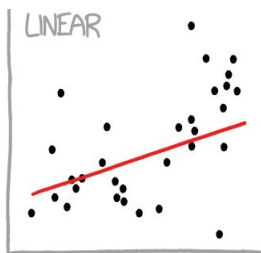


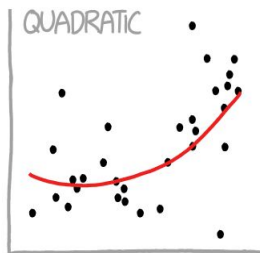
Validation & Evaluation in Physics:
How we do statistics in HEP, how we need to
change it for ML purposes
What we don't know and what we do wrong

Lydia Brenner

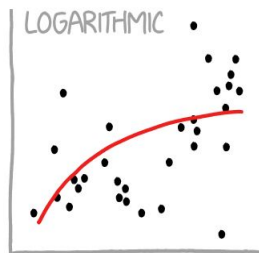
CURVE-FITTING METHODS AND THE MESSAGES THEY SEND



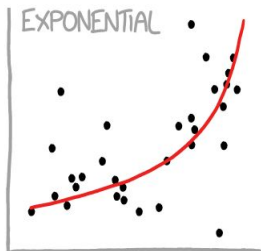
"HEY, I DID A REGRESSION."



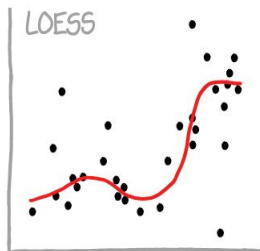
"I WANTED A CURVED LINE, SO I MADE ONE WITH MATH."



"LOOK, IT'S TAPERING OFF!"



"LOOK, IT'S GROWING UNCONTROLLABLY!"



"I'M SOPHISTICATED, NOT LIKE THOSE BUMBLING POLYNOMIAL PEOPLE."

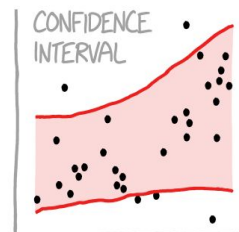


"I'M MAKING A SCATTER PLOT BUT I DON'T WANT TO."

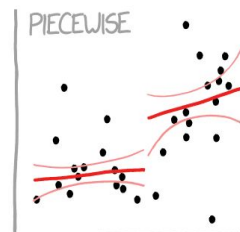
Badness of fit?



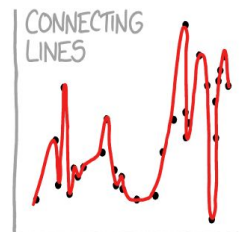
"I NEED TO CONNECT THESE TWO LINES, BUT MY FIRST IDEA DIDN'T HAVE ENOUGH MATH."



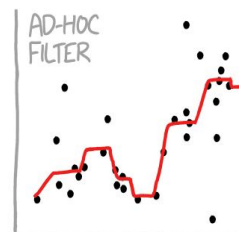
"LISTEN, SCIENCE IS HARD. BUT I'M A SERIOUS PERSON DOING MY BEST."



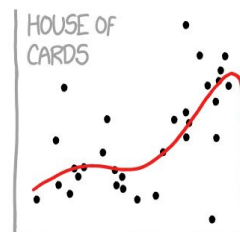
"I HAVE A THEORY, AND THIS IS THE ONLY DATA I COULD FIND."



"I CLICKED 'SMOOTH LINES' IN EXCEL."



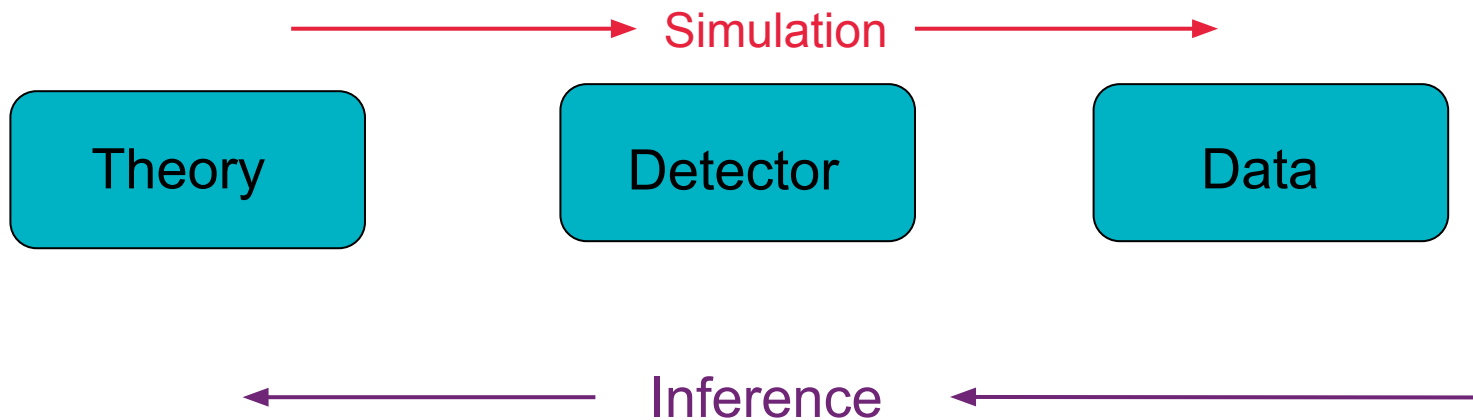
"I HAD AN IDEA FOR HOW TO CLEAN UP THE DATA. WHAT DO YOU THINK?"



"AS YOU CAN SEE, THIS MODEL SMOOTHLY FITS THE- WAIT NO NO DON'T EXTEND IT AAAAAA!!!"

