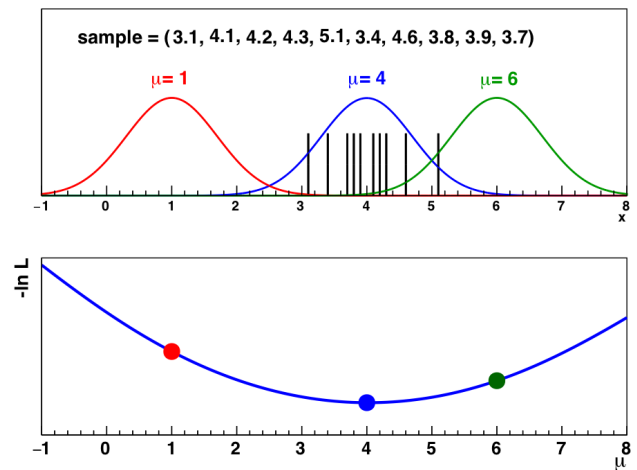
~> Let X be a gaussian r.v. with known σ and unknown mean μ

Suppose we make a measurement for the determination of μ providing 10 values $(x_1, \dots, x_{10})$

Of the 3 gaussian distribution presented

the BLUE ONE has the minimum

value of $-\ln \mathcal{L}$



The Maximum Likelihood estimator $\hat{\theta}_{ML}$ is the parameter that maximizes the Likelihood Function

$$\left. \frac{\partial \mathcal{L}}{\partial \theta} \right|_{\theta = \theta_{ML}} = 0 \quad \text{et}$$

$$\left. \frac{\partial^2 \mathcal{L}}{\partial \theta^2} \right|_{\theta = \hat{\theta}_{ML}} < 0$$

equivalently

$$\left. \frac{\partial (-\ln \mathcal{L})}{\partial \theta} \right|_{\theta = \hat{\theta}_{ML}} = 0 \quad \text{et}$$

$$\left. \frac{\partial^2 (-\ln \mathcal{L})}{\partial \theta^2} \right|_{\theta = \hat{\theta}_{ML}} > 0$$

Properties of estimators

The ML method provides estimators with nice properties

- Functional Invariance $\hat{\zeta}_{ML} = \zeta(\hat{\theta}_{ML})$
- ML estimators are consistent
- If an unbiased and efficient estimator exist $\Rightarrow$ its variance is

$$\text{var}[\hat{\zeta}_{ML}] = \frac{(\partial \zeta / \partial \theta)^2}{\frac{\partial^2 \ln \mathcal{L}}{\partial \theta^2} \Big|_{\theta = \theta_{ML}}}$$

- ML method provides **POINT ESTIMATION**

$\leadsto$ It's possible to have an error bar in addition to a single value?

$\Rightarrow$ YES! **GRAPHICAL METHOD**

GRAPHICAL METHOD

$$\ln \mathcal{L}(\theta) = -\ln \mathcal{L}(\hat{\theta}_{ML}) + (\theta - \hat{\theta}_{ML}) \frac{\partial(-\ln \mathcal{L})}{\partial \theta} \Big|_{\theta = \hat{\theta}_{ML}} + \frac{1}{2} (\theta - \hat{\theta}_{ML})^2 \frac{\partial^2(-\ln \mathcal{L})}{\partial \theta^2} \Big|_{\theta = \hat{\theta}_{ML}} + \dots$$

$$\text{but } \frac{\partial(-\ln \mathcal{L})}{\partial \theta} \Big|_{\theta = \hat{\theta}_{ML}} = 0 \quad \text{thus}$$

$$\ln \mathcal{L}(\theta) = -\ln \mathcal{L}(\hat{\theta}_{ML}) + \frac{1}{2} (\theta - \hat{\theta}_{ML})^2 \frac{\partial^2(-\ln \mathcal{L})}{\partial \theta^2} \Big|_{\theta = \hat{\theta}_{ML}} + \dots$$

$$\leadsto \text{If estimator unbiased and efficient: } \text{var}[\hat{\theta}_{ML}]^{-1} = -\frac{\partial^2 \ln \mathcal{L}}{\partial \theta^2} \Big|_{\theta = \hat{\theta}_{ML}}$$

