

Υπολογιστική Φυσική Στοιχειωδών Σωματιδίων

Εκτιμητές–Estimators

Υπολογισμός των παραμέτρων & προσαρμογή (fitting) των δεδομένων

- Ενδιαφέρει ο υπολογισμός των παραμέτρων της συνάρτησης πυκνότητας πιθανότητας (pdf), π.χ της $f(x, \theta)$ από ένα set μετρήσεων
- Οι συχνότερα χρησιμοποιούμενες μέθοδοι είναι:
 - της μεγιστοποίησης της πιθανότητας
(Maximum Likelihood)
 - των ελαχίστων τετραγώνων (Least squares)
- Πλεονεκτήματα και μειονεκτήματα των δύο μεθόδων
- Θα αναφερθούμε στα πολύ βασικά κάθε μεθόδου

Σύγκριση δεδομένων με την θεωρία Εκτιμητές των παραμέτρων κατανομών

- Περί της διαδικασίας 'εκτίμησης' :

Τι μπορούμε να πουμε για τα
δεδομένα με βάση τις
παραμέτρους της θεωρίας



Θεωρία

Δεδομένα



Τι μπορούμε να πουμε για τις
παραμέτρους της θεωρίας με
βάση τα δεδομένα

Τι είναι ένας εκτιμητής

$$\hat{\mu}(\{x\}) = \frac{1}{N} \sum_i x_i$$

$$\hat{\mu}(\{x\}) = \frac{x_{\max} + x_{\min}}{2}$$

$$\hat{V}(\{x\}) = \frac{1}{N} \sum_i (x_i - \hat{\mu})^2$$

$$\hat{V}(\{x\}) = \frac{1}{N-1} \sum_i (x_i - \hat{\mu})^2$$

- Η διαδικασία που μας παρέχει την τιμή για μία παράμετρο ή μία ιδιότητα μιάς κατανομής, σαν συνάρτηση των δεδομένων μας: $\hat{\theta}$

Πότε ένας εκτιμητής είναι αξιόπιστος?

- Όταν έχει συνέπεια (**consistent**) $\lim_{N \rightarrow \infty} \langle \hat{a} \rangle = a$
- Όταν είναι αμερόληπτος (**unbiased**) $b = E[\hat{\alpha}] - a$
 $\langle \hat{a} \rangle = \iiint \dots \hat{a}(x_1, x_2, \dots) P(x_1; a) P(x_2; a) P(x_3; a) \dots dx_1 dx_2 \dots = a$
- Όταν η απόκλιση είναι ελάχιστη
$$V(\hat{a}) = \left\langle (\hat{a} - \langle \hat{a} \rangle)^2 \right\rangle \quad (\text{minimum})$$
- Θα πρέπει να είναι γρήγορος (**efficient**)

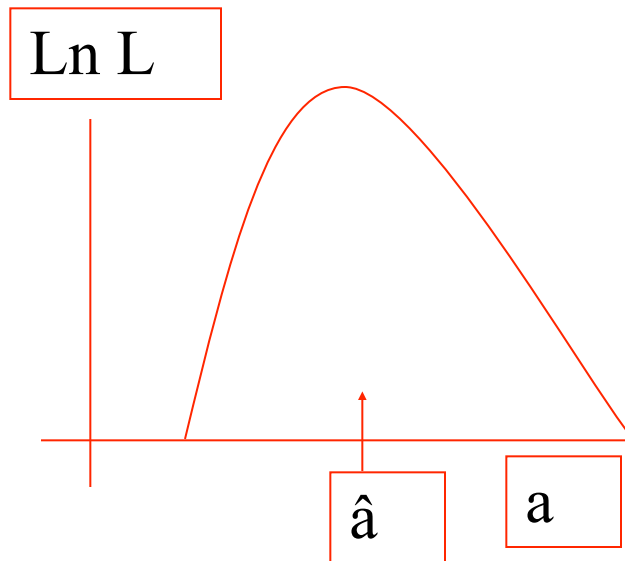
Η συνάρτηση πιθανότητας (Likelihood function)

- Σύνολο δεδομένων $\{x_1, x_2, x_3, \dots, x_N\}$ (κάθε x_i μπορεί να είναι πολυδιάστατο)
- Η πιθανότητα για κάθε x_i εξαρτάται από παράμετρο α που επίσης μπορεί να είναι πολυδιάστατη
- Η συνολική πυκνότητα πιθανότητας είναι:
$$P(x_1; \alpha) P(x_2; \alpha) P(x_3; \alpha) \dots P(x_N; \alpha) = L(x_1, x_2, x_3, \dots, x_N; \alpha)$$

Likelihood

Υπολογισμός της Maximum Likelihood (ML)

- Από το σύνολο των δεδομένων $\{x_1, x_2, x_3, \dots, x_N\}$ υπολογίζουμε την τιμή της παραμέτρου α που **μεγιστοποιεί την Likelihood**
- Στην πράξη συνήθως μεγιστοποιούμε τον λογάριθμο της L

