

Constrained Curve Fitting

P. Lepore

+ P. Mackenzie

C. Morningstar

K. Hornbostel

J. Shigenitsu

C. Davies

B. Clark

See talks: M. Wingate
M. Gürtler

The Problem

An Example

M.C. Sim'n $\Rightarrow$ $\bar{G}(t)$ for $t=0,1,2,3$
eg

Theory $\Rightarrow G(t) = \sum_{n=1}^{\infty} A_n e^{-E_n t}$

24 pieces of
data



∞ number of
parameters to fit

How?

$\hookrightarrow$ systematic errors?

$\hookrightarrow$ high-precision!

Sample Data

- 840 γ ^{local-local} propagators at $\beta = 6.0$
- Fit: minimize $\chi^2(A_m, E_m)$ where

$$\chi^2 \equiv \sum_{t, t'} \Delta G(t) \sigma_G^{-2}(t, t') \Delta G(t')$$

and

$$\Delta G(t) \equiv \overline{G}(t) - G_{\text{fit}}(t; A_m, E_m)$$

M.C. $\nearrow$ $\nwarrow$ $\sum_m A_m e^{-E_m t}$

$$\sigma_G^2(t, t') \equiv \text{correl' matrix}$$

(from MC)



Sol'n 1

Fit all t with $G(t) = \sum_{n=1}^M A_n e^{-E_n t}$

- increase M to get good fit
- Where do you stop?

M too small $\Rightarrow E_n$'s ... biased

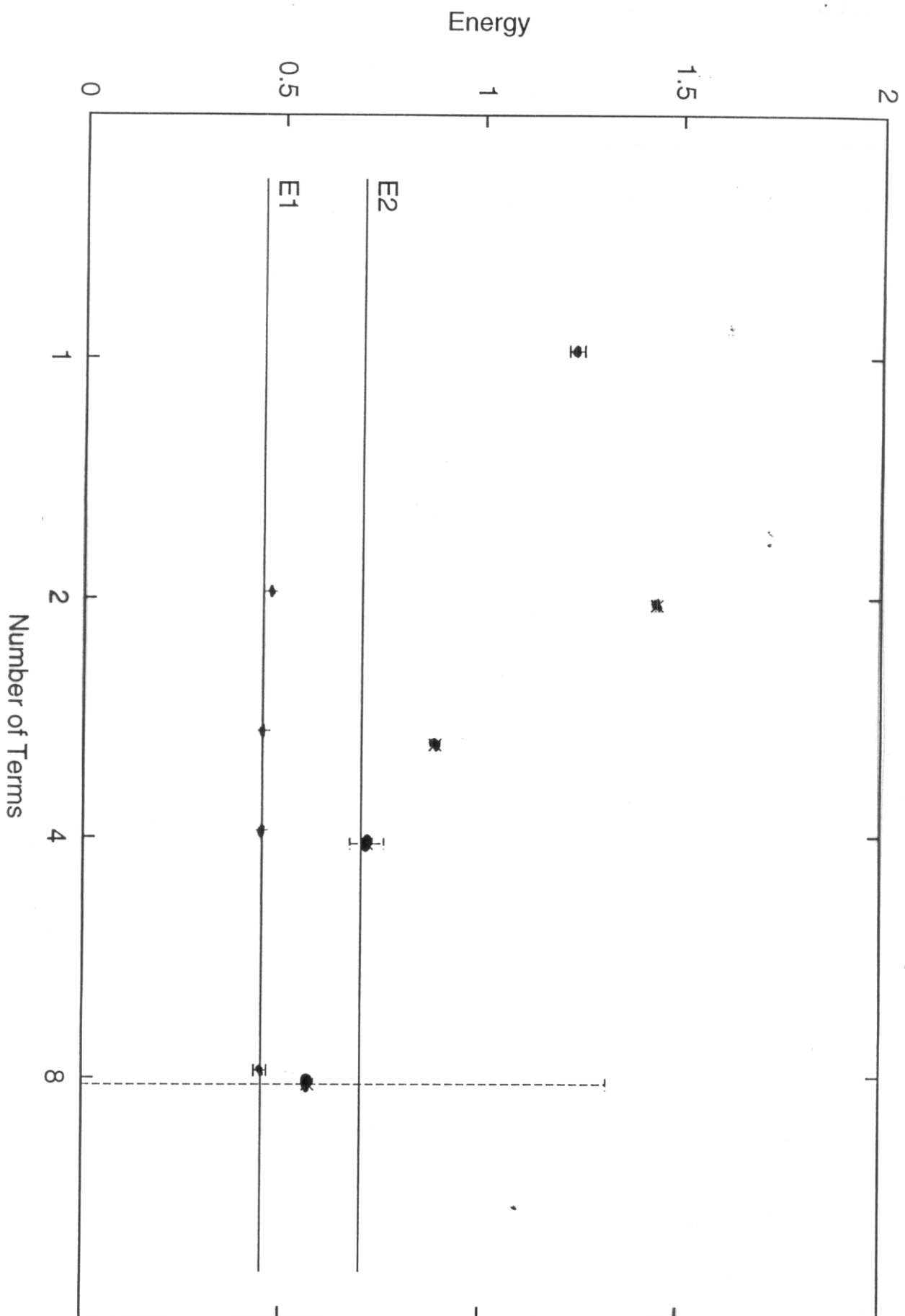
M too large $\Rightarrow \sigma_{E_n}$'s ... too large



$\rightarrow \infty$ if

$2M > \# \text{ data points.}$

E1, E2 vs Number of Terms (unconstrained)



Sol'n 2 (Standard)

Truncate data ($t \geq t_{\min}$)

and theory ($M = 2$ terms)

