

Statistics in Data Analysis

All you ever wanted to know about statistics but never dared to ask

part 5

Paweł Brückman de Renstrom
(pawel.bruckman@ifj.edu.pl)

April 02, 2025

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

Question from the previous lecture

Suppose a beam of particles is known to consist of charged pions and muons. For each particle in the beam we measure a variable t , whose distribution for pions (π) and muons (μ) is

$$f(t; \pi) = \frac{1}{\sqrt{2\pi}\sigma} e^{-(t-\mu_\pi)^2/2\sigma^2}, \quad f(t; \mu) = \frac{1}{\sqrt{2\pi}\sigma} e^{-(t-\mu_\mu)^2/2\sigma^2},$$

where $\mu_\pi = 0$, $\mu_\mu = 2$ and $\sigma = 1$. For each particle we want to test the hypothesis H_0 that it is a pion against the alternative H_1 that it is a muon. The critical region of the test is given by $t > t_c$ where t_c is a given constant.

- 1 Suppose we want to have a test of size $\alpha = 0.05$. Illustrate where the critical region lies and what α means on a sketch of the p.d.f.s $f(t|\pi)$ and $f(t|\mu)$ and show that t_c is numerically about 1.64.
- 2 Suppose a sample of particles is known to consist of 99% pions and 1% muons. What is the purity of the muon sample selected by $t > t_c$? Here, purity means the probability to be a muon given that the particle had $t > t_c$ (i.e., it was rejected as a pion and thus selected as a muon candidate).

Solution to be sent to me before the next lecture

Solution

1 $\alpha = 0.05$ (test size) means that:

$$1 - F(t_c; 0.0, 1.0) = \int_{t_c}^{\infty} f(t|\pi) dt = 0.05$$

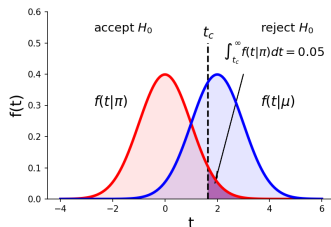
$$1 - \text{erf}(t_c) = 0.05 \implies t_c = \text{erf}^{-1}(0.95) = 1.64485\dots$$

2 For the second part, we need to calculate:

$$1 - \beta = 1 - F(t_c; 2.0, 1.0) = \int_{t_c}^{\infty} f(t|\mu) dt = 0.6388\dots$$

$$N_{\pi} = 0.99 * N, \quad N_{\mu} = 0.01 * N$$

$$\text{purity} = \frac{(1 - \beta)N_{\mu}}{(1 - \beta)N_{\mu} + \alpha N_{\pi}} \simeq \frac{0.01 * 0.639}{0.01 * 0.639 + 0.99 * 0.05} = 0.114\dots$$



The choice of the critical point t_c (selection working point) results in muon selection efficiency of 64% with purity (given the beam composition) of 11.4% .

Maximum Likelihood estimator - reminder

Let $f(x; \theta)$ is a p.d.f. of a known form but unknown parameter θ (more generally $\theta = (\theta_1, \dots, \theta_m)$). Let $x_1, x_2, \dots, x_n$ be a sample of n events drawn from the above p.d.f. Generally, x_i may be a multidimensional vector. We define:

$$L(\theta) = \prod_{i=1}^n f(x_i; \theta) \quad (1)$$

called the **likelihood function**.

- L is technically a joint p.d.f. of x but, assuming a fixed data sample, represents a function of θ .
- The **maximum likelihood** (ML) estimator $\hat{\theta}$ is given by:

$$\frac{\partial L}{\partial \theta_i} = 0, \quad i = 1, \dots, m. \quad (2)$$

- **Log-likelihood function** is commonly used:

$$\log L(\theta) = \sum_{i=1}^n \ln f(x_i; \theta) \quad (3)$$

Variance of ML estimator

RCF bound

Rao-Cramér-Frechet (RCF) inequality states that:

$$V[\hat{\theta}] \geq \left(1 + \frac{\partial b}{\partial \theta}\right)^2 \bigg/ E \left[-\frac{\partial^2 \log L}{\partial \theta^2} \right]. \quad (4)$$

In case of equality (i.e. minimum variance) the estimator is said to be **efficient**.

- E.g., $\hat{\tau} = \frac{1}{n} \sum_{i=1}^n x_i$ is an efficient estimator for the parameter τ . (show!)
- ML are always efficient in the large sample limit!