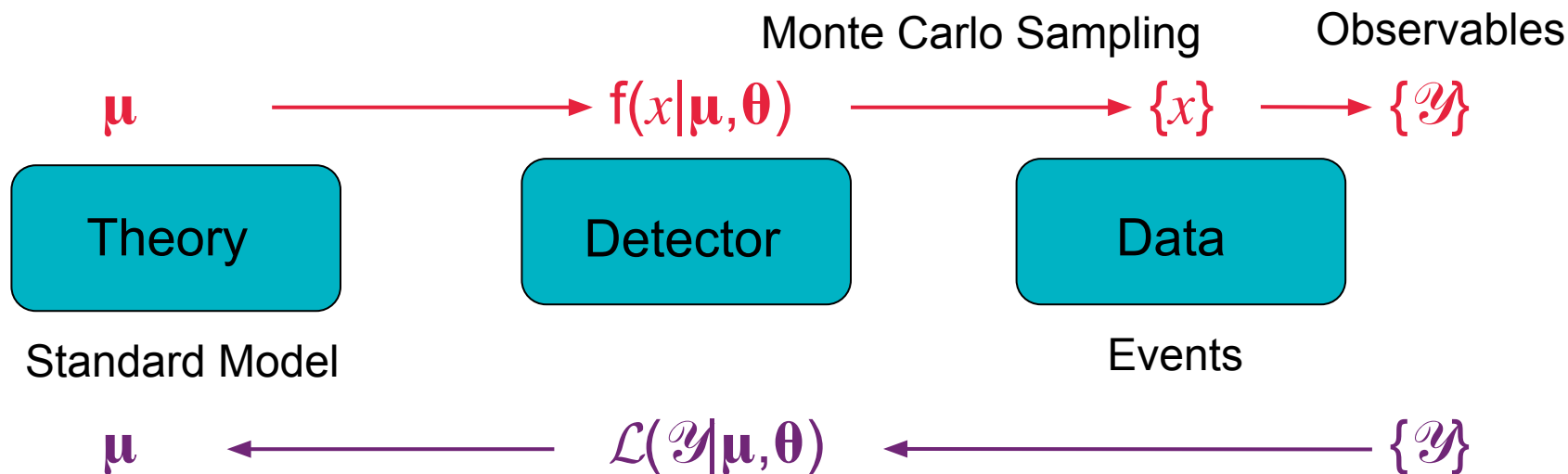
Physics analysis in a nutshell

Simulation based inference: the classical way



Physics analysis in a nutshell

Simulation based inference: the classical way



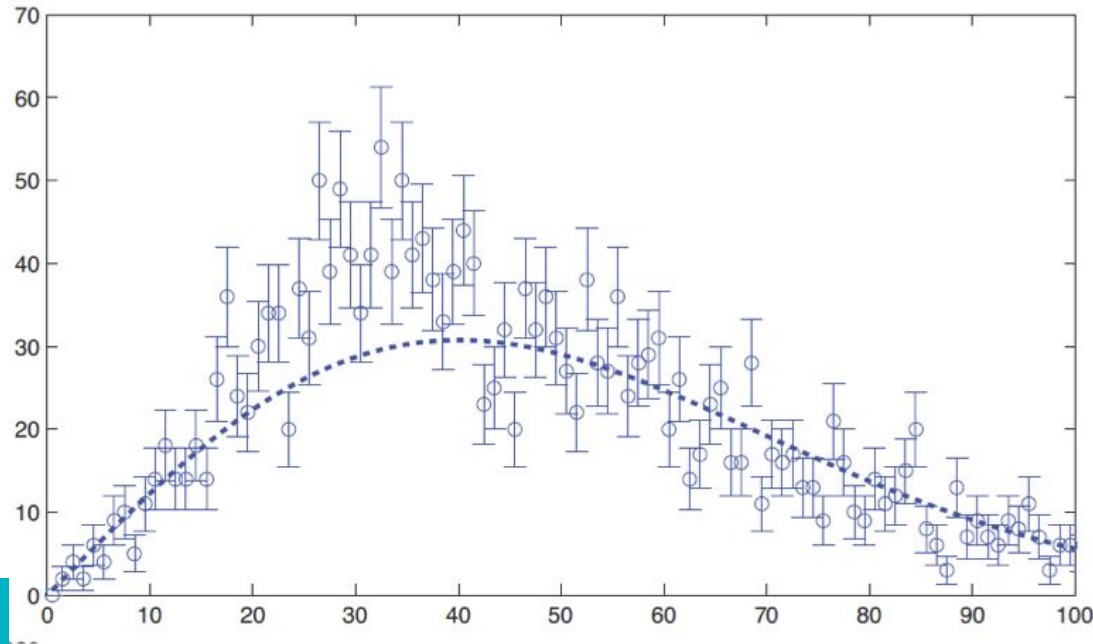
The question: Does my hypothesis describe the data

Basic question: what do we want from the answer?

- Would like a clearly understandable number
- Would like it to match with visual input
- Would like it to have a meaningful interpretation in terms of the likelihood

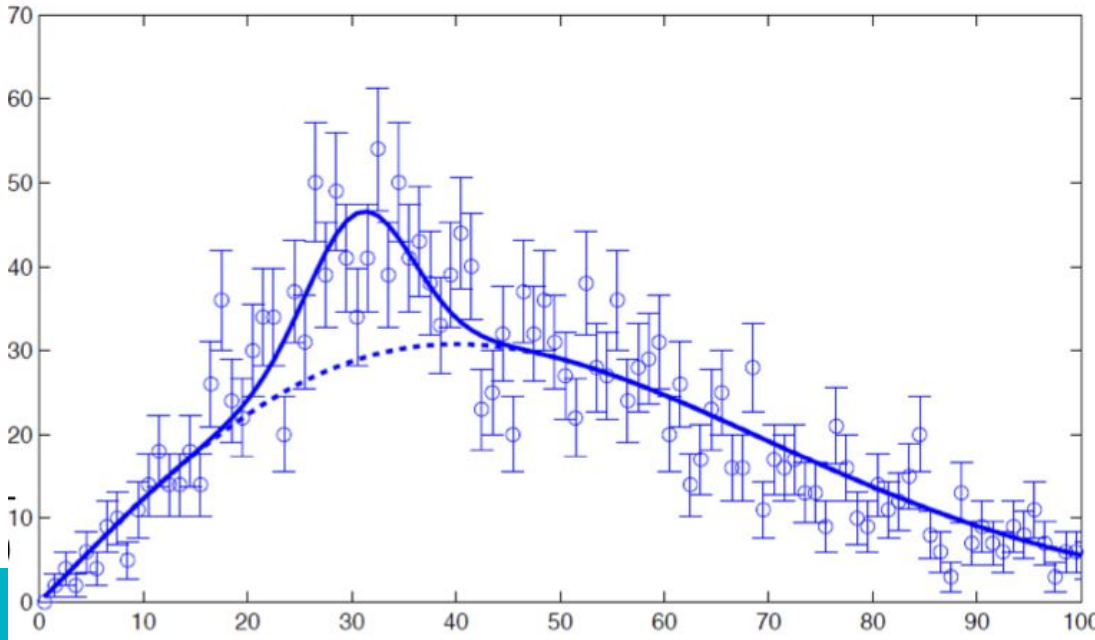
Is there a signal?