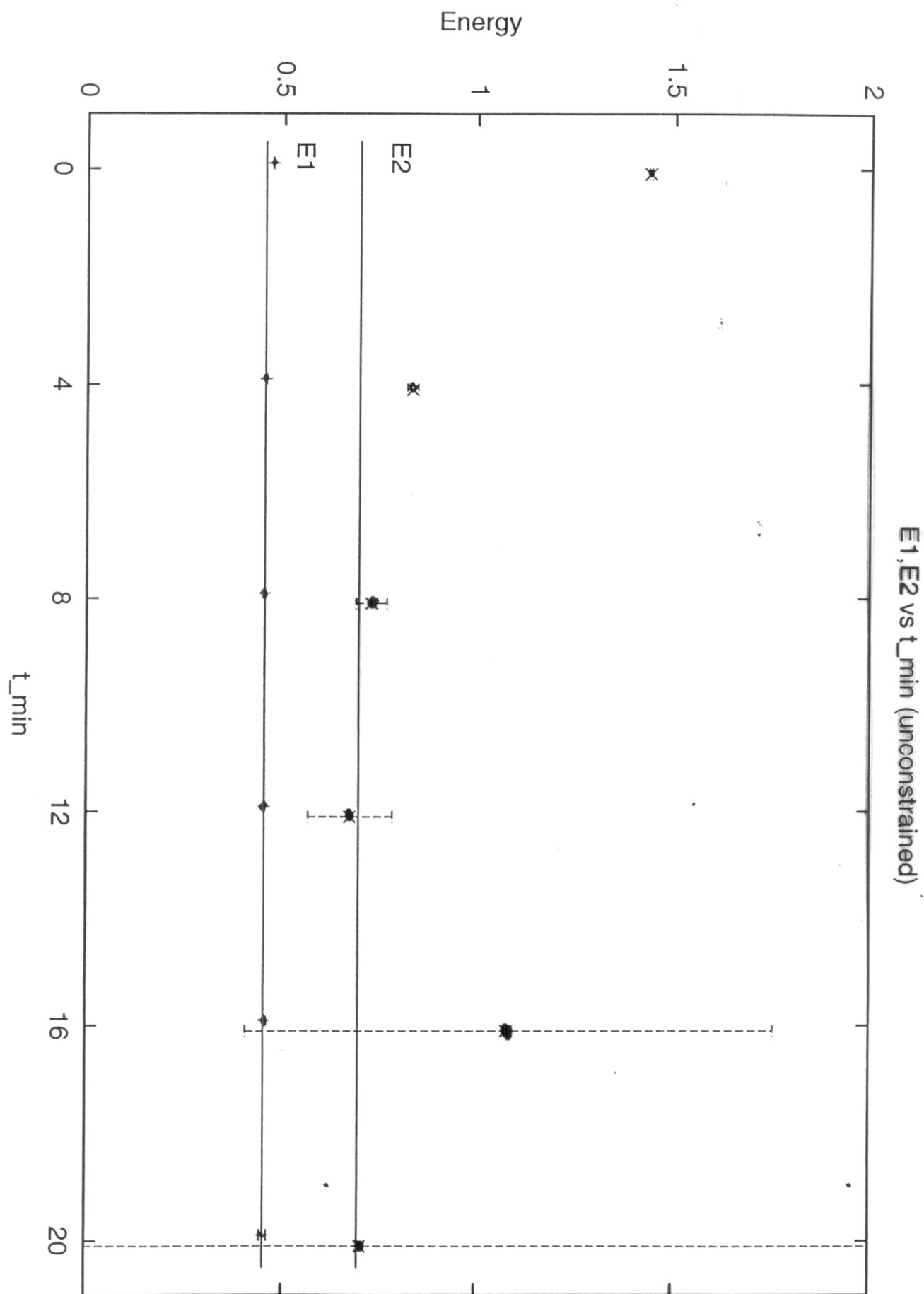
- Increase $t_{\min}$ to get good fit
- Where do you stop?

$t_{\min}$ too small $\Rightarrow E_m$'s ... biased

$t_{\min}$ too large $\Rightarrow \sigma_{E_m}$'s ... too large



usually truncate data until stat. errors large enough to swamp sys. error due to truncating theory (hopefully).



Better Solution

- Want to fit $\forall$ data ($t_{\min}=0$)
with $\sum^M A_m e^{-E_m t}$ for $M \rightarrow \infty$.

Problem: e.g., $M=8$

$$\Rightarrow A_4 \sim 5 \text{ to } 10 \times A_1$$

$$\text{But Physics} \Rightarrow A_4 \lesssim A_1$$

Need to teach Physics to
fitting software!

Constrained Fit

Replace $\chi^2 \rightarrow$ augmented χ_a^2 :

$$\chi_a^2 \equiv \chi^2 + \chi_p^2$$

where

$$\chi_p^2 \equiv \sum_m \frac{(A_m - \tilde{A}_m)^2}{\tilde{\sigma}_{A_m}^2} + \sum_m \frac{(E_m - \tilde{E}_m)^2}{\tilde{\sigma}_{E_m}^2}$$

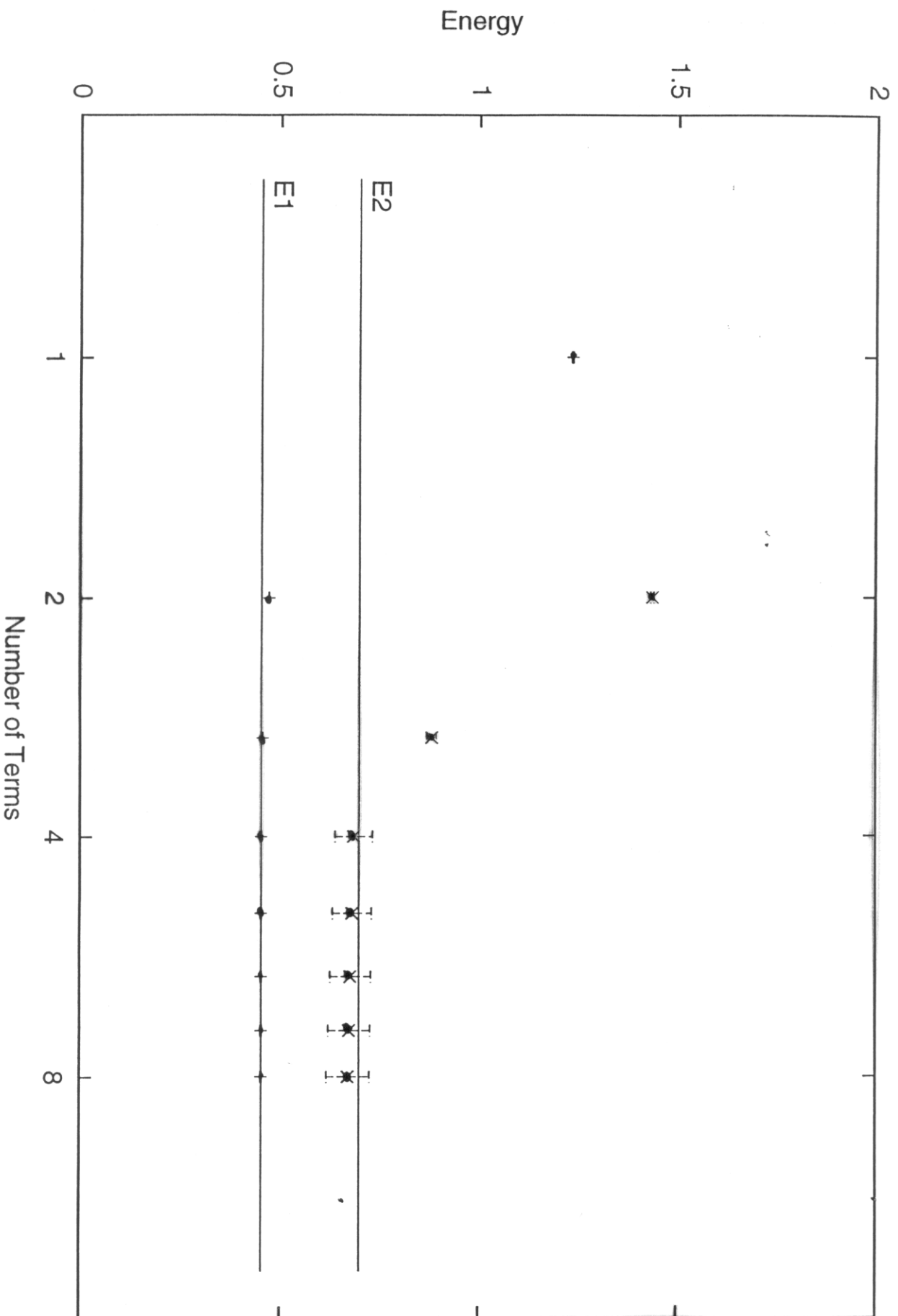
Favors A_m in interval $\tilde{A}_m \pm \tilde{\sigma}_{A_m}$
 E_m " " $\tilde{E}_m \pm \tilde{\sigma}_{E_m}$

$\tilde{A}_m, \tilde{\sigma}_{A_m} \dots$ are inputs;

choose reasonable values
on basis of prior knowledge

"PRIORS"

E1, E2 vs Number of Terms (constrained)



Method:

Increase $M = \# \text{ terms}$ until
answers (and errors) converge.