Assuming efficiency and zero bias, for a general case when $\theta = (\theta_1, \dots, \theta_m)$ we get:

$$(V^{-1})_{ij} = E \left[-\frac{\partial^2 \log L}{\partial \theta_i \partial \theta_j} \right] \longrightarrow (\widehat{V^{-1}})_{ij} = -\frac{\partial^2 \log L}{\partial \theta_i \partial \theta_j} \bigg|_{\theta=\hat{\theta}} \quad (5)$$

For a single parameter θ this reduces to: $\widehat{\sigma^2_{\hat{\theta}}} = \left(-1 / \frac{\partial^2 \log L}{\partial \theta^2} \right) \bigg|_{\theta=\hat{\theta}}$

Variance of ML estimator

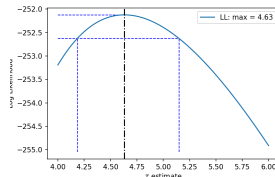
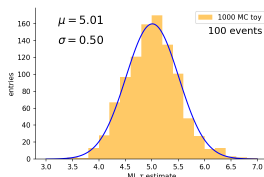
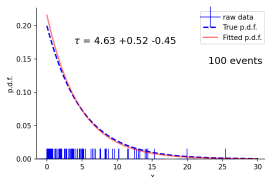
Taylor expansion

Let's take a single parameter θ and expand log-likelihood in Taylor series about the ML estimate:

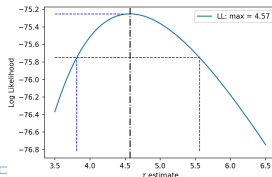
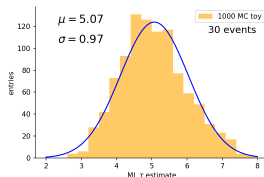
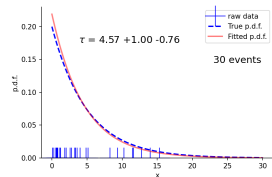
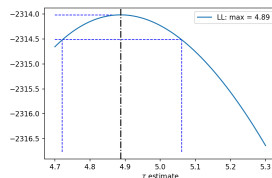
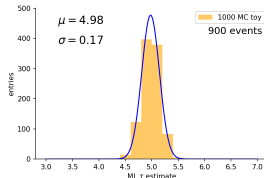
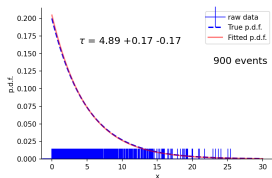
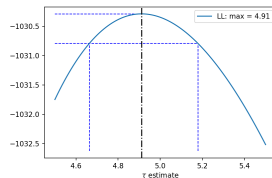
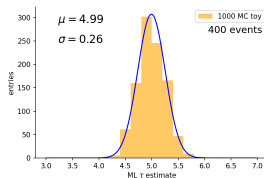
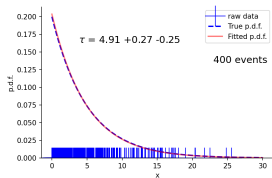
$$\log L(\theta) = \log L(\hat{\theta}) + \left[-\frac{\partial \log L}{\partial \theta} \right] \bigg|_{\theta=\hat{\theta}} (\theta - \hat{\theta}) + \frac{1}{2!} \left[\frac{\partial^2 \log L}{\partial \theta^2} \right] \bigg|_{\theta=\hat{\theta}} (\theta - \hat{\theta})^2 + \dots \quad (6)$$

By definition, $\log L(\hat{\theta}) = \log L_{\max}$ and $\frac{\partial \log L}{\partial \theta} \bigg|_{\theta=\hat{\theta}} = 0$, and so:

$$\log L(\theta) = \log L_{\max} - \frac{(\theta - \hat{\theta})^2}{2\sigma_{\hat{\theta}}^2} \quad \text{or} \quad \log L(\theta \pm \hat{\sigma}_{\hat{\theta}}) = \log L_{\max} - \frac{1}{2}. \quad (7)$$



Variance of ML estimator



Extended Maximum Likelihood

Let $f(x; \theta)$ is a p.d.f. of a known form but unknown parameter θ (more generally $\theta = (\theta_1, \dots, \theta_m)$). Let $x_1, x_2, \dots, x_n$ be a sample of n events drawn from the above p.d.f. Now, let us treat n as a Poisson random variable with mean ν .

- The likelihood takes the form:

$$L(\theta) = \frac{\nu^n}{n!} e^{-\nu} \prod_{i=1}^n f(x_i; \theta) = \frac{e^{-\nu}}{n!} \prod_{i=1}^n \nu f(x_i; \theta) \quad (8)$$

called the **extended likelihood function**.

- If ν is a function of θ then log-likelihood is:

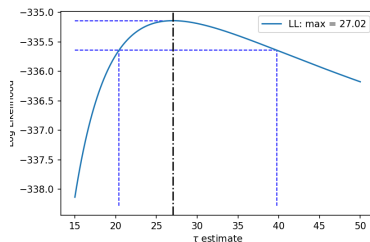
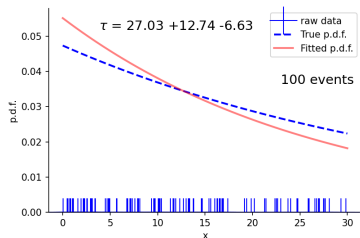
$$\log L(\theta) = n \ln \nu(\theta) - \nu(\theta) + \sum_{i=1}^n \ln f(x_i; \theta) = -\nu(\theta) + \sum_{i=1}^n \ln (\nu(\theta) f(x_i; \theta)), \quad (9)$$

where terms not depending on θ have been dropped.