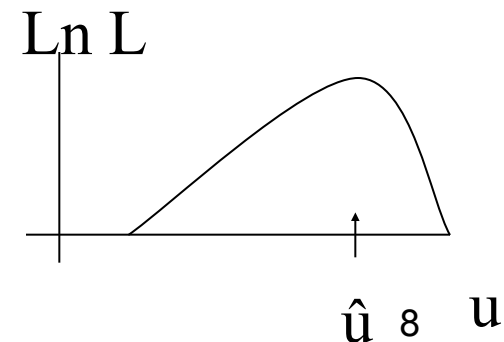
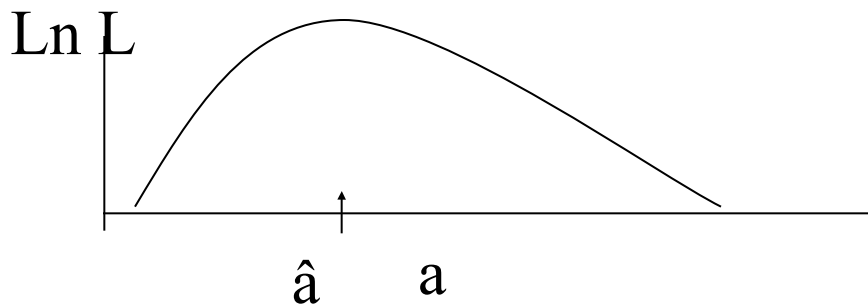


$$\left. \frac{dL}{dA} \right|_{a=\hat{a}} = 0$$

Η ML έχει μερικές χρήσιμες ιδιότητες

Ιδιότητες της ML

- Έχει συνέπεια
- Για μεγάλο σύνολο δεδομένων N είναι ‘αμερόληπτη’ (**unbiased**)
- Είναι αποδοτική (**efficient**) για μεγάλο N (η απόκλιση $\forall$ είναι ελάχιστη)
- Για μικρό N πρέπει να είμαστε προσεκτικοί (**biased**)
- Αν χρησιμοποιήσουμε την συνάρτηση $u(a)$, τότε και $\hat{u}=u(\hat{a})$ (**invariant, but does not follow that $\hat{u}$ unbiased if $\hat{a}$ unbiased**)



Περισσότερα για την (ML)

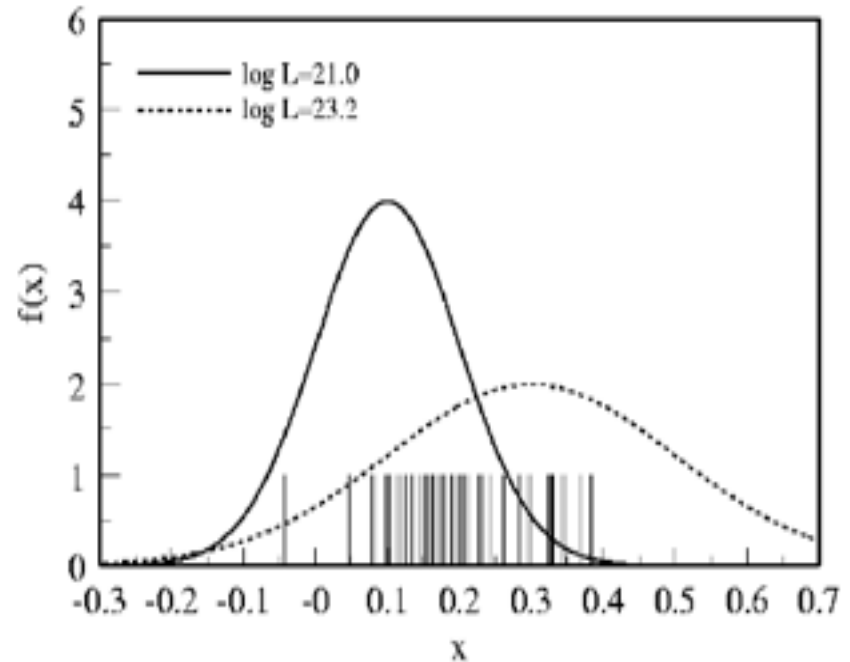
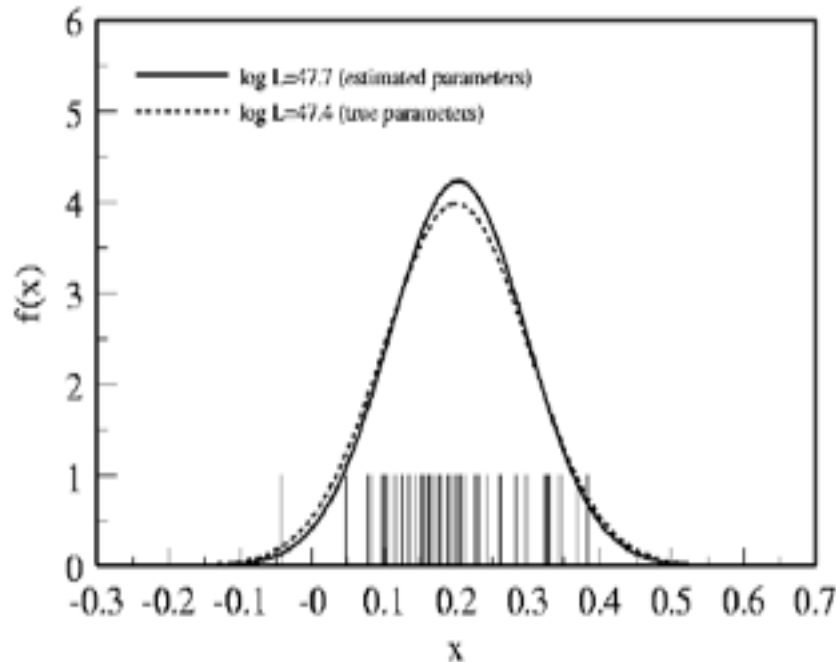
- Δίνει την πιο πιθανή τιμή για τα δεδομένα που έχουμε -όχι κατ'ανάγκη και την πραγματική τιμή της παραμέτρου που προσαρμόζουμε
- Συνήθως απαιτούνται αριθμητικές μέθοδοι για την ελαχιστοποίηση/μεγιστοποίηση
- Για περισσότερες από μία μεταβλητές η ελαχιστοποίηση/μεγιστοποίηση δεν είναι εύκολη
- Χρησιμοποιούμε MINUIT αλλά προσοχή στο αρνητικό πρόσημο του λογαρίθμου!
- Αν η συνάρτηση $P(x;\alpha)$ της πιθανότητας που χρησιμοποιούμε είναι λάθος ο υπολογισμός της α ΔΕΝ έχει νόημα!