N.B.

- Can't make M too large;
c.f. $t_{\min}$ which must be optimized!

- Errors = statistical due to M.C.
+ systematic due to poorly
constrained high- E
states.

→ assumptions are
explicit & intelligible

- Constraints $\Rightarrow$ better numerics:
e.g., can fit 100 term G

~~NB~~

Convergence (as M increases)
is key criterion.

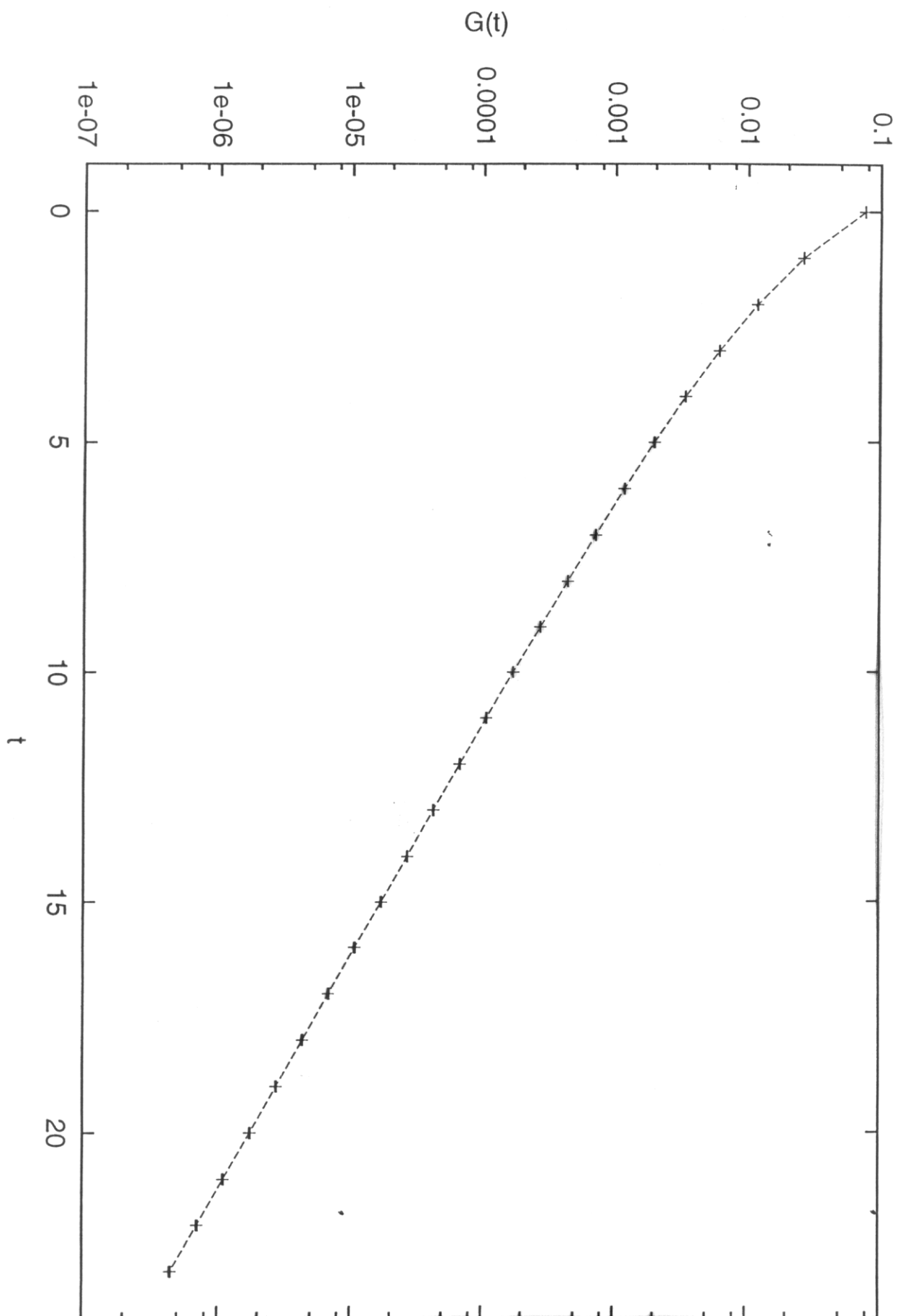
χ^2 , Q ... play secondary role

Fits here have

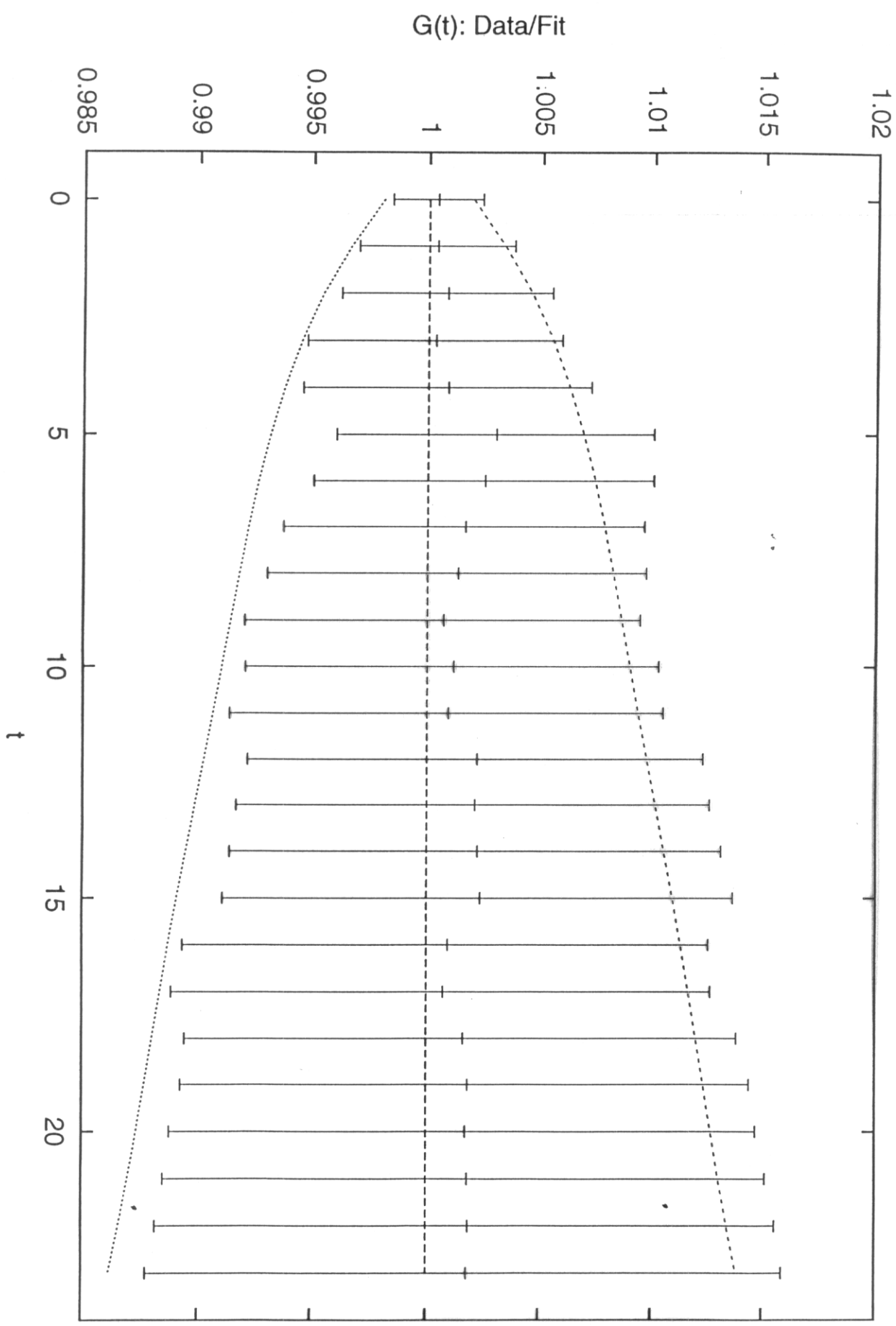
$$\frac{\chi^2_{a, \min}}{\# \text{ Data Points}} \sim 0.7 - 0.8$$