Thus

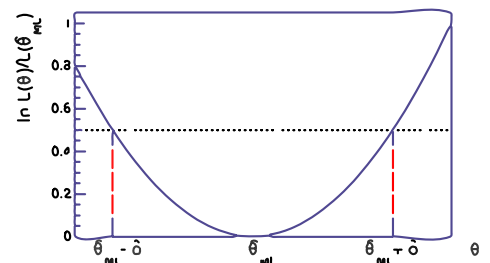
$$-\ln \mathcal{L}(\theta) \simeq -\ln \mathcal{L}(\hat{\theta}_{ML}) + \frac{(\theta - \hat{\theta}_{ML})^2}{2 \text{var}[\hat{\theta}_{ML}]}$$

Conclusion

$$\theta = \hat{\theta}_{ML} \pm \sqrt{\text{var}[\hat{\theta}_{ML}]}$$

$$\text{then } -\ln \mathcal{L}(\theta) \simeq -\ln \mathcal{L}(\hat{\theta}_{ML}) + \frac{1}{2}$$

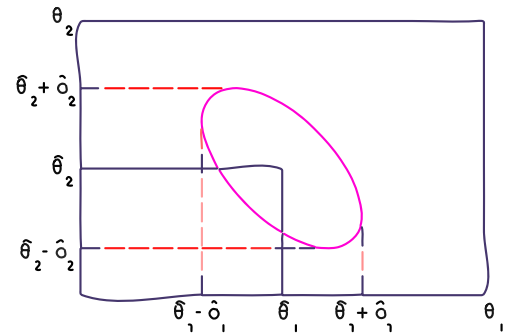
□ This leads to the following graphical method:



Note: doing this we build a so-called "**confidence interval**" for the parameter of interest

→ We'll come back to this notion later

The graphical method also works with 2 parameters of interest



□ Using this method we can:

- Estimate uncertainties on the 2 parameters
- Estimate correlation between the two estimates

2. Least Squares Method

LSM is useful when one is interested in dependence of one variable with one or more other variables

Notations

Y_i : observation i

$m_i(\theta) = \mathbb{E}[Y_i]$ expected value of observation i

V : covariance matrix of the Y_i 's

$$V = \begin{pmatrix} \sigma_1^2 & \dots & \dots & \text{cov}(Y_1, Y_n) \\ \vdots & \ddots & \text{cov}(Y_i, Y_j) & \vdots \\ \text{cov}(Y_n, Y_1) & \dots & \dots & \sigma_n^2 \end{pmatrix}$$

Method

Least Squares estimators are values that minimize $\chi^2(\theta) = (Y - m)^T V^{-1} (Y - m)$

Thus $\frac{\partial^2 \chi}{\partial \theta} \Big|_{\theta = \hat{\theta}_{LS}} = 0$ et $\frac{\partial^2 \chi^2}{\partial \theta^2} \Big|_{\theta = \hat{\theta}_{LS}} > 0$

→ We will consider independent observations

$$\chi^2(\theta) = \sum_{i=1}^n \frac{(Y_i - m_i(\theta))^2}{\sigma_i^2}$$

Least squared and Max likelihood method

For gaussian observations, the likelihood is

$$\mathcal{L} = e^{-\frac{1}{2} (y-m)^T V^{-1} (y-m)}$$

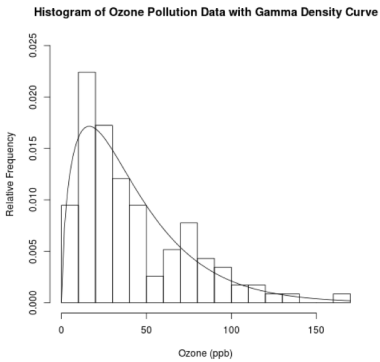
For independent observations

$$\mathcal{L} = \prod_{i=1}^n e^{-\frac{(y_i - m(x_i, \theta))^2}{2\sigma_i^2}}$$