Παράδειγμα ML

- Έστω 50 σημεία μιας μεταβλητής x με κατανομή gauss: $G(\mu=0.2, \sigma=0.1)$

Οι καλύτερες τιμές για μ και σ και οι πραγματικές

Τιμές των μ και σ μακριά από το μέγιστο της $\log L$



Maximum Likelihood

- Το μέγιστο βρίσκεται είτε γραφικά από τις τιμές του $\log L$ για διάφορα α είτε από διαφόριση της $\log L$ που άλλοτε μπορεί να γίνει αναλυτικά και άλλοτε με αριθμητικές μεθόδους
- Παράδειγμα : ο υπολογισμός του χρόνου ζωής τ μίας κατάστασης από σύνολο N μετρήσεων του χρόνου $\{t_i\}$. Η $\ln L$ συνάρτηση είναι:

$$\ln L = \sum \ln\left(\frac{1}{\tau} e^{-t_i/\tau}\right) = \sum -\frac{t_i}{\tau} - \ln \tau$$

- Διαφορίζουμε την $\ln L$ και βρίσκουμε την τιμή τ που μηδενίζει την παράγωγο

$$\frac{d \ln L}{d \tau} = \sum \left(\frac{t_i}{\tau^2} - \frac{1}{\tau}\right), \sum \left(\frac{t_i}{\tau^2} - \frac{1}{\tau}\right) = 0 \Rightarrow$$

$$\hat{\tau} = \frac{1}{N} \sum t_i$$

Υπολογισμός σφαλμάτων στον εκτιμητή ML

- Αν το πλήθος των μετρήσεων N πολύ μεγάλο τότε η διαφορά $(\hat{a} - a_0)$ έχει κατανομή gauss με $\mu=0$ και απόκλιση σ :

$$\sigma_a^2 = V(\hat{a}) = -\frac{1}{\left(\frac{d^2 \ln L}{da^2} \Big|_{a_0} \right)} = -\frac{1}{\left\langle \frac{d^2 \ln L}{da^2} \right\rangle}$$

=minimum variance bound (MVB)=το όριο ακριβείας του εκτιμητή δεν μπορεί να είναι μικρότερο του V

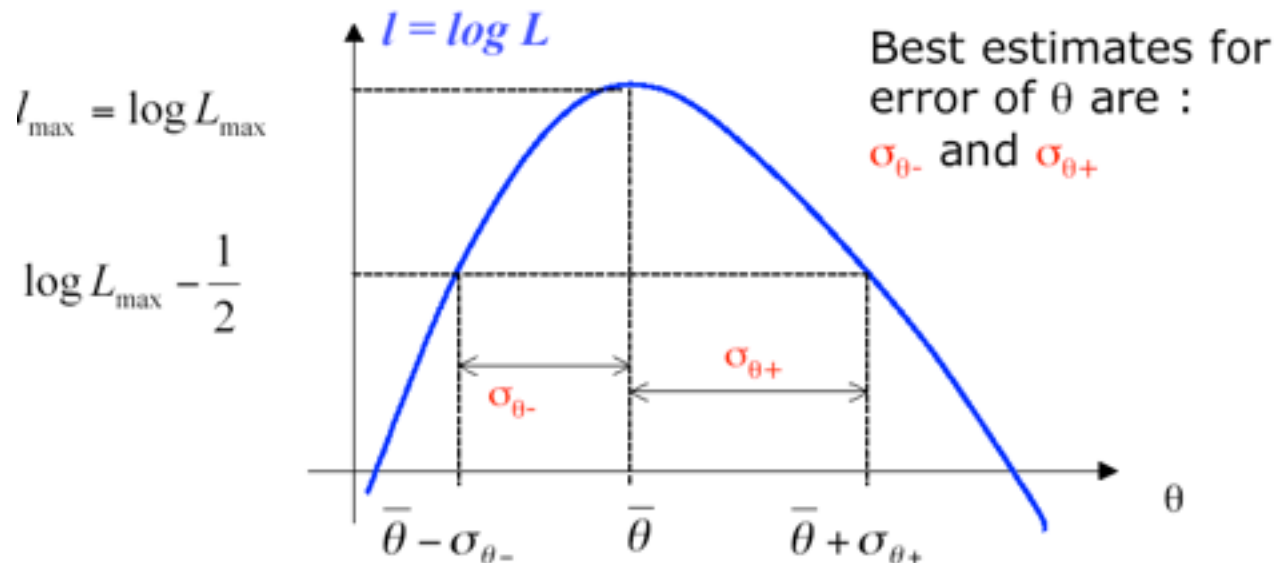
- Η κατανομή πιθανότητας για το $\hat{a}$ είναι gauss αποδεικνύεται από το CLT

Υπολογισμός σφαλμάτων στον εκτιμητή ML

- Αποδεικνύεται ότι για μεγάλο N η συνάρτηση L (likelihood) είναι *gauss* και άρα η συνάρτηση $\log L$ είναι *παραβολή*.
- Επομένως το σφάλμα στον εκτιμητή το βρίσκουμε γραφικά : από την γραφική παράσταση της $\log likelihood$ που πήραμε από τα δεδομένα
- 1σ από την κορυφή $\log L = \log L_{\max} - 0.5$
 2σ από την κορυφή $\log L = \log L_{\max} - 2.0$
 3σ από την κορυφή $\log L = \log L_{\max} - 4.5$...κλπ