↑ correct ratio to monitor.

$G(t)$ vs t (constrained fit -- 5 terms)



G(t): Data/Fit vs t (constrained fit -- 5 terms)



B

Actual fit parameters :

$$a_m \equiv \log A_m \Rightarrow A_m > 0$$

$$E_m \equiv \log(E_m - E_{m-1}) \Rightarrow E_m > E_{m-1}$$

Priors :

$$\left. \begin{array}{l} \tilde{a}_m = \log 0.02 \\ \tilde{\sigma}_{a_m} = 2 \end{array} \right\} \Rightarrow A_m = 0.02 \begin{array}{l} +0.02 \\ -0.01 \end{array}$$

$$\left. \begin{array}{l} \tilde{E}_m = \log 0.2 \\ \tilde{\sigma}_{E_m} = 2 \end{array} \right\} \Rightarrow E_m - E_{m-1} = 0.2 \begin{array}{l} +0.2 \\ -0.1 \end{array}$$

- Results insensitive to precise values of priors (within reason)?

Double $\tilde{\sigma}_{am}$ & $\tilde{\sigma}_{Em}$ $\Rightarrow$

$$E_1 = 0.4526 (15) \rightarrow 0.4528 (14)$$

$$E_2 = 0.683 (49) \rightarrow 0.697 (50)$$

$$E_3 = 1.05 (12) \rightarrow 1.10 (21)$$

$\Rightarrow E_1, E_2$ largely determined
by the M.C. data

E_3 strongly affected
by priors

13. Other methods are special cases:

$t_{\min} > 0$ method $\Rightarrow$

$$\tilde{\sigma}_{A_m}, \tilde{\sigma}_{E_m} = \begin{cases} \infty & \text{for 1st terms} \\ 0 & \text{for all rest} \end{cases}$$

$\hookrightarrow$ forcing $A_m = 0$

$t_{\min} = 0$ $M \rightarrow \infty$ method $\Rightarrow$

$$\tilde{\sigma}_{A_m}, \tilde{\sigma}_{E_m} = \infty \text{ for } \forall \text{ terms}$$

$\Rightarrow \infty$ fit errors on all
 A_m 's & E_m 's (as $M \rightarrow \infty$)

$\hookrightarrow$ correct answer if really don't
have any (prior) idea of
sizes of A_m 's & E_m 's.

Lots of other fits

→ Matrix G 's up to 4×3 (30-40-... parameters)

→ simultaneous fits of multiple channels (eg, π & ρ)

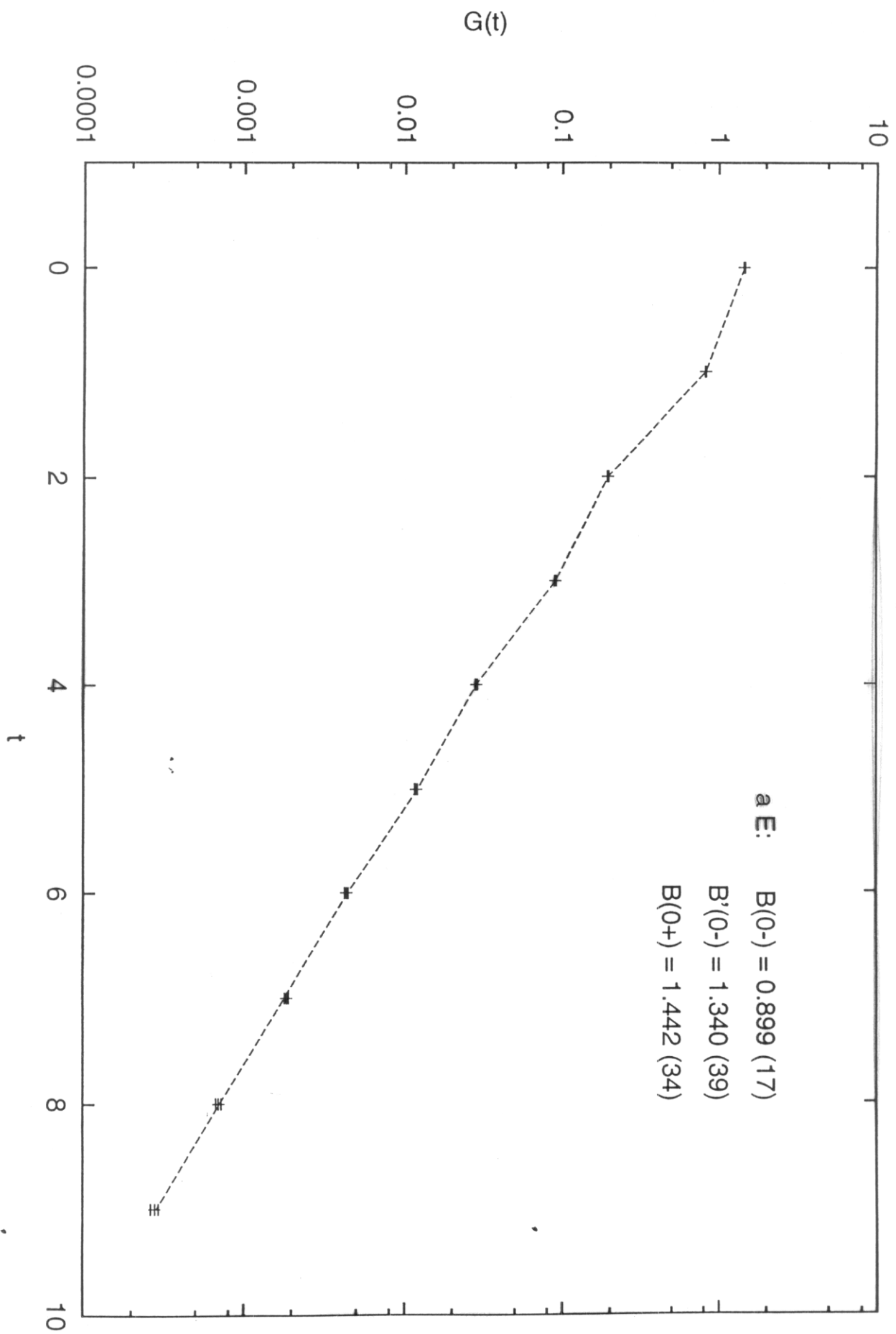
→ $V(r)$, glueballs

→ staggered quarks

‘
.
.