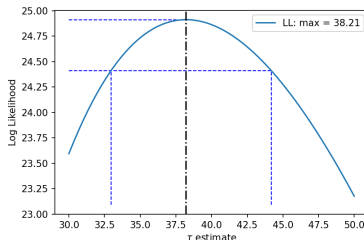
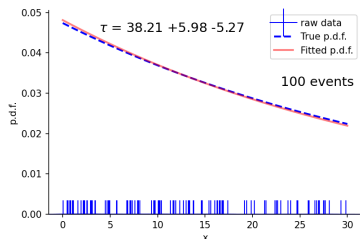
Extended Maximum Likelihood

mean lifetime: $\tau = 40$, $\nu = \nu_0(1 - e^{-T/\tau})$

$$\log L = \sum \ln f(x_i)$$



$$\log L = -\nu + \sum \ln(\nu f(x_i))$$

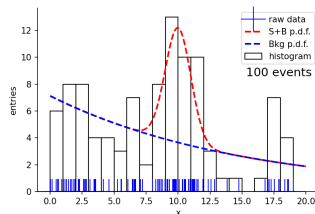


Extended Maximum Likelihood

sample composition

In a trivial case when ν is independent from θ , the derivative of $\log L$ w.r.t. ν gives the estimator $\hat{\nu} = n$. (show!)

However, it is often the case that different classes of events with known p.d.f.'s contribute to the observed distribution and our task is to estimate yields of the individual components.



In a general case of having m contributions, a $\log L$ can be defined as:

$$\log L(\boldsymbol{\mu}) = - \sum_{j=1}^m \mu_j + \sum_{i=1}^n \ln \left(\sum_{j=1}^m \mu_j f_j(x_i) \right), \quad (10)$$

where the vector $\boldsymbol{\mu} = (\mu_1, \dots, \mu_m)$ represents directly the yields of individual contributions. Of course, the fit is capable of estimating all μ_j parameters only if the corresponding p.d.f.'s are different. In general, such a fit results in correlated estimates.

Binned Maximum Likelihood fit

ML fit as we've discussed it, uses measured quantities directly and provides *efficient* estimators. But this approach is not always practical. When a dataset is large it may be very computationally intense. The commonly used alternative is the **binned maximum likelihood** fit.

- The data sample of n_{tot} events has to be histogrammed, yielding a certain number of entries $\mathbf{n} = (n_1, \dots, n_N)$ in N bins.
- The expectation values $\boldsymbol{\nu} = (\nu_1, \dots, \nu_N)$ for the bin are given by:

$$\nu_i(\boldsymbol{\theta}) = n_{\text{tot}} \int_{x_i^{\min}}^{x_i^{\max}} f(x; \boldsymbol{\theta}) dx, \quad (11)$$

where $x_i^{\min/\max}$ are the bin limits.

- If n_{tot} is fixed and one is interested in the *shape* of the distribution, the joint p.d.f. is given by the multinomial distribution:

$$f_{\text{joint}}(\mathbf{n}, \boldsymbol{\nu}) = \frac{n_{\text{tot}}!}{n_1! \dots n_N!} \left(\frac{\nu_1}{n_{\text{tot}}} \right)^{n_1} \dots \left(\frac{\nu_N}{n_{\text{tot}}} \right)^{n_N}, \quad (12)$$

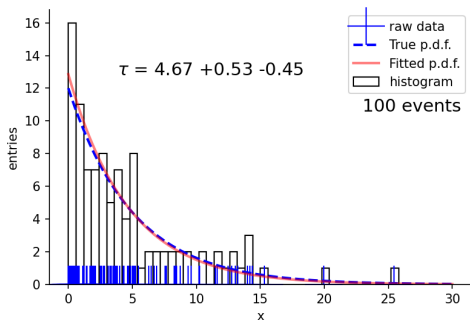
Binned Maximum Likelihood fit

- The resulting log-likelihood (ignoring spurious terms) reads:

$$\log L(\boldsymbol{\theta}) = \sum_{i=1}^N n_i \ln \nu_i(\boldsymbol{\theta}). \quad (13)$$

- In the limit of large number of bins, the binned ML approaches the standard one. However, for coarser binning may result in suboptimal results.