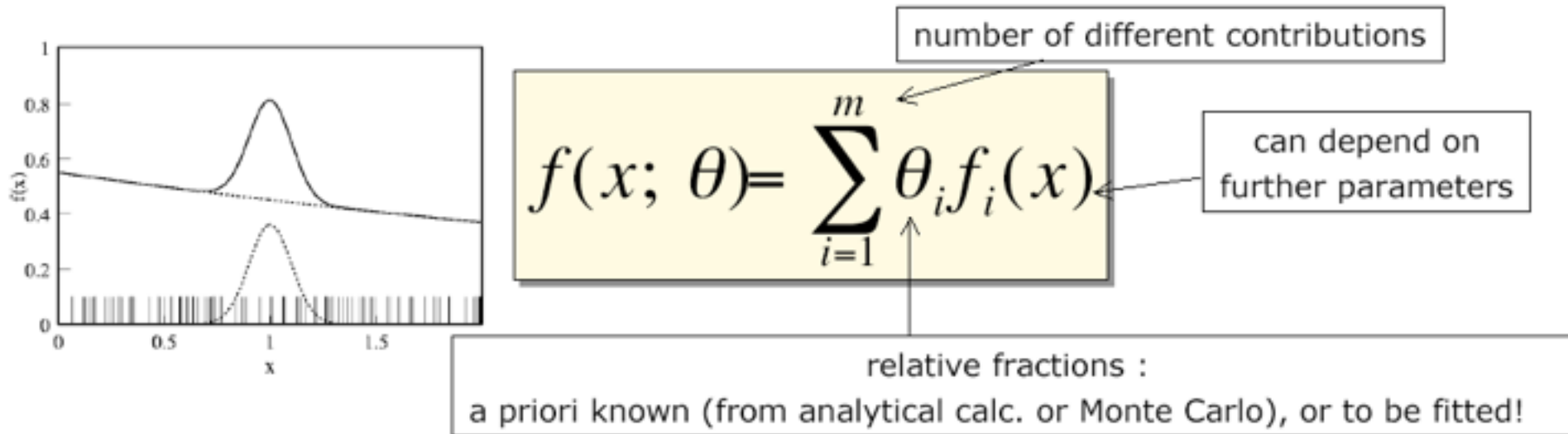
Υπολογισμός σφαλμάτων στον εκτιμητή ML (ασυμμετρικά σφάλματα)

- Εάν το πλήθος των γεγονότων N είναι μικρό, η L ΔEN είναι gaussian και η $\log L$ ΔEN είναι παραβολή
- Αποδεικνύεται ότι το 1σ του εκτιμητή βρίσκεται πάλι στο $\log L = \log L_{\max} - 0.5$



Ειδικές περιπτώσεις

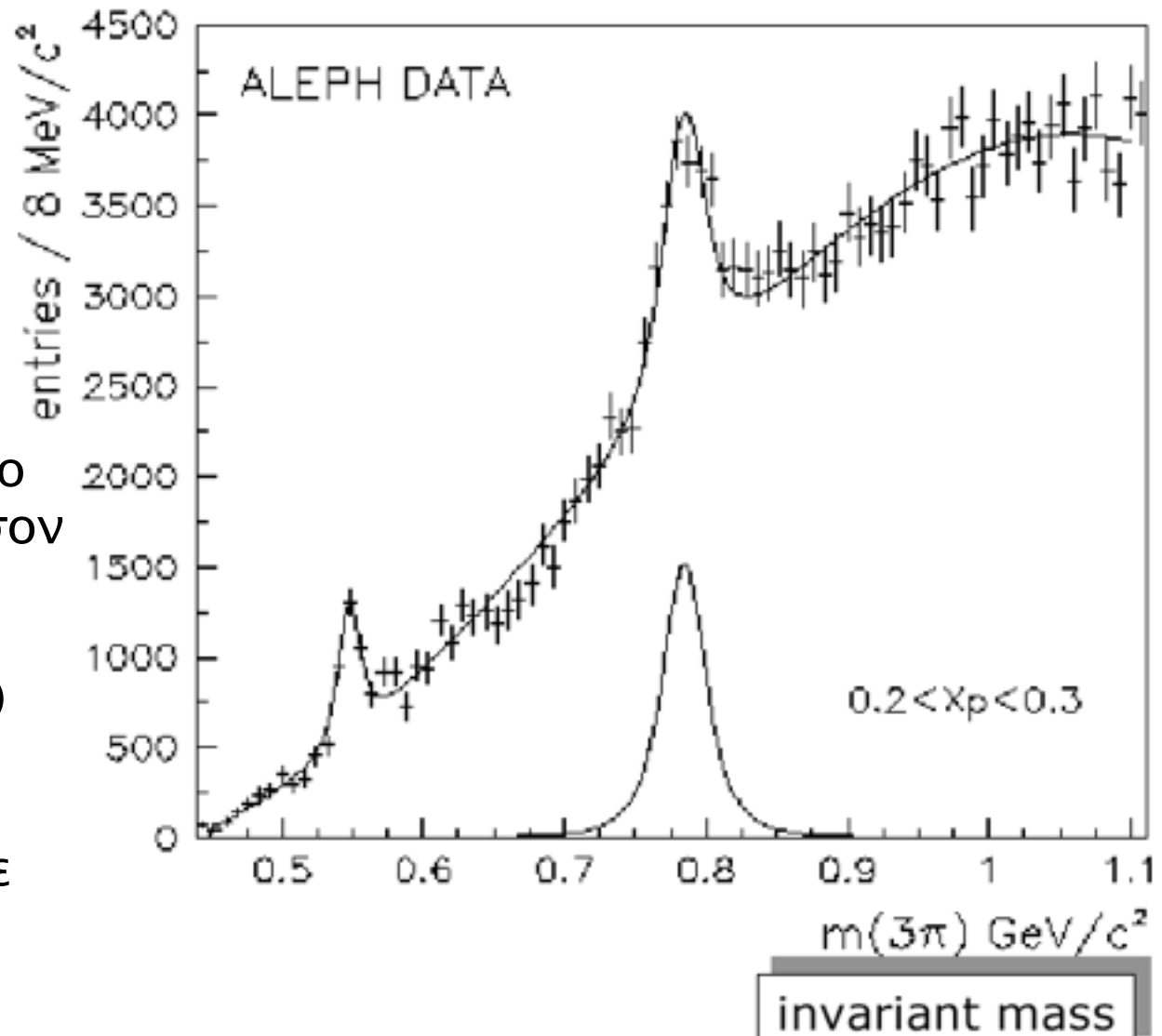
- Στα δεδομένα υπάρχει συνεισφορά από διάφορες πηγές. (Συνεισφορά από διάφορες pdfs) π.χ.
 - Στην αναζήτηση νέων σωματιδίων τα δεδομένα περιέχουν – μετά από επιλογή των γεγονότων– σήμα και υπόβαθρο
 - Στον υπολογισμό της μάζας ενός συντονισμού υπάρχει και συμβολή από υπόβαθρο λόγω πολλαπλών συνδυασμών