→ All worked well

$G(t)$ vs t (staggered-quark B: 3 even + 3 odd terms)



Theory & Interpretation

Interpretation : Bayesian Statistics

→ Calculus of probabilities
provides framework

Not'n :

$P(A|B)$ = probability of A
given B
(i.e. assuming B true)

θ_i = $\{A_i, E, \dots\}$ = parameters

η_i = $\{\tilde{\sigma}_{A_i}, \tilde{A}_i, \dots\}$ = priors

$\bar{G}$ = M.C. data average

Assumption #1:

M.C. data set sufficiently large
that $\bar{G}$ has Gaussian statistics
(Central Limit Thm)

$\Rightarrow \left\{ \begin{array}{l} P(\bar{G} | \mathbf{g}) \equiv \text{prob. of a particular } \bar{G} \\ \text{given a theory with} \\ \text{param's } \mathbf{g} \end{array} \right.$

$$\propto e^{-\chi^2(\mathbf{g})/2}$$

$$\chi^2 \equiv \sum_t \frac{(\bar{G} - G_{\text{th}}(\mathbf{g}))^2}{\sigma_G^2}$$

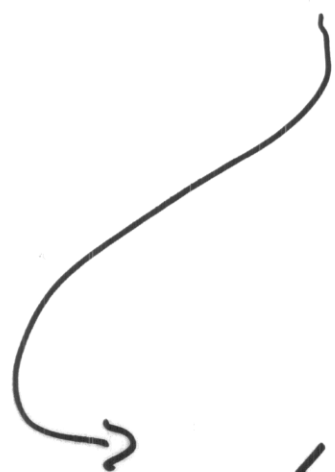
$\bar{G}, \sigma_G$ from M.C.

⚠ N.B. σ_G actually $\mathbf{g}$ -dependent ⚡

Goal: Don't want $P(\bar{G} | \mathcal{P})$!

Need

$P(\mathcal{P} | \bar{G})$ = Prob. of particular parameter values $\mathcal{P}$ given $\bar{G}$



$$\langle \mathcal{P}_i \rangle \equiv \int \pi d\mathcal{P} P(\mathcal{P} | \bar{G}) \mathcal{P}_i$$

= ~~estimate~~ for param. $\mathcal{P}_i$

$$\sigma_{\mathcal{P}_i}^2 \equiv \langle (\mathcal{P}_i - \langle \mathcal{P}_i \rangle)^2 \rangle$$

= ~~variance of estimate~~

$$\langle f(\mathcal{P}) \rangle = \int \pi d\mathcal{P} P(\mathcal{P} | \bar{G}) f(\mathcal{P})$$

= ~~estimate for fcn~~ $f(\mathcal{P})$

...

Bayes' th =

Probab. Product Rule $\Rightarrow$

$$\begin{aligned} P(\bar{G} | \mathcal{P}) &= P(\bar{G} | \mathcal{P}) P(\mathcal{P}) \\ &= P(\mathcal{P} | \bar{G}) P(\bar{G}) \end{aligned}$$

$\Rightarrow$

$$\begin{aligned} P(\mathcal{P} | \bar{G}) &= \frac{P(\mathcal{P}) P(\bar{G} | \mathcal{P})}{P(\bar{G})} \xrightarrow{\text{indep of } \mathcal{P}} \\ &\propto P(\mathcal{P}) P(\bar{G} | \mathcal{P}) \end{aligned}$$

$\rightarrow P(\mathcal{P}) =$ probab. $\mathcal{P}$ is correct given no $\bar{G}$
"Prior" $\swarrow$
 $=$ Prior Assumptions

Assumption #2:

Prior dist'n $P(\theta)$ can be approximated by a Gaussian:

$$P(\theta) = e^{-\chi_P^2(\theta)/2}$$

where

$$\chi_P^2(\theta) = \sum_i \frac{(\theta_i - \tilde{\theta}_i)^2}{\tilde{\sigma}_{\theta_i}^2}$$

$\Rightarrow$ a priori assumption is
that $\theta_i \sim \tilde{\theta}_i \pm \tilde{\sigma}_{\theta_i}$

$$\Rightarrow \left\{ \begin{aligned} P(\theta | \bar{G}) &\propto e^{-\chi^2/2} e^{-\chi_P^2/2} \\ &\propto e^{-\chi_a^2(\theta)/2} \end{aligned} \right.$$

$\rightarrow$ augmented χ^2 from before.

Why Gaussian's for $P(p)$?

- Least biased choice for $P(p)$ minimizes information content
 $\Rightarrow$ maximizes entropy

$$S \equiv \int -P(p) \log P(p) dp$$

- Given $\langle p \rangle = \tilde{p}$, $\langle (p - \langle p \rangle)^2 \rangle = \tilde{\sigma}^2$

compute

$$\frac{\delta}{\delta P} \left[S + \lambda \int P + \mu \int P p + \nu \int P p^2 \right]$$

$$= 0$$

$$\Rightarrow \left\{ P(p) = \frac{1}{\sqrt{2\pi} \tilde{\sigma}} e^{-\frac{(p - \tilde{p})^2}{2\tilde{\sigma}^2}} \right.$$

↪ "Least Informative Prior"

Error Estimation

Best Fit

- From integrals:

$$\langle p_i \rangle = \frac{\int \pi dp \, e^{-\chi_a^2/2} p_i}{\int \pi dp \, e^{-\chi_a^2/2}}$$

$$\sigma_{p_i}^2 = \langle (p_i - \langle p_i \rangle)^2 \rangle \dots$$

- Evaluate using

1) Vegas

2) M.C. - Metropolis, HMC, HMD...