Try to distinguish background fluctuations from signals.



Is there a signal?

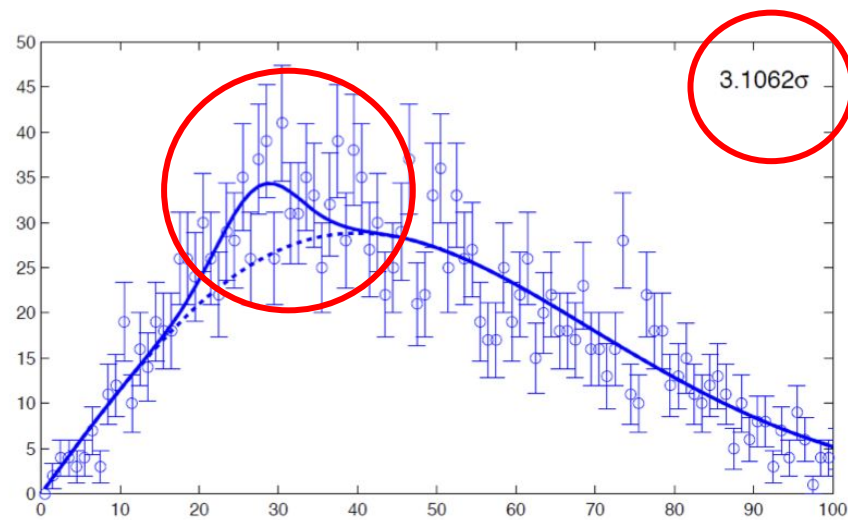
Looks like a signal around $m=30$ maybe



So is there a signal?

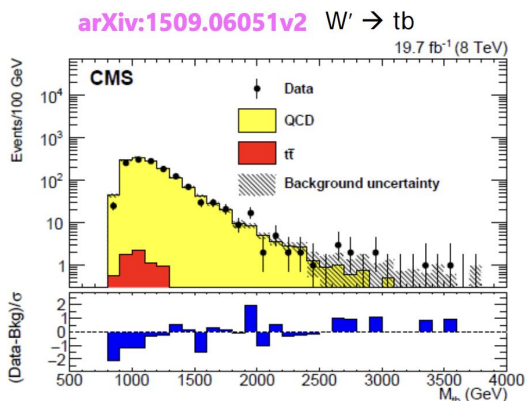
By eye; does not *really* look like a signal but: >3 sigma!

$$q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}_s(m=30) + b)}$$



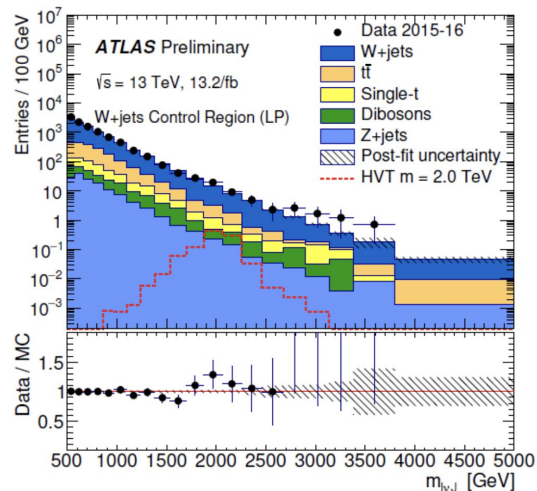
Do we use by eye?

Yes!

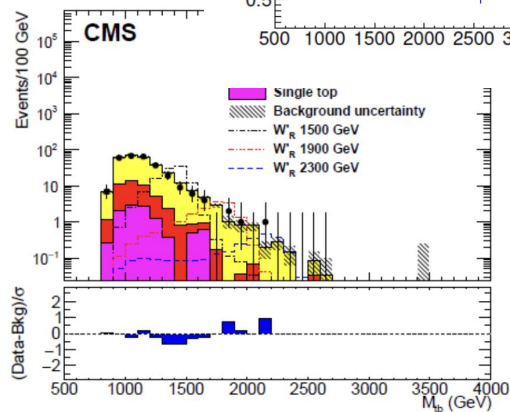


"This test (optical inspection) shows good agreement"