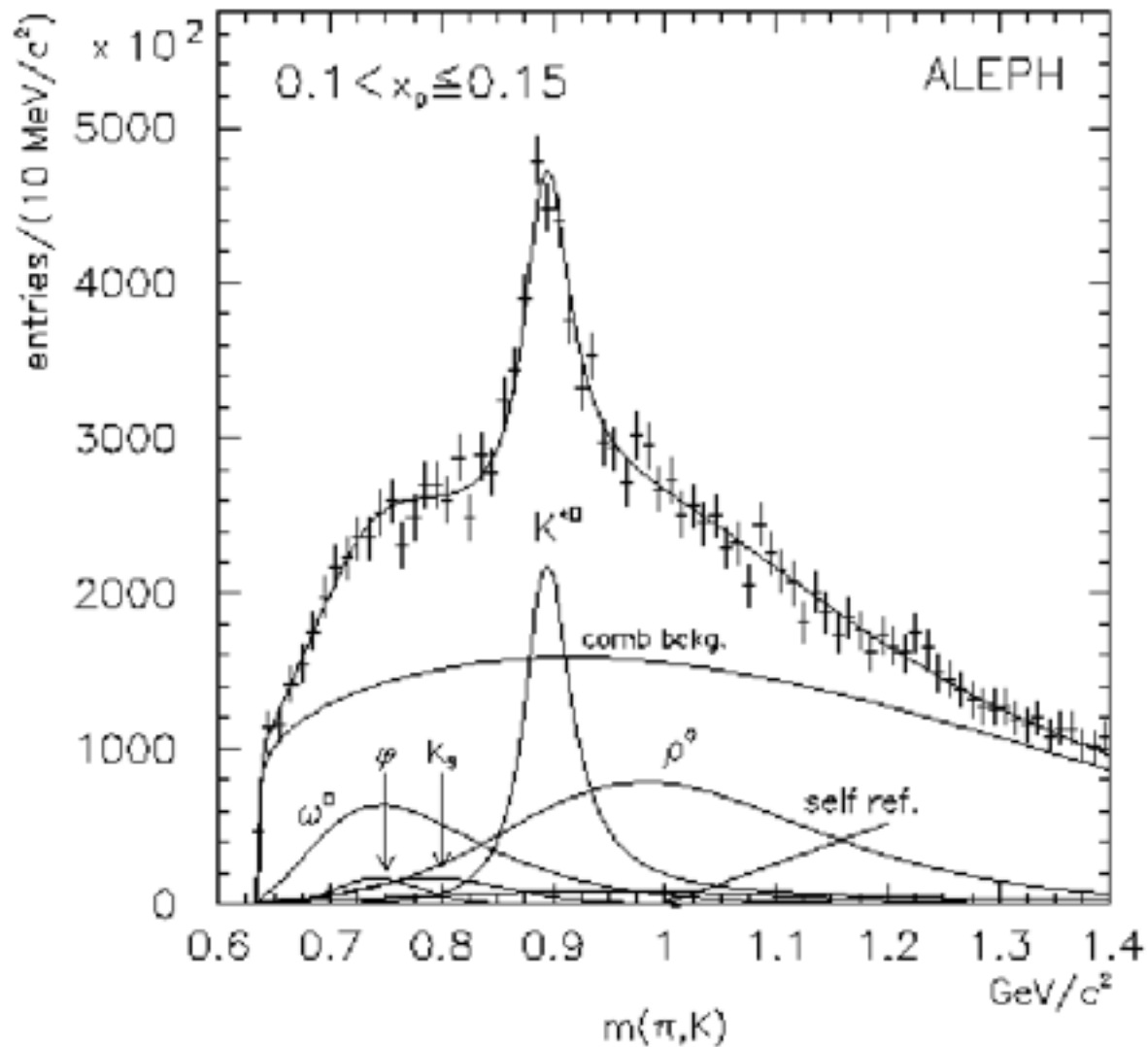
Ειδικές περιπτώσεις

- Παράδειγμα συντονισμού με υπόβαθρο από πολλαπλούς συνδυασμούς

Στα αδρονικά γεγονότα σε e^+e^- στο LEP έχουμε κατά μέσον όρο 20 φορτισμένα σωματίδια. Ο συντονισμός $\omega(782)$ που μπορεί να δημιουργηθεί διασπάται συχνά σε $\pi^+\pi^-\pi^0$.