- $\tau = 5.0$
- $n_{\text{tot}} = 100$
- $\sigma_{\hat{\tau}} = 0.5$
- $N = 50$



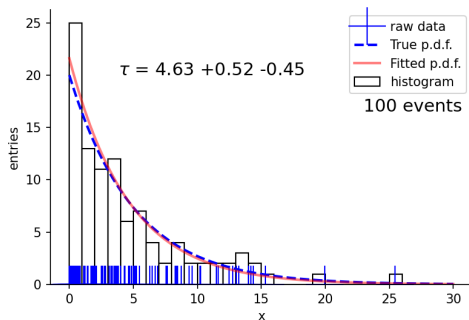
Binned Maximum Likelihood fit

- The resulting log-likelihood (ignoring spurious terms) reads:

$$\log L(\boldsymbol{\theta}) = \sum_{i=1}^N n_i \ln \nu_i(\boldsymbol{\theta}). \quad (14)$$

- In the limit of large number of bins, the binned ML approaches the standard one. However, for coarser binning may result in suboptimal results.

- $\tau = 5.0$
- $n_{\text{tot}} = 100$
- $\sigma_{\hat{\tau}} = 0.5$
- $N = 30$



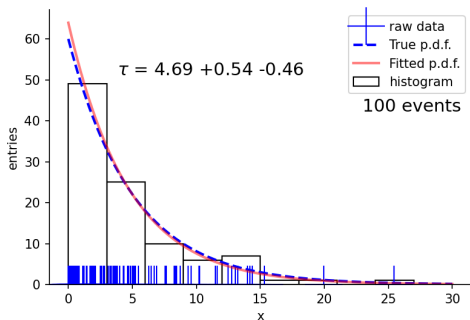
Binned Maximum Likelihood fit

- The resulting log-likelihood (ignoring spurious terms) reads:

$$\log L(\boldsymbol{\theta}) = \sum_{i=1}^N n_i \ln \nu_i(\boldsymbol{\theta}). \quad (15)$$

- In the limit of large number of bins, the binned ML approaches the standard one. However, for coarser binning may result in suboptimal results.

- $\tau = 5.0$
- $n_{\text{tot}} = 100$
- $\sigma_{\hat{\tau}} = 0.5$
- $N = 10$



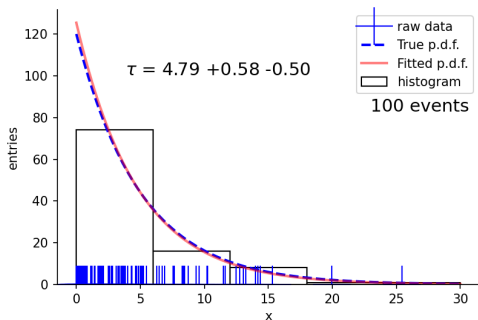
Binned Maximum Likelihood fit

- The resulting log-likelihood (ignoring spurious terms) reads:

$$\log L(\boldsymbol{\theta}) = \sum_{i=1}^N n_i \ln \nu_i(\boldsymbol{\theta}). \quad (16)$$

- In the limit of large number of bins, the binned ML approaches the standard one. However, for coarser binning may result in suboptimal results.

- $\tau = 5.0$
- $n_{\text{tot}} = 100$
- $\sigma_{\hat{\tau}} = 0.5$
- $N = 5$



Binned Maximum Likelihood fit

extended binned log-likelihood

We may take an alternative approach whereby the total number of events is itself considered a Poisson-distributed random variable with the mean ν_{tot} .

- The expectation values $\boldsymbol{\nu} = (\nu_1, \dots, \nu_N)$ now depend on ν_{tot} and $\boldsymbol{\theta}$:

$$\nu_i(\nu_{\text{tot}}, \boldsymbol{\theta}) = \nu_{\text{tot}} \int_{x_i^{\min}}^{x_i^{\max}} f(x; \boldsymbol{\theta}) dx. \quad (17)$$

- Content of each bin becomes a Poisson random variable. We are now concerned both with the *shape* and the *normalisation* of the sample distribution. The joint p.d.f. is given by:

$$f_{\text{joint}}(\boldsymbol{n}, \boldsymbol{\nu}) = \prod_{i=1}^N \frac{\nu_i^{n_i}}{n_i!} e^{-\nu_i} \quad (18)$$

Binned Maximum Likelihood fit

extended binned log-likelihood

- Taking the logarithm of the joint p.d.f. and dropping terms that do not depend on the parameters gives:

$$\log L(\nu_{\text{tot}}, \boldsymbol{\theta}) = -\nu_{\text{tot}} + \sum_{i=1}^N n_i \ln \nu_i(\nu_{\text{tot}}, \boldsymbol{\theta}), \quad (19)$$

- which is the binned version of the extended log-likelihood function.
- Here, similar conclusions apply as the ones drawn for the unbinned ML fit.