ATLAS-CONF-2016-062 Diboson resonance

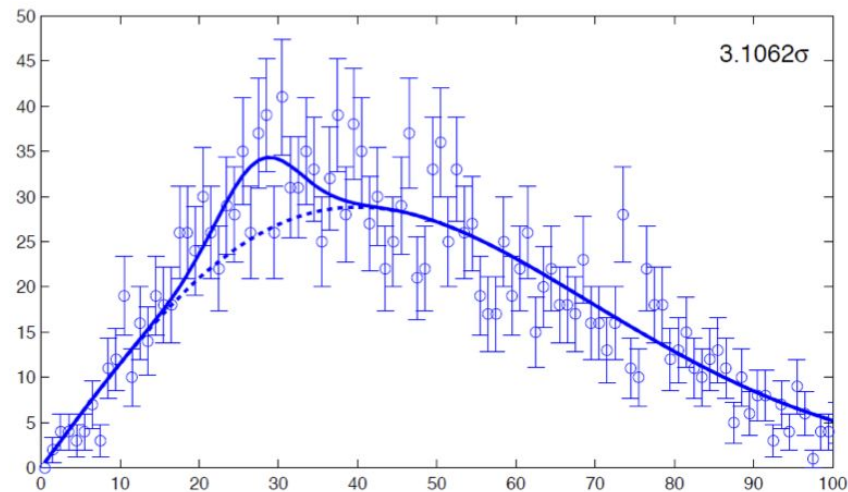
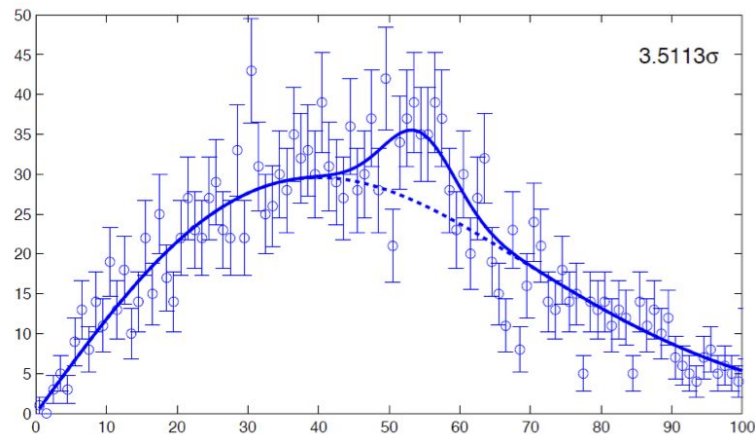
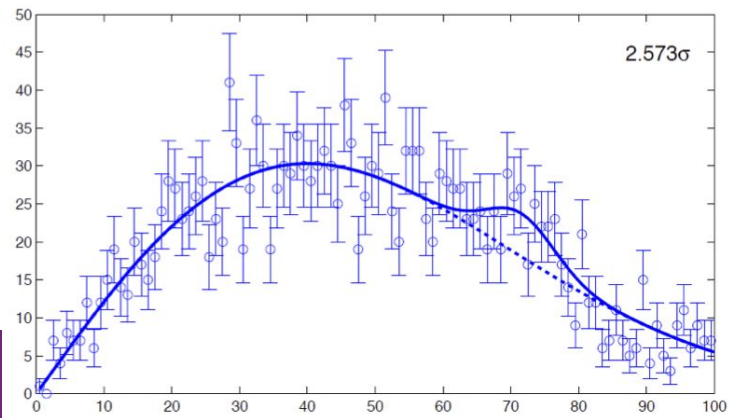
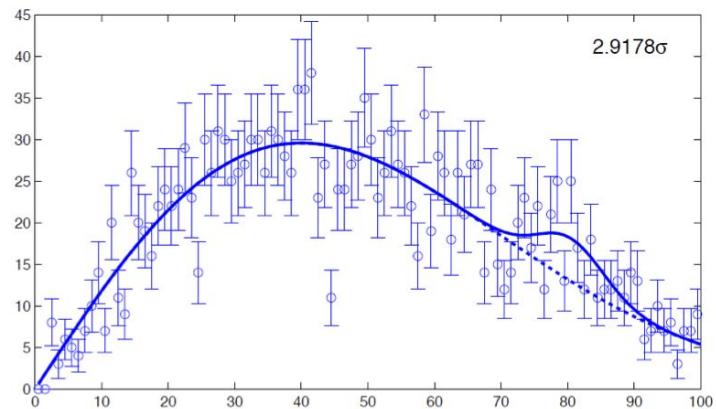


"Good agreement is observed (optical inspection of pulls) between the data and the background prediction"

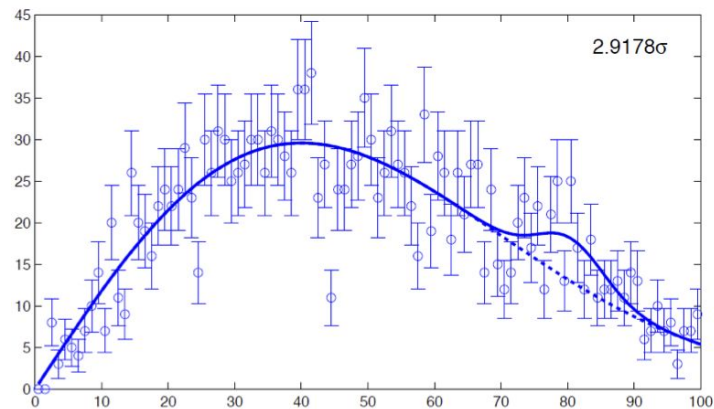


"Good agreement (optical inspection) between data and expectation from SM"

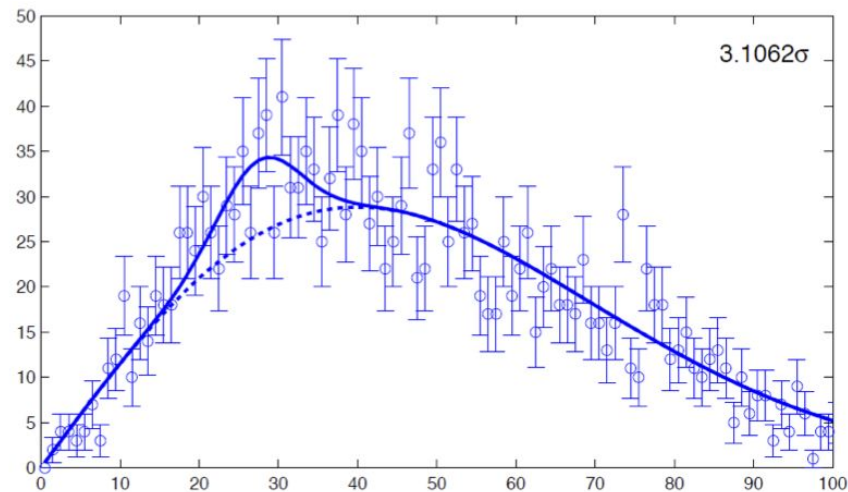
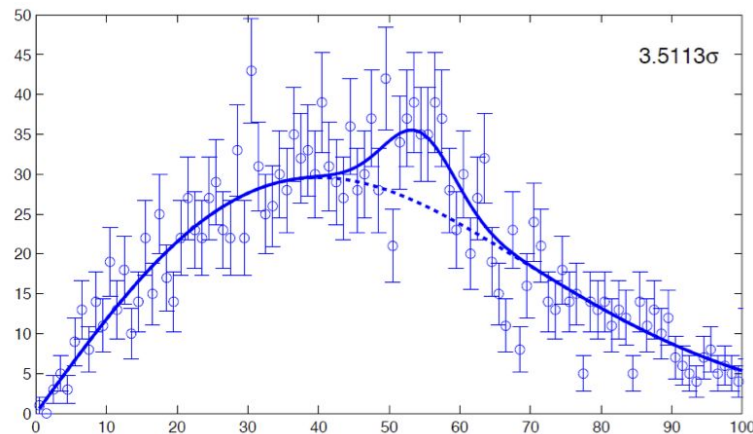
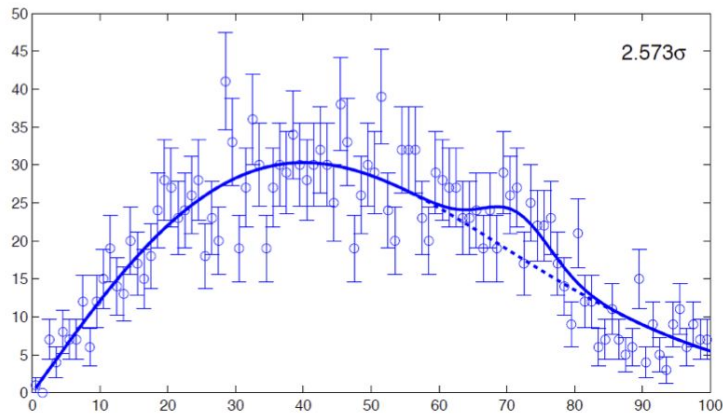
What about these signals?