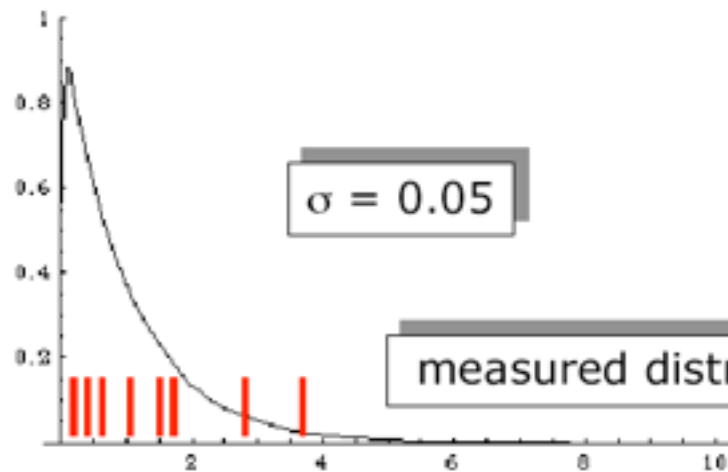
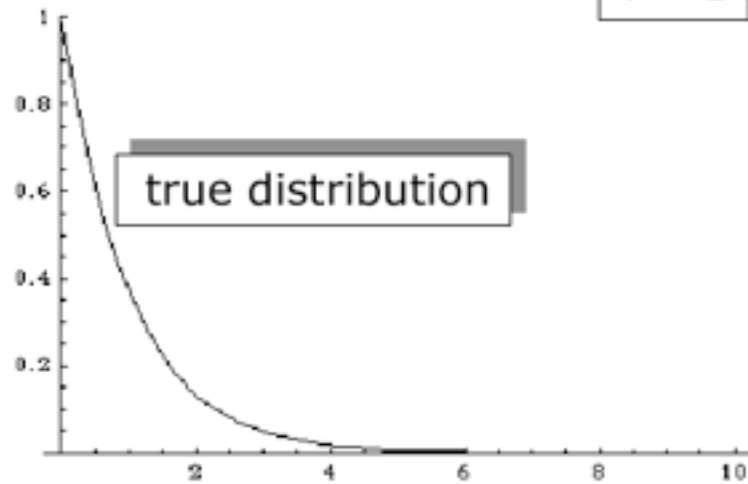
Ειδικές περιπτώσεις



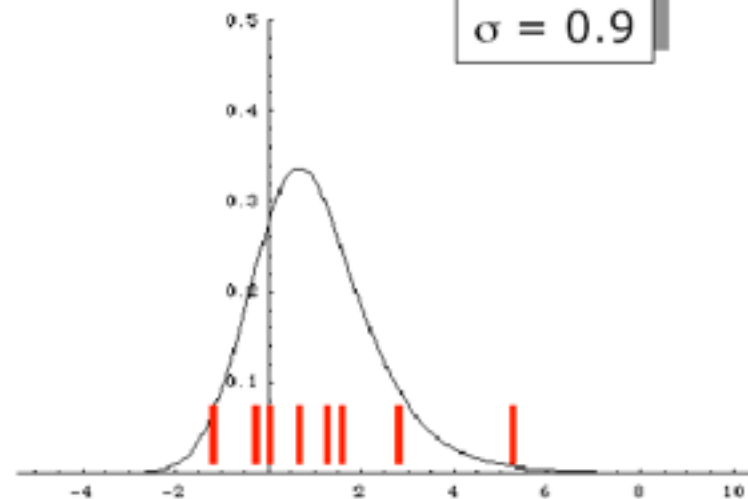
Ειδικές περιπτώσεις

- πως παίρνουμε υπόψη τον ανιχνευτή στα αποτελέσματα? (π.χ. **διακριτική ικανότητα**)
- Παράδειγμα : η μέτρηση του **μέσου μήκους διάσπασης** σωματιδίου γίνεται με πεπερασμένη ακρίβεια
- Η ‘αληθινή’ pdf είναι εκθετική με μέση τιμή τ
- Η διακριτική ικανότητα του ανιχνευτή είναι κατανομή gauss με μέση τιμή t_i (την μέτρηση) και σφάλμα σ
- Η κατανομή που μετρούμε είναι η συνέλιξη των δύο κατανομών