↳ $\chi_a^2 \leftrightarrow \text{Action}$
 $p_i \leftrightarrow \text{Field}$

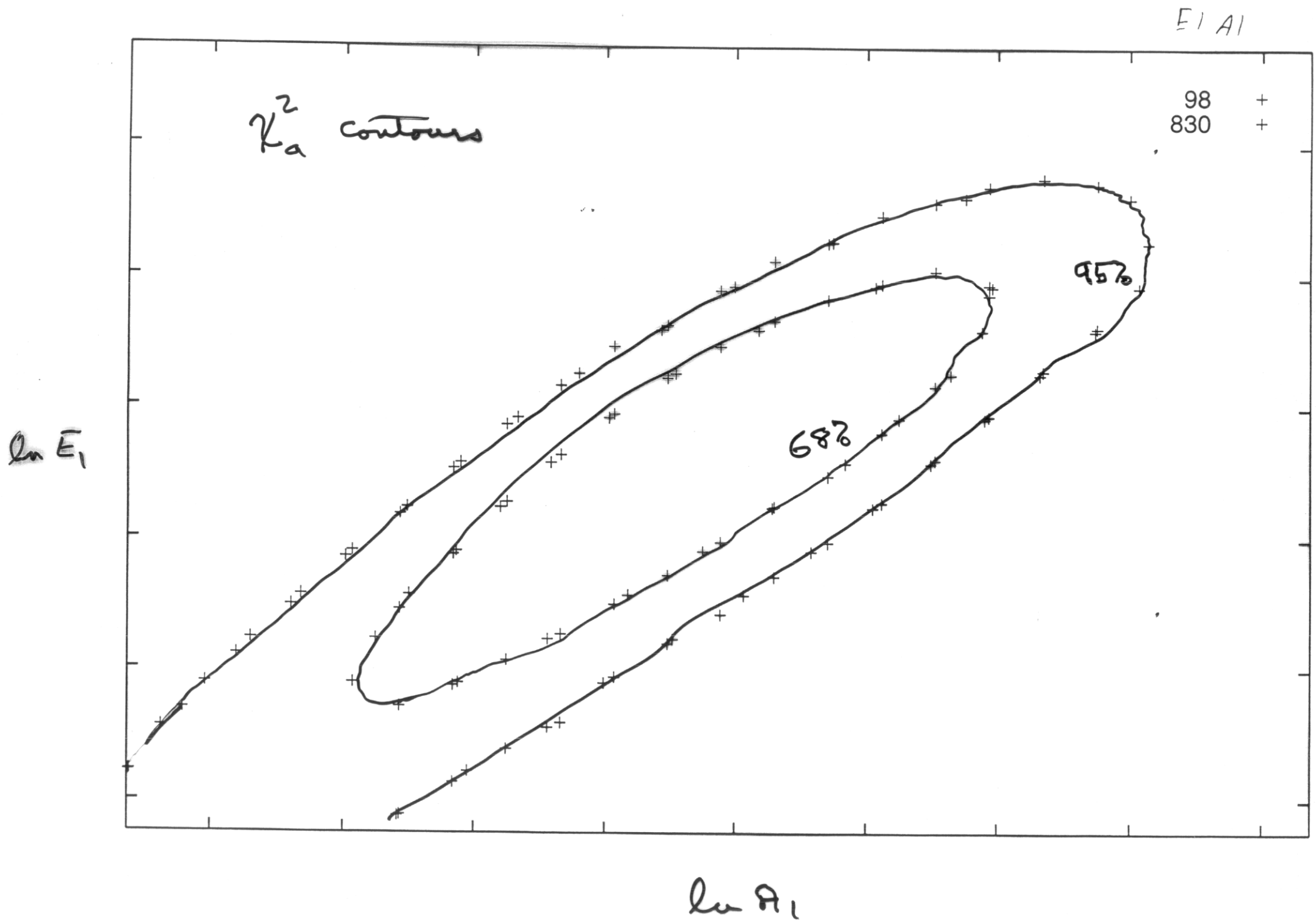
N.B.

Integrals can be extremely
difficult & costly

→ integrand extremely
sharply peaked

→ has long, narrow,
high wandering
ridges $\Rightarrow$ huge
auto correlation times!

$\Rightarrow$ approx's useful.



E2A2

χ^2_{a} contours

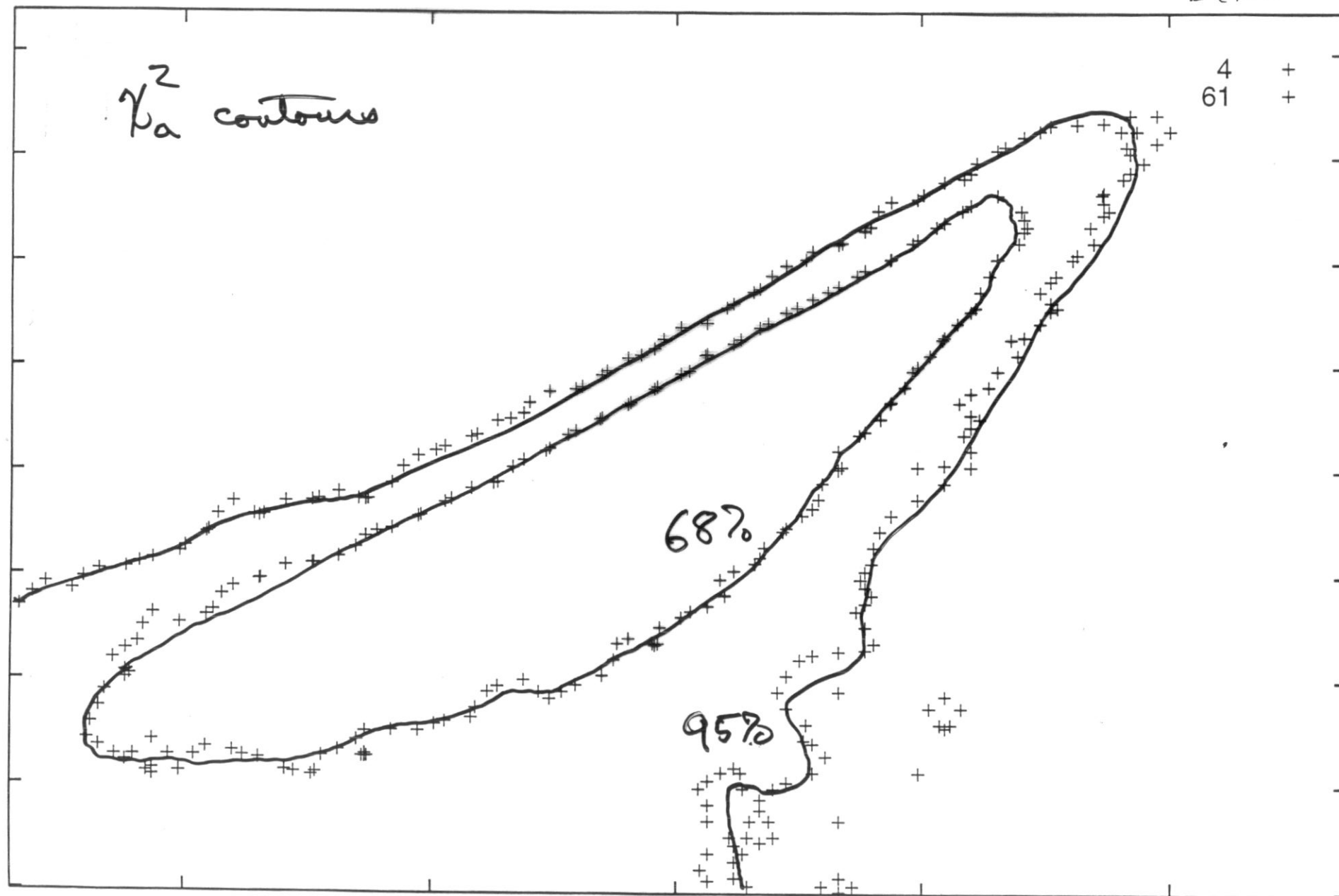
4 +
61 +

ln E₂

68%

95%

ln A₂



Gaussian Approx (High statistics)

$$\chi_a^2(p) \approx \chi_{a,\min}^2 + \frac{1}{2} (p - p^*)_i (p - p^*)_j \partial_i \partial_j \chi_a^2$$

$$\Rightarrow \langle p_i \rangle \approx p_i^*$$
$$\sigma_{p_i}^2 \approx C_{ii} \quad \text{where } C_{ij}^{-1} = \frac{\partial_i \partial_j \chi_a^2}{2}$$

...