Answer: All background fluctuations:

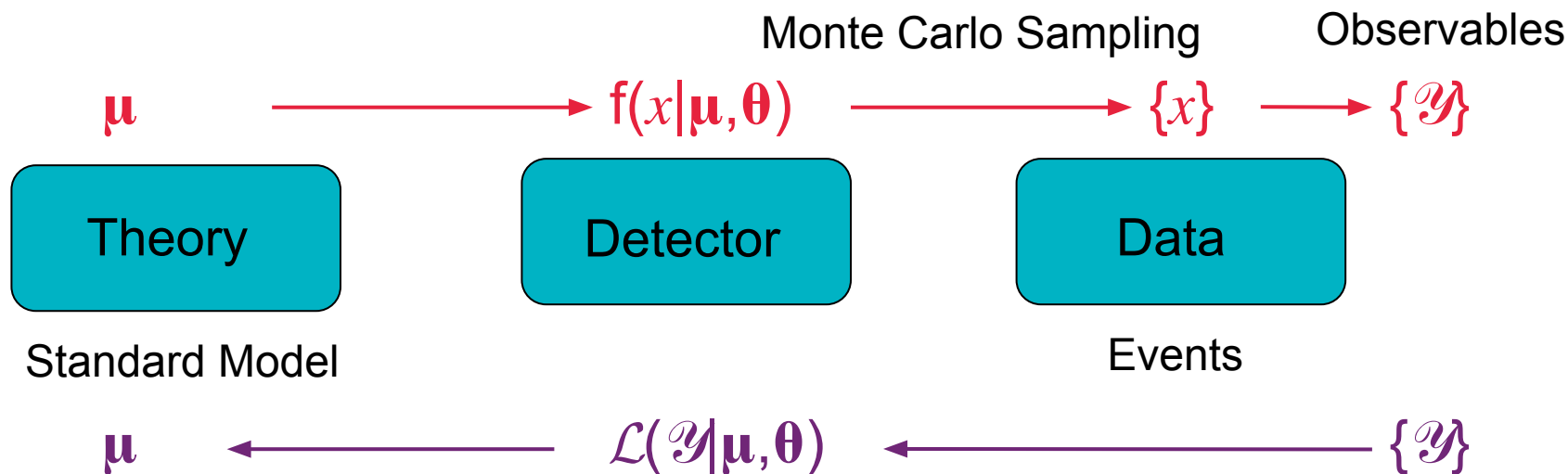


The look
elsewhere
effect



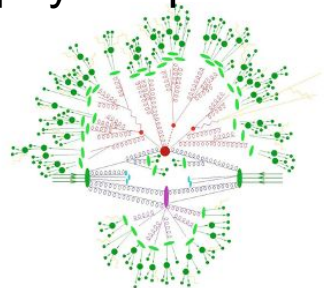
Physics analysis in a nutshell

Simulation based inference

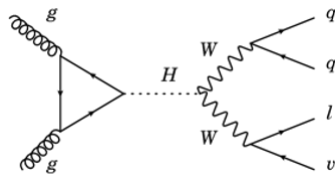


Sources of systematic uncertainties

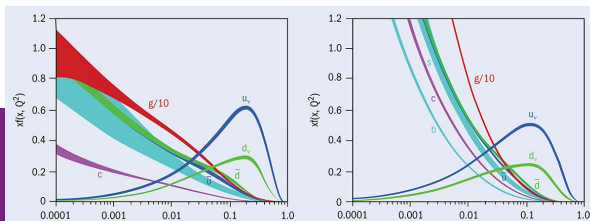
Soft physics process simulation



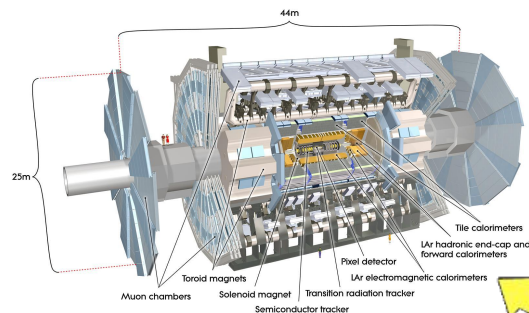
Physics process simulation



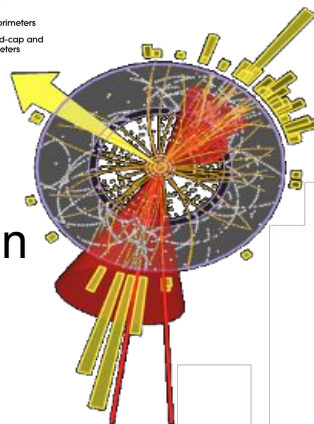
Proton structure functions



Detector simulation



Detector reconstruction



Experimental

$$\mathcal{L}(\mathcal{Y}|\mu, \theta)$$



The top quark 'discovery' at UA1

$W \rightarrow tb$ and $t \rightarrow bl\pm\nu$

→ 2 b jets, charged lepton
missing energy

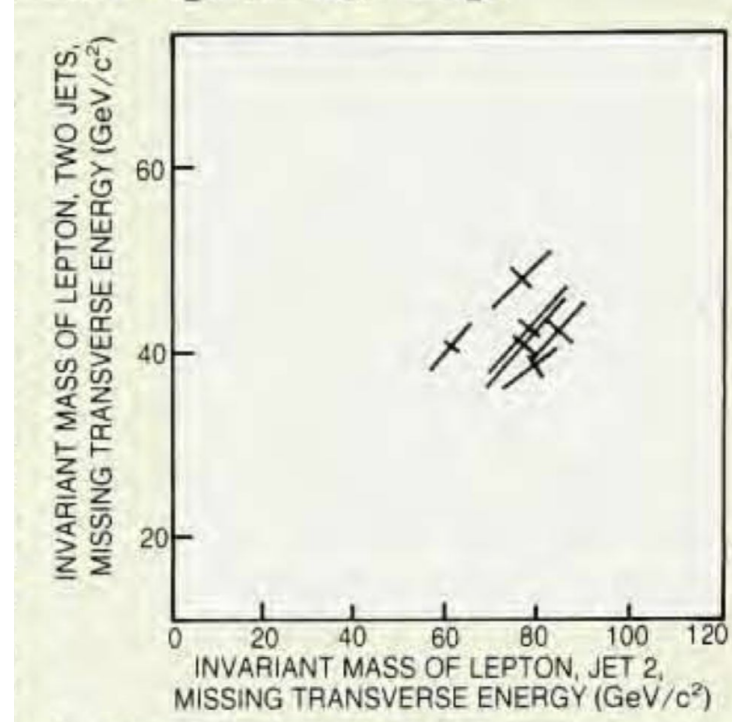
Find 6 events.

→ Plot total mass against $bl\pm\nu$ mass
(ν from missing energy/momentum)

→ W mass in right place

→ t mass around 40 GeV (correct mass approximately 173 GeV)

Turned out to be background - and very creative selection cuts





Wilks theorem

- H_0 : Additional parameters (as predicted by H_1) not needed