We see that $-2 \ln \mathcal{L} = \chi^2$

$\Rightarrow$ Minimizing χ^2 is equivalent to maximizing the likelihood

Least Squared Method: binned sample



□ **Hypothesis:** the Y_i 's are independent and $Y_i \sim \text{Pois}(m_i)$

$$\Rightarrow \sigma_i = \sqrt{m_i}$$

□ **Remark:** m_i depends on θ

$$m_i(\theta) = v p_i(\theta) = v \int_{C_i} f_X(x; \theta) dx$$

(v =total expected number of events= $\mathbb{E}[\sum Y_i]$)

$$\Rightarrow \sigma_i(\theta) = \sqrt{m_i(\theta)}$$



The χ^2 writes:

$$\chi^2 = \sum_{i=1}^n \frac{(Y_i - m_i(\theta))^2}{m_i(\theta)} \quad (\text{Pearson's } \chi^2)$$

which is often simplified to ("modified LS method")

$$\chi^2 = \sum_{i=1}^n \frac{(Y_i - m_i(\theta))^2}{Y_i} \quad (\text{Neyman's } \chi^2)$$

Remark: approximation $\sigma_i(\theta) = \sqrt{m_i(\theta)} \simeq \sqrt{Y_i}$ often done in various contexts

$\rightarrow$ Equivalent to using for $\sigma(\theta)$ not the true value (unknown) but an estimate:

$$\hat{\sigma}_i(\theta) = \sqrt{Y_i}$$

Bayesian way of estimating parameters

Based on the posterior distribution:

$$P(\theta; x) = \frac{P(x; \theta) P(\theta)}{P(x)} \quad (P: \text{pmf or pdf})$$

Labels: Posterior distribution, Stat. model (likelihood), Prior distribution, Normalization constant

$$\Rightarrow P(\text{model} | \text{data}) = \frac{P(\text{data} | \text{model}) P(\text{model})}{P(\text{data})}$$

Once the posterior $f(\theta|x)$ is found, parameters can be estimated using either the

Mean: $\hat{\theta} = \int \theta f(\theta|x) d\theta$

Mode: $\hat{\theta} = \max_{\theta} f(\theta|x)$

Median: $F_{\theta}(\hat{\theta}) = \int_{-\infty}^{\hat{\theta}} f(\theta|x) d\theta = 1/2$

Any other location measurement quantity

→ Impact of prior on final result decreases as sample size increases.

2D Example

Poissonian measurement with signal and background

$$P(N|s, b) = \frac{(s+b)^N}{N!} e^{-(s+b)}$$

→ b is measured and it's $b = b_0 \pm \sigma$ and you take the following normal prior

$$g(b|b_0, \sigma) = \frac{1}{\sqrt{2\pi} \sigma} e^{-\frac{(b-b_0)^2}{2\sigma^2}}$$

→ Suppose s is totally unknown prior to the measurement and you take a uniform prior: $g(s) \propto \text{constant}$

2D posterior:

$$f(s, b|N) \propto P(N|s, b) \underbrace{g(b|b_0, \sigma) g(s)}_{\text{Priors}}$$

$$\propto \frac{(s+b)^N}{N!} e^{-(s+b)} \times \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(b-b_0)^2}{2\sigma^2}}$$

Also bayesian inference offers an easy way to describe how knowledge increases as measurements are performed

$$g(\theta) \xrightarrow{x} f(\theta|x)$$

Suppose you make two measurements with likelihoods $\mathcal{L}_1$ and $\mathcal{L}_2$

→ After first measurement: $f_2(\theta|x_1) \propto \mathcal{L}_1 \times g(\theta)$

→ After second measurement: we could use

$$f_2(\theta|x_2) \propto \mathcal{L}_2 \times g(\theta)$$

but better to use knowledge acquired from first measurement:

$$f_2(\theta|x_2, x_1) \propto \mathcal{L}_2 \times f_2(\theta|x_1) \propto \mathcal{L}_2 \times \mathcal{L}_1 \times g(\theta)$$

→ Posterior of first measurement used as prior in second one

Generalization to n measurements:

$$\underline{f_n(\theta|x_n, \dots, x_1) \propto \mathcal{L}_n \times \mathcal{L}_{n-1} \times \dots \times \mathcal{L}_1 \times g(\theta)}$$

→ As the number of measurement increases, the likelihood term "wins" and the prior becomes more and more subdominant in final result