$\Rightarrow$ constrained curve fitting
based on χ_a^2 minimization
as in 1st part of talk.

Bootstrap Approx

Modified Bootstrap Procedure:

1) Replace $\bar{G} \rightarrow \bar{G}' = \text{avg of bootstrap copy of original } G's$

$\rightarrow$ Replace $\tilde{p}_i \rightarrow \tilde{p}_i' = \text{Gauss. random \# from dist'n with mean } \tilde{p}_i \text{ \& s-dev } \tilde{\sigma}_{p_i}$

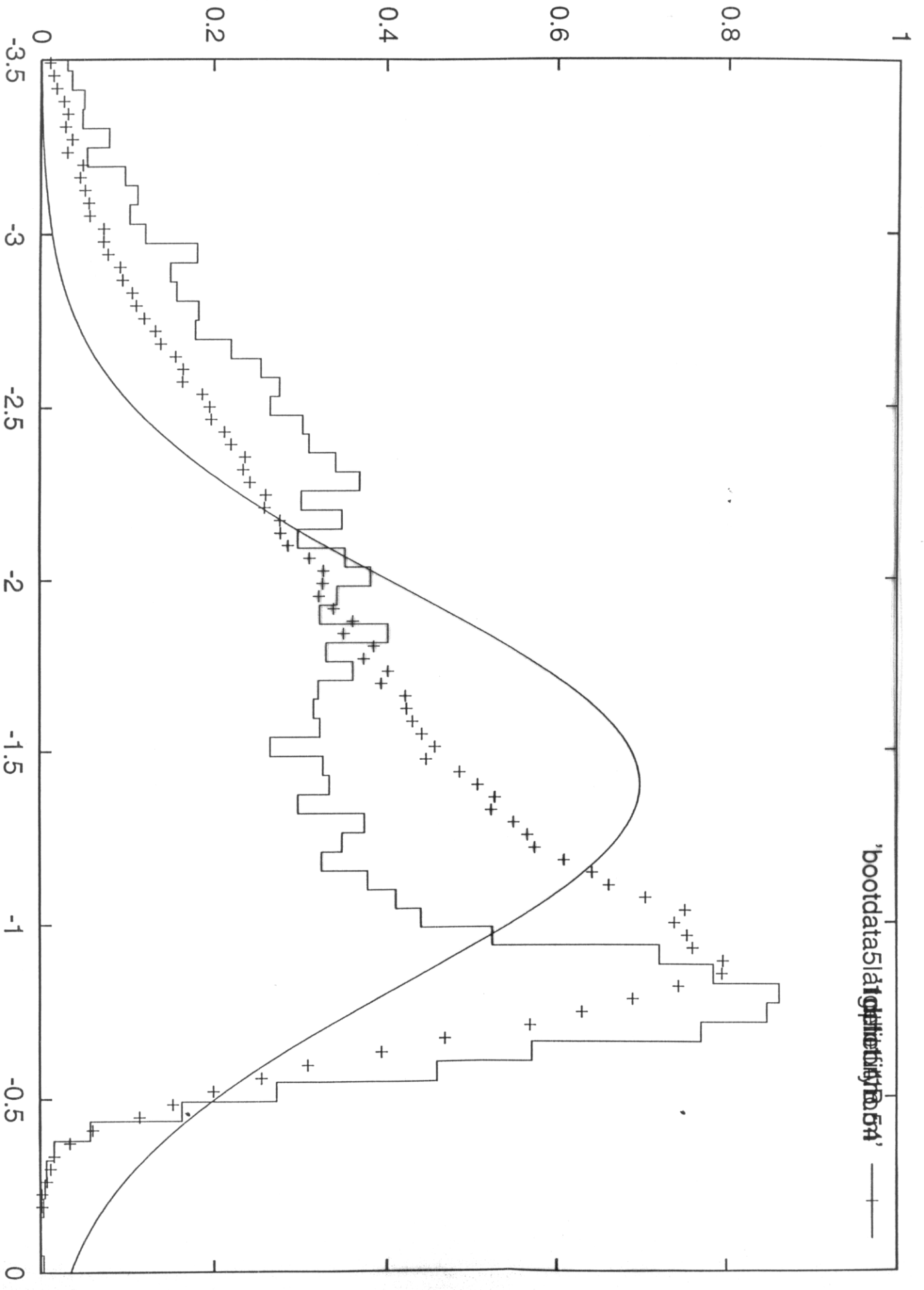
2) Compute $p_i^* = \min \text{ of } \chi^2_a$
using $\bar{G}', \tilde{p}_i'$

3) Repeat steps 1 & 2 100-1000 times.

$\Rightarrow$ ensemble of p_i^* 's

$\Rightarrow$ avg's on p_i^* $\approx$ Bayes avg's on p_i

$\rightarrow$ small add'l corr'n
but usually not important (?)



Other Applications

Pert'n Theory from High- β Simulations

Trotter GPL Mackenzie
Shobarene

- Measure $W(m, n)$ for 9 lattice spacings $\Rightarrow \alpha_p(3.4/a) = 0.01 - 0.07$.

- Fit M.C. values $-\frac{\log W(m, n)}{2m + 2n}$

$$\overline{C_0} \sum_{n=1}^{M \geq 5} C_n \alpha_p^n(g^*) + \text{prior} \quad \begin{matrix} \tilde{C}_n = 0 \\ \tilde{\sigma}_{C_n} = 5 \end{matrix}$$

- $W(5, 5)$ Results (e.g.):

$$C_1 = 1.7693 (6)$$

$$C_2 = -1.201 (40)$$

$$C_3 = 4.3 (7)$$

Exact

$$1.7690$$

$$-1.177$$

?

$\rightarrow$ Fix C_1, C_2 to exact val
 $\Rightarrow C_3 = 3.91 (23)$

N.B.

Some c_n 's fixed by MC data
($n \leq 3$) ; others by priors.

As statistics improve more c_n 's
fixed by data $\Rightarrow$

order of fit increases
automatically as
warranted by statistics

$\hookrightarrow$ i.e. never agonizing
about linear vs.
quadratic fit!
Set $M = 10$ or 20
and forget about order!