- If H_0 correct then according to Wilks' theorem:
 $-\Delta\chi^2 = -2\ln[L(H_1)/L(H_0)]$ should follow for $n \rightarrow \infty$ χ^2 function with
 $\text{ndf} = \text{\#added parameters}$

Wilks' theorem only applies for nested hypotheses:

H_0 : 1st order polynomial H_1 : 2nd order polynomial ✓

H_0 : 1st order polynomial H_1 : $a \cdot \exp(bx + cx^2)$ ✗

Profiling over Nuisance Parameters (θ)

So if we are interested in the μ that maximises $\mathcal{L}(\mathcal{Y}|\mu, \theta)$ we can use Wilk's theorem and instead consider

$$\text{Max } \mathcal{L}(\mu, \theta | \mathcal{Y}) = -2\ln(\mathcal{L}(\mu, \hat{\theta} | \mathcal{Y}) / \mathcal{L}(\hat{\mu}, \hat{\theta} | \mathcal{Y})) = \text{NLL}(\mu, \theta | \mathcal{Y})$$

In the large N limit this is a chi-squared distribution with k degrees of freedom.

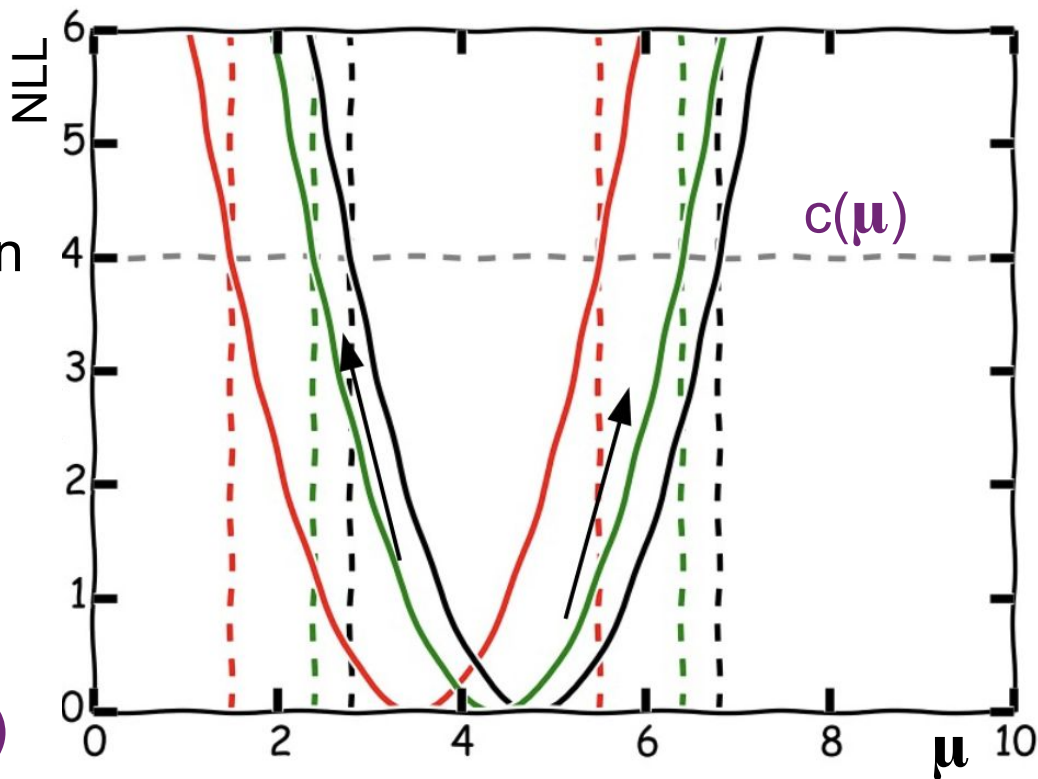
$$\text{Max } \mathcal{L}(\mu, \theta | \mathcal{Y}) = \text{NLL}(\mu, \theta | \mathcal{Y}) = \chi^2_k$$

Confidence intervals

Wilks' theorem (if it applies) makes it relatively simple to construct confidence intervals in a multi-dimensional model-parameter space.