Example of a 2D ML fit

Data events have been collected in an experiment yielding a scalar random variable x :

- The sample consists of a mixture of *signal* and *background* events with known p.d.f.'s.
- Signal p.d.f.: $f_S(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-(x-\mu)^2/2\sigma^2}$,
- Background p.d.f.: $f_B(x; \lambda) = \frac{x}{\lambda^2} e^{-x/\lambda}$, ($\lambda = 5$),
- The variable x has been recorded in the range $(0, 30)$.
- No assumption is made about the yields N_S & N_B .

We want to:

- 1 Estimate number of signal N_S and background N_B events in the observed sample,
- 2 assess the error on the fitted N_S , N_B .

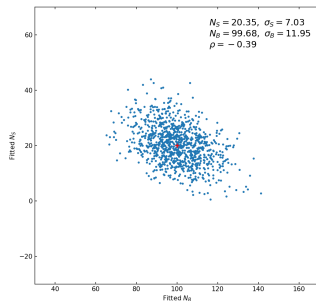
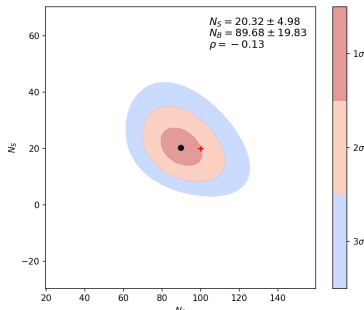
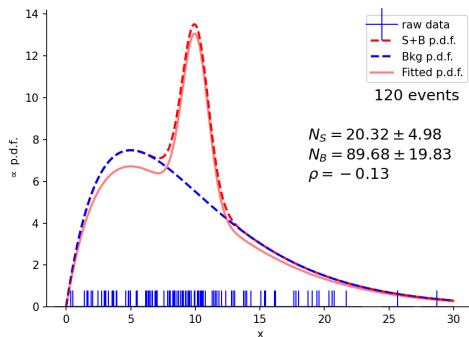
We use the extended ML fit:

$$\log L(N_S, N_B) = -N_S - N_B + \sum_{i=1}^n \ln (N_S f_S(x_i; \mu, \sigma) + N_B f_B(x_i; \lambda)). \quad (20)$$

Example of a 2D ML fit

- $\lambda = 5.0, \mu = 10.0, \sigma = 1.0$.
- Uncertainties from the $\log L$ profile ↗
- Scatter plot of the fit for 1000 MC events ↘

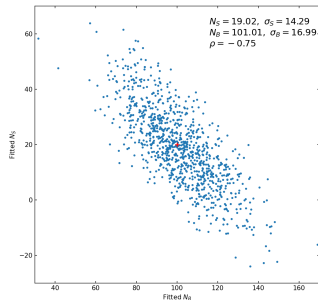
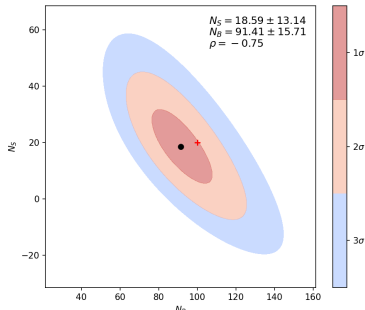
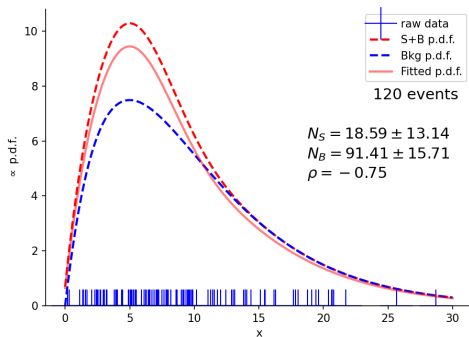
■ Example fit:
$$\hat{V} = \left[-\frac{\partial^2 \log L}{\partial \theta_i \partial \theta_j} \bigg|_{\theta = \hat{\theta}} \right]^{-1}$$



Example of a 2D ML fit

- $\lambda = 5.0, \mu = 5.0, \sigma = 3.0$.
- Uncertainties from the $\log L$ profile ↗
- Scatter plot of the fit for 1000 MC events ↘

■ Example fit:
$$\hat{V} = \left[-\frac{\partial^2 \log L}{\partial \theta_i \partial \theta_j} \Big|_{\theta = \hat{\theta}} \right]^{-1}$$



The method of least-squares (LS)

Consider fitting a function to a set of n measurements (x_i, y_i) , where y_i is assumed to be Gaussian random centered about the true value $\lambda_i(x_i, \boldsymbol{\theta})$ with standard deviation σ_i . $\boldsymbol{\theta} = (\theta_1, \dots, \theta_m)$ are the function parameters we seek to estimate. The measurements are assumed to be mutually independent. The joint p.d.f. is given by:

$$g(\mathbf{y}; \boldsymbol{\lambda}, \boldsymbol{\sigma}) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma_i} \exp\left(-\frac{(y_i - \lambda_i)^2}{2\sigma_i^2}\right) \quad (21)$$