$$\tau = 1$$

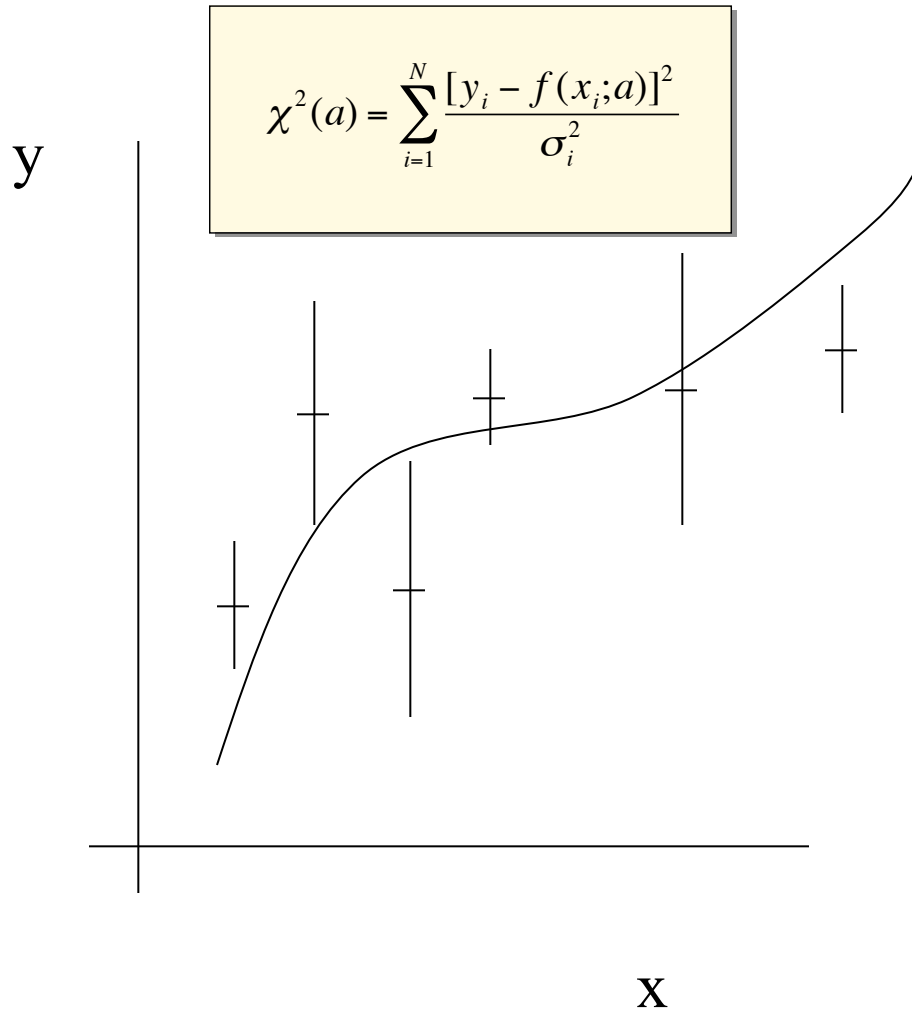


$$\sigma = 0.9$$



so: can measure negative values!
Take this into account when fitting for τ

Μεθοδος ελαχίστων τετραγώνων (least squares)

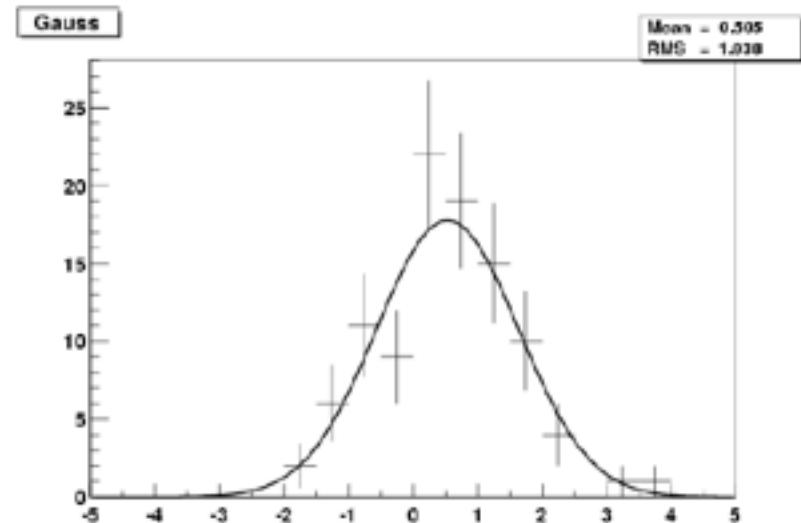
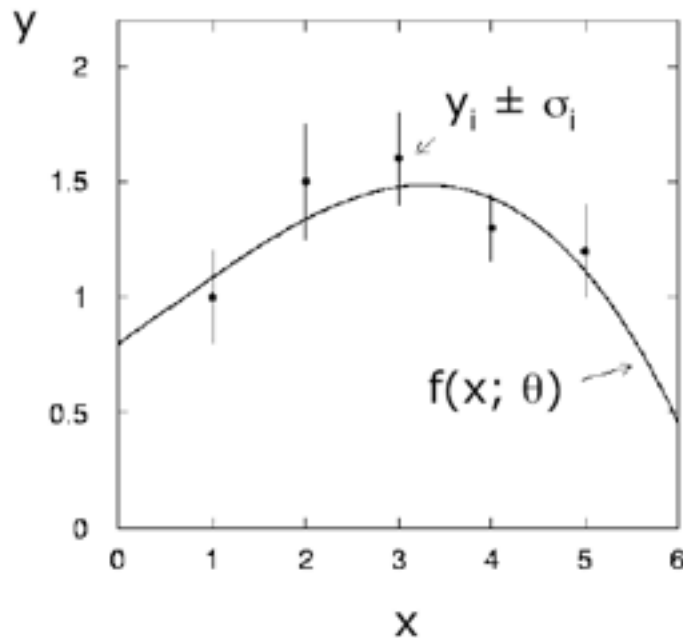


- Μέτρηση του y σε διάφορα x με σφάλμα σ και πρόβλεψη $f(x; a)$
- Η πιθανότητα:
$$\propto e^{-(y-f(x; a))^2 / 2\sigma^2}$$
- (y_i gaussian random variables)
- Η $\ln L(a)$ είναι:

$$\log L(a) = -\frac{1}{2} \sum_{i=1}^N \frac{[y_i - f(x_i; a)]^2}{\sigma_i^2}$$

- Maximize $\ln L \Rightarrow$ minimize χ^2

Μεθοδος ελαχίστων τετραγώνων- παραδείγματα



```
*****
*                               *
*   W E L C O M E  to  G A U S S   *
*                               *
*   Version:  0.02/00   21 January 2000  *
*   You are welcome to visit our Web site:  *
*   http://www.cran.ch   *
*                               *
*****
```

Compiled for Linux with shared support.