Confidence interval is defined as

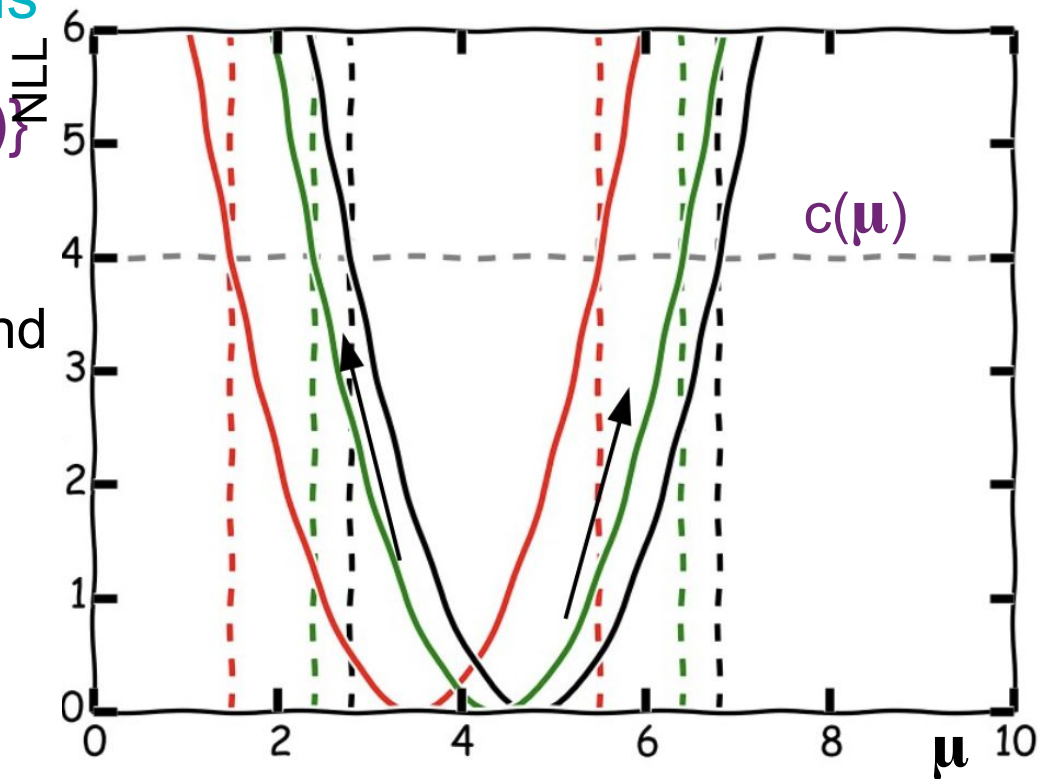
$$I(\mathcal{Y}) = \{\mu \mid \text{NLL}(\mu, \theta \mid \mathcal{Y}) < c(\mu)\}$$



Multidimensional likelihoods

$$I(\mathcal{Y}) = \{\mu \mid \text{NLL}(\mu, \theta | \mathcal{Y}) < c(\mu)\}^{\text{NLL}}$$

The threshold values c for different confidence level depend on the dimensionality of the confidence interval!

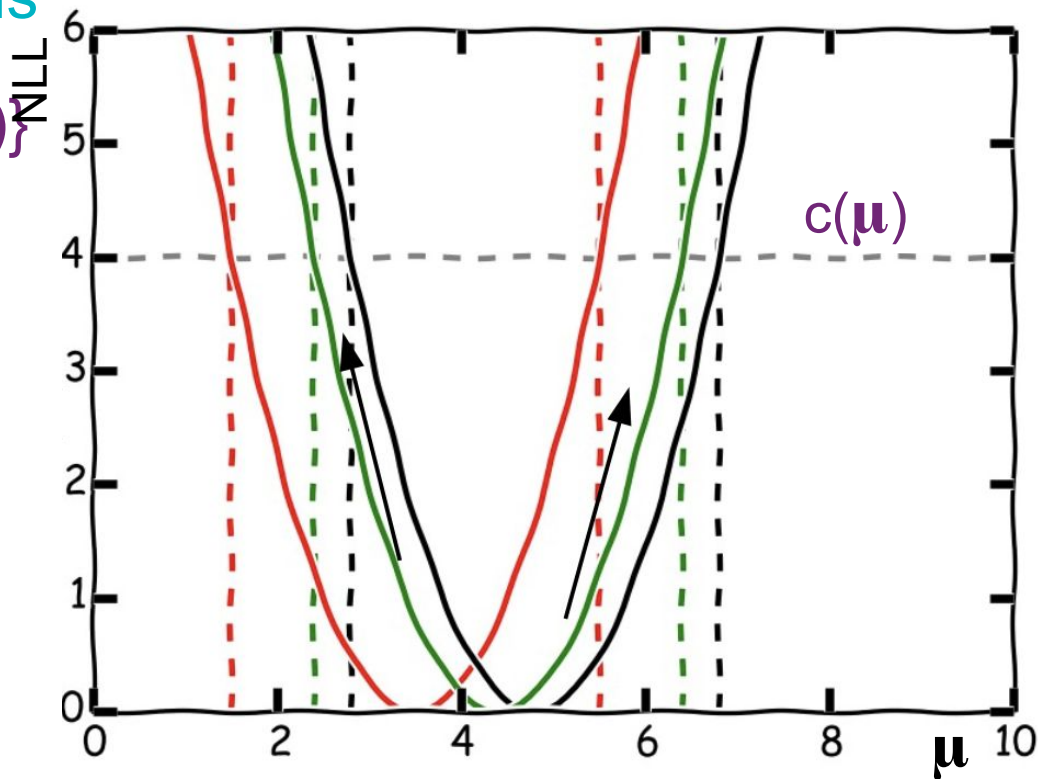


Multidimensional likelihoods

$$I(\mathcal{Y}) = \{\mu \mid \text{NLL}(\mu, \theta \mid \mathcal{Y}) < c(\mu)\}$$