Taking the logarithm and ignoring terms not dependent on $\boldsymbol{\lambda}$ gives the log-likelihood function:

$$\log L(\boldsymbol{\theta}) = -\frac{1}{2} \sum_{i=1}^n \frac{(y_i - \lambda(x_i; \boldsymbol{\theta}))^2}{\sigma_i^2} \quad (22)$$

$$\text{or: } \chi^2(\boldsymbol{\theta}) = \sum_{i=1}^n \frac{(y_i - \lambda(x_i; \boldsymbol{\theta}))^2}{\sigma_i^2}, \quad (23)$$

where the latter χ^2 needs to be minimized in order to find LS estimators $\hat{\boldsymbol{\theta}}$.

The method of least-squares (LS)

General case of correlated Gaussian

The $\chi^2(\boldsymbol{\theta})$ function can be defined also for more general case of correlated measurements, as long as they can be described by a n -dimensional Gaussian:

$$\chi^2(\boldsymbol{\theta}) = \sum_{i,j=1}^n (y_i - \lambda(x_i; \boldsymbol{\theta})) (V^{-1})_{ij} (y_j - \lambda(x_j; \boldsymbol{\theta})) \quad (24)$$

which reduces to E.q. 23 when V is diagonal.

Linear LS fit

Let us consider a special case when predictions λ depend *linearly* on θ .

$$\lambda(x_i; \theta) = \sum_{j=1}^m a_j(x_i) \theta_j = \sum_{j=1}^m A_{ij} \theta_j = (A\theta)_i \quad (25)$$

The χ^2 can be written as:

$$\chi^2(\theta) = (\mathbf{y} - \boldsymbol{\lambda})^T V^{-1} (\mathbf{y} - \boldsymbol{\lambda}) = (\mathbf{y} - A\theta)^T V^{-1} (\mathbf{y} - A\theta), \quad (26)$$

and the χ^2 minimum w.r.t. θ can be found **analytically!**

$$\frac{\partial \chi^2}{\partial \theta} = -2(A^T V^{-1} \mathbf{y} - A^T V^{-1} A \theta) = 0 \quad (27)$$

$$\hat{\theta} = (A^T V^{-1} A)^{-1} A^T V^{-1} \mathbf{y}, \quad (28)$$

i.e. the solution is a linear combination of the measurements y_i , provided matrix $A^T V^{-1} A$ can be inverted (is not singular).

Gauss-Markov theorem: $\hat{\theta}$ is unbiased and has minimum variance independently of n and individual measurements p.d.f.'s.

Linear LS fit

The $\text{cov}[\hat{\boldsymbol{\theta}}]$ is obtained using error propagation:

$$\text{cov}[\hat{\boldsymbol{\theta}}] \equiv U = \frac{\partial \hat{\boldsymbol{\theta}}}{\partial \mathbf{y}} V \frac{\partial \hat{\boldsymbol{\theta}}^T}{\partial \mathbf{y}} = (A^T V^{-1} A)^{-1} \quad (29)$$

Note that we reproduce the RCF bound:

$$(U^{-1})_{ij} = (A^T V^{-1} A)_{ij} = \frac{1}{2} \left[\frac{\partial^2 \chi^2}{\partial \theta_i \partial \theta_j} \right]_{\boldsymbol{\theta}=\hat{\boldsymbol{\theta}}} = - \left[\frac{\partial^2 \log L}{\partial \theta_i \partial \theta_j} \right]_{\boldsymbol{\theta}=\hat{\boldsymbol{\theta}}} \quad (30)$$

$$\begin{aligned} \chi^2(\boldsymbol{\theta}) &= \chi^2(\hat{\boldsymbol{\theta}}) + \frac{1}{2} \sum_{i,j=1}^m \left[\frac{\partial^2 \chi^2}{\partial \theta_i \partial \theta_j} \right]_{\boldsymbol{\theta}=\hat{\boldsymbol{\theta}}} (\theta_i - \hat{\theta}_i)(\theta_j - \hat{\theta}_j) \\ &= \chi^2(\hat{\boldsymbol{\theta}}) + (\boldsymbol{\theta} - \hat{\boldsymbol{\theta}})^T U^{-1} (\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}), \end{aligned} \quad (31)$$

gives us the contour of one standard deviation from LS estimates:

$$\chi^2(\hat{\boldsymbol{\theta}} \pm \boldsymbol{\sigma}) = \chi^2(\hat{\boldsymbol{\theta}}) + 1 = \chi_{\min}^2 + 1 \quad (32)$$

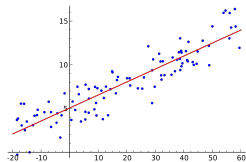
The contour holds even for non-linear parameters in analogy to the $\log L$ Taylor expansion.