Copyright 2000 C/C++ International version 1.35.10, Jan 9
Type ? for help. Commands must be C++ statements.
Enclose multiple statements between { }.

Processing /usr/local/bin/gauss...

```
FCN=1.07334 DATA=1336AD  STATUS=COMPLETED  ADT CALLS  200 TOTAL
                                EDH=1.11130e-10  STRATEGY=1  ERRORS  MATRIX UNCERTAINTY  1.9 per cent
                                -----
EXT PARAMETER  NO.  NAME  VALUE  ERROR  SIZE  STEP  DERIVATIVE
-----
1  p0  1.77555e+01  2.41407e+00  4.07572e+04  7.42657e+05
2  p1  5.94870e+03  3.27126e+03  1.94122e+01  1.07716e+04
3  p2  1.11777e+00  1.51120e+01  -4.33746e+01  7.90081e+01
```

$$\chi^2 = \text{FCN}$$

- Definitions, Errors on parameters

Μεθοδος ελαχίστων τετραγώνων

Η ποιότητα της προσαρμογής

- Σύγκριση της προσαρμογής (fit) με τα δεδομένα
- Πρέπει το χ^2 της προσαρμογής: $\chi^2 \approx 1$ per data point οι βαθμοί ελευθερίας του συστήματος
ndof :

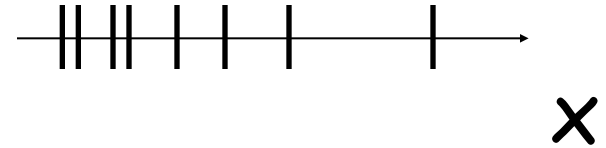
$$N_{\text{degrees Of Freedom}} = N_{\text{data points}} - N_{\text{parameters}}$$

- Αν τα δεδομένα σε bins:

$$N_{\text{degrees Of Freedom}} = N_{\text{data bins}} - N_{\text{parameters}}$$

- Αν $\chi^2 \gg \text{ndof}$: bad fit, bad theory, underestimated errors, bad luck!
- Αν $\chi^2 \ll \text{ndof}$: errors too big, good luck!

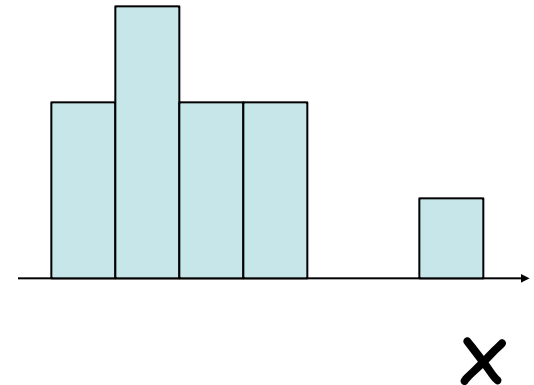
Fitting Histograms



- Συνήθως βάζουμε τα $\{x_i\}$ σε bins
- Οπότε τα data είναι $\{n_j\}$, n_j Poisson distributed

$$\text{mean } f(x_j) = P(x_j) \Delta x$$

- 4 ΤΕΧΝΙΚΕΣ:
 - Full ML
 - Binned ML
 - Proper χ^2
 - Simple χ^2



Τι μεγιστοποιούμε/ελαχιστοποιούμε

- Full ML

$$\ln L = \sum_i \ln P(x_i; a)$$

- Binned ML

$$\ln L = \sum_j \ln \text{Poisson}(n_j; f_j) \approx \sum_j n_j \ln f_j - f_j$$

- Proper χ^2

approx. by Poisson $\sigma^2 = \mu$)
$$\sum_j \frac{(n_j - f_j)^2}{f_j}$$
 (contents of each bin are

- Simple χ^2

number of entries observed
$$\sum_j \frac{(n_j - f_j)^2}{n_j}$$
 (approx. variance with

Τι χρησιμοποιούμε?

- **Full ML**: χρησιμοποιεί όλη την πληροφορία αλλά είναι πιο περίπλοκη. Την προτιμούμε για μικρό αριθμό γεγονότων
- **Binned ML**: λιγότερο περίπλοκη. Χάνουμε πληροφορία αν το bin size πολύ μεγάλο
- **Proper χ^2** : λιγότερο περίπλοκη και δίνει την ποιότητα του fit απ'ευθείας. Αν n_j large Poisson $\rightarrow$ Gauss
- **Simple χ^2** : n_j πρέπει να είναι μεγάλο

Τι χρησησιμοποιούμε?

- **Likelihood ratio** για να υπολογίσουμε πόσο : signal-like ή background-like είναι το πειραματικό μας αποτέλεσμα (π.χ. Higgs discovery) κατασκευάζουμε έναν estimator (estimator function) X
- Θεωρούμε ότι όταν ο estimator είναι “positive” είναι “more background-like” όταν είναι “negative” είναι “more signal-like”

$$X(x) = -2 \ln Q(x) = -2 \ln \frac{\mathcal{L}(x|s+b)}{\mathcal{L}(x|b)}$$

$\mathcal{L}(x|b)$ → Η Likelihood του x με την υπόθεση ότι μόνο το background μπορεί να μας δώσει x

- Το x είναι χαρακτηριστική ιδιότητα των δεδομένων (συνήθως το πλήθος των γεγονότων)

$$\mathcal{L}(x|b) = \prod_{i=1}^{\text{number of bins}} \frac{e^{-b_i} b_i^{x_i}}{x_i!}$$

Είναι γινόμενο Poisson πιθανοτήτων για όλα τα bins της κατανομής x