Extrapolations in Lattice Size L Trotter

Same program for Polyakov line

$$\Rightarrow \text{static-quench energy } E_0 = \sum_n c_n \alpha_p^n(g^*)$$

$$\Rightarrow \text{M.C. values for } c_3(L) \text{ for } L=3, \dots, 11$$

Fit to

$$c_3(L) = c_3 + p_1 \frac{a_n^2 L^2}{L} + p_2 \frac{a_n L^2}{L} + \frac{p_3}{L} \\ + \frac{p_4}{L^2} + \frac{p_5}{L^3} + \dots$$

$$+ \text{ priors: } \tilde{\sigma}_{c_3} = \tilde{\sigma}_{p_i} \approx 6$$

and extrapolate to $L = \infty$.

Result : $c_3 = 3.8 \pm 1.1$

c.f. Burgio et al $\Rightarrow 3.2 \pm 0.6$

~~N.B.~~

$$E_0 = 1.0701 \alpha_p(q^*) + 0.117 \alpha_p^2 \\ + 3.3(5) \alpha_p^3 \\ + \dots$$

57

$$E_0^{\text{unimp}} = 2.1173 \alpha_{\text{ext}} + 11.152 \alpha_{\text{ext}}^2 \\ + 86.3(5) \alpha_{\text{ext}}^3 \\ + \dots$$

N.B.

Analogous to chiral
extrapolation

$\Rightarrow$ "all orders" extrapolations
in m_b possible!

Variations

Empirical Bayes - Tuning Priors

$P(\bar{G} | \eta)$ = probab of $\bar{G}$ given $\{\eta_i\}$ priors
↓

$$\propto \int \pi \frac{d\beta}{\sqrt{2\pi} \tilde{\sigma}_\beta} e^{-\eta^2/2}$$

$$\rightarrow \left[\pi \frac{1}{\tilde{\sigma}_{\eta_i}} \right] \sqrt{\det C_{ii}} e^{-\eta_{a, \min}^2/2}$$

in Gaussian approx

$\Rightarrow$ vary η_i to increase $P(\bar{G} | \eta)$

→ using $\bar{G}$ data to suggest reasonable values for η_i

→ OK for 1 or 2 parameters
but { complete nonsense if every η_i optimized }

Non-Parametric Fits - Max. Ent.

see Oevers et al hep-lat/0009031

Yamazaki et al " /0105030

Eg/

Fit

$\bar{G}$

to

spectral
fcn

$$G(t) \equiv \int d\omega e^{-\omega t} p(\omega)$$

- Fit parameters:

$$p(\omega_j) \quad \text{for } \omega_j = 0.01 \times j$$
$$\& \quad j = 0, 1, 2, \dots, 750$$

2) "non parametric"

2) # fit parameters $\gg$ # data points

$\Rightarrow$ Prior is essential

2) $p(\omega)$ has peaks at b.d.
state energies

- Max Ent $\Rightarrow$ useful framework for
designing priors

- Properties of $p(\omega)$ (source = sinh.)
 - 1) $p(\omega)$ has random noise (M.C.)
 - 2) $p(\omega) > 0$
 - 3) $p(\omega) d\omega =$ physically interesting (additive)

$\Rightarrow$ think of $p(\omega)$ as a probability density

- Most probable probab. density is one with largest entropy

" $\Rightarrow$ "

$$P(p) \propto e^{\alpha S(p)}$$

$\alpha = \text{input parameter}$

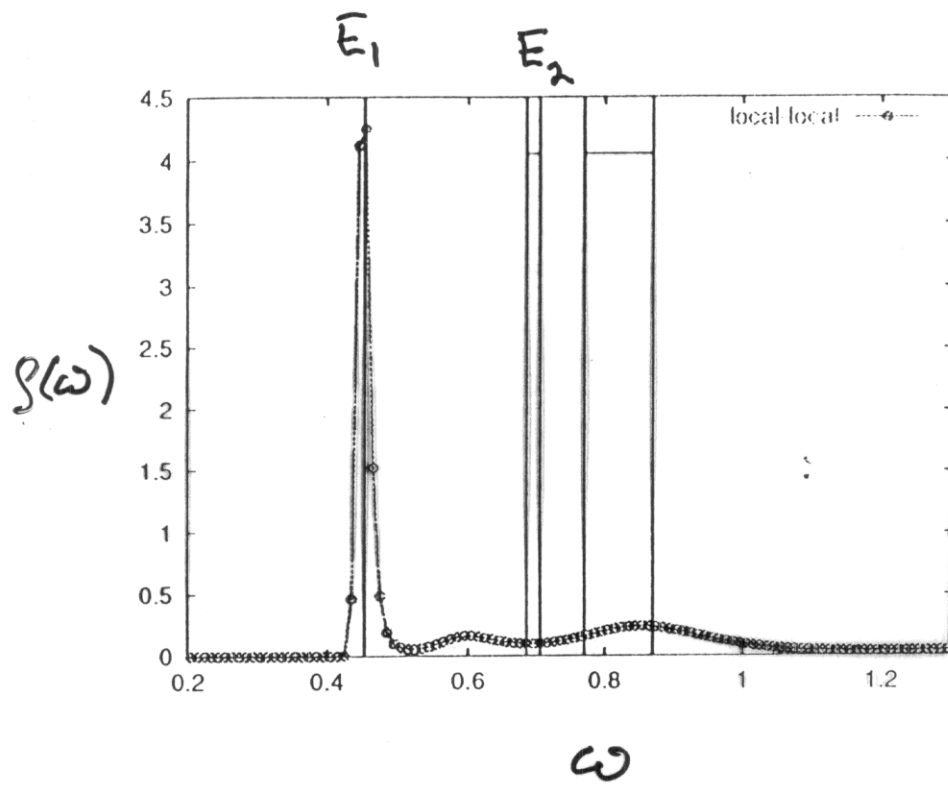
where

$$S \equiv \int d\omega \left[p(\omega) - m(\omega) - p(\omega) \ln \frac{p}{m} \right]$$

$p = m$ if no data.

$\hookrightarrow m(\omega) = \text{input.}$

$\Rightarrow$ constrained fit gives $p(\omega_i) \neq p_j$.



Oevers et al

↳ same τ data
as earlier

N.B.

- MaxEnt much less accurate than constrained fits to $G = \sum A_n e^{-E_n^t}$
 - uses much less prior information, ~~but~~ should use info if have it.
- Useful when there is no simple parametrization — e.g., in image processing
 - ↳ many variations on $e^{\alpha s}$ used in applications

Conclusions

- Constrained / Bayes fits: major improvement
 - fits all data $\Rightarrow$ more info, smaller σ 's
- ***
 - includes sys. errors from
as of unconstrained terms
 - can combine info from other
sim's via priors

- useful for
 - correlator analysis
 - Part'n the fits to M.C.
 - Extrapolations in m_π or L or a

 $\hookrightarrow$ order of extrapol'n
increases automatically
as warranted by increased
statistics.

- Easy & Fun for Lattice experts.

Books :

D. Sivia "Data Analysis - A Bayesian Tutorial" (easy)

B. Carlin "Bayes & Empirical Bayes Methods for Data Analysis" (less easy)