Linear LS fit

A super-simple example - linear regression

Assume a straight line fit through n measurements (x_i, y_i) , all with the same error σ :

$$\chi^2(p_0, p_1) = \sum_{i=1}^n \frac{(y_i - (p_0 + p_1 x_i))^2}{\sigma^2} \quad (33)$$



$$A = \begin{pmatrix} 1 & x_1 \\ \cdot & \cdot \\ \cdot & \cdot \\ 1 & x_n \end{pmatrix}, \quad V = \begin{pmatrix} \sigma^2 & \dots & 0 \\ \dots & \sigma^2 & \dots \\ 0 & \dots & \sigma^2 \end{pmatrix}, \quad U = \frac{\sigma^2}{n \sum x_i^2 - (\sum x_i)^2} \begin{pmatrix} -\sum x_i^2 & -\sum x_i \\ -\sum x_i & n \end{pmatrix} \quad (34)$$

When unknown, σ can be estimated from the fit itself: $\hat{\sigma}^2 = \sum_{i=1}^n \frac{(y_i - (\hat{p}_0 + \hat{p}_1 x_i))^2}{n-2}$

Linear LS approximation (linear expansion) has a lot of applications, allowing for solutions to problems with very large number of parameters using linear algebra. For convergence, it often involves multiple iterations of the fit!

Least Squares with binned data

LS is also used to fit p.d.f. to data binned in a histogram. Let $f(x; \boldsymbol{\theta})$ be the hypothetical p.d.f. and y_i contents of bin i . The number of entries predicted in the bin $\lambda_i = E[y_i]$ is:

$$\lambda_i(\boldsymbol{\theta}) = n \int_{x_i^{\min}}^{x_i^{\max}} f(x; \boldsymbol{\theta}) dx = np_i(\boldsymbol{\theta}). \quad (35)$$

Making an approximation of Gaussian distribution of y_i the χ^2 is given as:

$$\chi^2(\boldsymbol{\theta}) = \sum_{i=1}^N \frac{(y_i - np_i(\boldsymbol{\theta}))^2}{np_i(\boldsymbol{\theta})}, \quad (36)$$

where $\sqrt{np_i(\boldsymbol{\theta})}$ is the standard deviation for the Poisson distribution. Sometimes the **modified** LS (MLS) is used instead:

$$\chi_M^2(\boldsymbol{\theta}) = \sum_{i=1}^N \frac{(y_i - np_i(\boldsymbol{\theta}))^2}{y_i}. \quad (37)$$

Careful! This is fine for large statistics. Think of poorly populated or empty bins!

Least Squares with binned data

It might be tempting to use LS method in order to fit the total yield of the distribution by introducing additional free parameter ν :

$$\lambda_i(\boldsymbol{\theta}, \nu) = \nu \int_{x_i^{\min}}^{x_i^{\max}} f(x; \boldsymbol{\theta}) dx = \nu p_i(\boldsymbol{\theta}). \quad (38)$$

However, minimising the χ^2 by setting $\frac{\partial \chi^2}{\partial \nu} = 0$ results in biased estimates of n . We get¹:

- $\hat{\nu}_{\text{LS}} = n + \frac{\chi^2}{2},$
- $\hat{\nu}_{\text{MLS}} = n - \chi^2.$

Note: This can be corrected for, but should be considered at all times.

¹Exercise: Show it by explicit calculation

Putting it to work...

1D fit of the signal yield

N data events have been collected in an experiment yielding a scalar random variable x :

- The sample consists of a mixture of *signal* and *background* events with known p.d.f.'s.
- Background p.d.f.: $f_B(x) = \frac{1}{\tau}e^{-x/\tau}$, $\tau = 5$,
- Signal p.d.f.: $f_S(x) = \frac{1}{\sqrt{2\pi\sigma^2}}e^{-(x-\mu)^2/2\sigma^2}$, $\mu = 10$, $\sigma = 3$,
- The variable x has been recorded in the range $(0, 30)$.
- No assumption about background yield can be made: $N = N_S + N_B$.

Our task is to:

- 1 Estimate number of signal events N_S in the observed sample,
- 2 assess the error of the N_S estimate from the $\log L$ or χ^2 profile.

MIND: This is NOT an extended fit.

Putting it to work...

1D fit of the signal yield

We shall use three strategies:

- 1 the *Unbinned Maximum Likelihood* fit:
[Unbinned ML Python notebook template in Colab](#)
- 2 the *Binned Maximum Likelihood* fit; 10 bins over (0, 30):
[Binned ML Python notebook template in Colab](#)
- 3 the *Binned Least Squares* (NOT modified) fit; 10 bins over (0, 30):
[Binned LS Python notebook template in Colab](#)

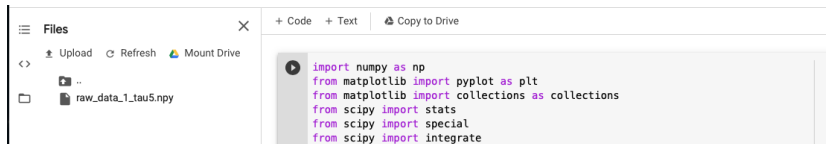
and four data samples:

- 1 [pickled data sample 1 from GitHub](#)
- 2 [pickled data sample 2 from GitHub](#)
- 3 [pickled data sample 3 from GitHub](#)
- 4 [pickled data sample 4 from GitHub](#)

All shall be executed on the *Google Colaboratory* platform.

Detailed instructions

- Click on one of the Python notebook links in order to open it in *Google Colaboratory*.
- Click to download the assigned dataset file from GitHub.
- Click on one of the *Files* icon on the left bar of your *Colab* interface. If you cannot see any datafiles, click on the *Upload* button and select previously downloaded file. As a result you should see:



Hints part 1: fit the N_s , $\log L$ or χ^2 function

→ You need p.d.f. normalization factors in the range (0,30). For this purpose calculate scales and scaleb just as it is done in the main part of the Python script. You will need `xmin`, `xmax`, `tau0`, `mu0` and `sigma0`.

→ You need the total number of collected events. For the unbinned ML this is just the length of the data vector (`len(data)`). For the binned methods loop over the `hist` array and sum all entries. `nbins=len(hist)` gives you the number of bins.

→ For the unbinned ML you need to loop over the data array, for each entry calculate the normalized p.d.f.'s (gauss & decay) and accumulate $\log L$ according to Eq. (25) of lecture 4 and using the combined S+B p.d.f.

→ For the binned methods you need to loop over the bins, the `hist` array (e.g. for `k` in `range(nbins)`). You need to get the prediction for the bin by integrating the normalized p.d.f., see Eq. (13) of lecture 5. Bin k is delimited by `binsy[k]` and `binsy[k+1]`.

→ For the binned ML increment the $\log L$ using Eq. (13) of lecture 5 and the combined S+B p.d.f.

→ For the binned LS increment the χ^2 using Eq. (36) of lecture 5 and the combined S+B p.d.f.

→ NOTE: We are fitting just one free parameter, N_S (coded as `mus`). Make sure you properly define the combined S+B p.d.f. using N_S and the total number of collected events.

Hints part 2: estimate the error on N_s using $\log L$ or χ^2 profile

- The code provides you with the `p1` & `ll_array` arrays which contain the estimated N_S and the corresponding value of either $\log L$ or χ^2 , respectively.
- Your task is to find values of N_S corresponding to $+1\sigma$ and -1σ about the fitted value (coded as `sigma_neg` & `sigma_pos`).
- For the purpose, recall Eq. (7) and Eq. (32) of lecture 5.

Thank you

Back-up

Least Squares with binned data

χ^2 estimators of the yield

The total yield of the distribution parameterised by an additional free parameter ν : $\lambda_i = \nu p_i$. We have n events distributed over N bins.

Minimising the χ^2 by setting $\frac{\partial \chi^2}{\partial \nu} = 0$ results in biased estimates of n :

LS:

$$\begin{aligned}\chi^2 &= \sum_{i=1}^N \frac{(y_i - \nu p_i)^2}{\nu p_i}, \quad \frac{\partial \chi^2}{\partial \nu} = -2 \sum_{i=1}^N \frac{(y_i - \nu p_i) p_i}{\nu p_i} - \sum_{i=1}^N \frac{(y_i - \nu p_i)^2 p_i}{(\nu p_i)^2} = 0 \\ &- 2 \sum_{i=1}^N (y_i - \nu p_i) - \sum_{i=1}^N \frac{(y_i - \nu p_i)^2}{\nu p_i} = -2n + 2\nu - \chi^2 = 0 \\ \implies \hat{\nu}_{\text{LS}} &= n + \frac{\chi^2}{2}\end{aligned}\tag{39}$$

MLS:

$$\begin{aligned}\chi^2 &= \sum_{i=1}^N \frac{(y_i - \nu p_i)^2}{y_i}, \quad \frac{\partial \chi^2}{\partial \nu} = -2 \sum_{i=1}^N \frac{(y_i - \nu p_i) p_i}{y_i} = 0 \implies \nu \sum_{i=1}^N \frac{p_i^2}{y_i} = 1 \\ \chi^2 &= \sum_{i=1}^N y_i - 2\nu \sum_{i=1}^N p_i + \nu^2 \sum_{i=1}^N \frac{p_i^2}{y_i} = n - 2\nu + \nu = n - \nu \\ \implies \hat{\nu}_{\text{MLS}} &= n - \chi^2\end{aligned}\tag{40}$$