One dimension $\text{NLL} = \chi^2_{k=1}$

68,3% CL	$c=1$
95,4%CL	$c=4$
99,7% CL	$c=9$



Multidimensional likelihoods

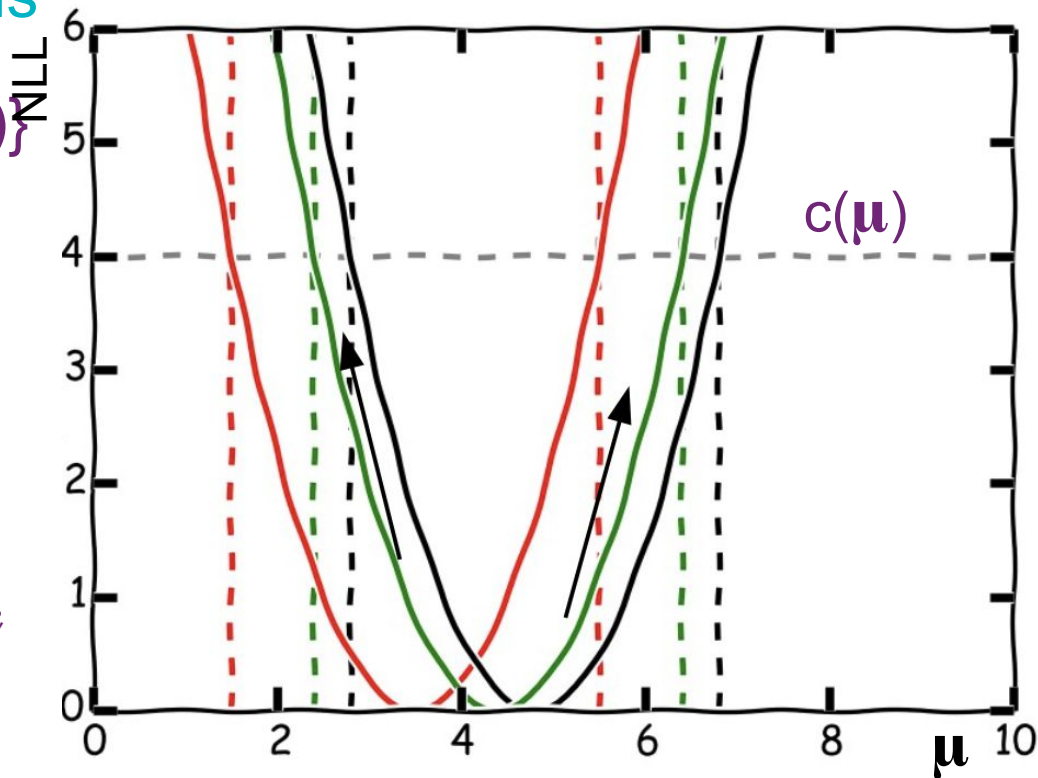
$$I(\mathcal{Y}) = \{\mu \mid \text{NLL}(\mu, \theta \mid \mathcal{Y}) < c(\mu)\}$$

One dimension $\text{NLL} = \chi^2_{k=1}$

68,3% CL	$c=1$
95,4%CL	$c=4$
99,7% CL	$c=9$

Two dimensions $\text{NLL} = \chi^2_{k=2}$

68,3% CL	$c=2,3$
95,4%CL	$c=6,2$
99,7% CL	$c=11,8$



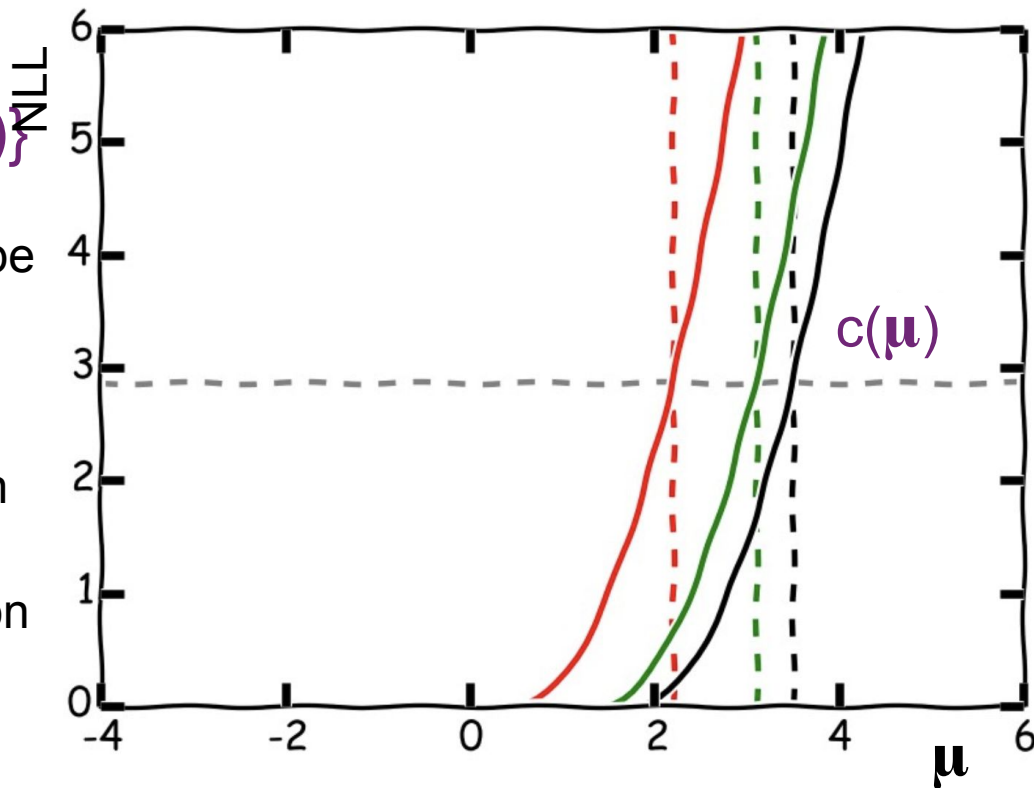
Setting limits

$$I(\mathcal{Y}) = \{\mu \mid \text{NLL}(\mu, \theta \mid \mathcal{Y}) < c(\mu)\}^{\text{NLL}}$$

Upper limits on a parameter (in contrast to two-sided intervals) can be obtained by setting the NLL to zero below the best-fit value.

The threshold c is a bit more tricky in this case, since the NLL follows a χ^2 distribution only half of the distribution. In this case you can show that

95,4%CL	$c=2,86$
---------	----------

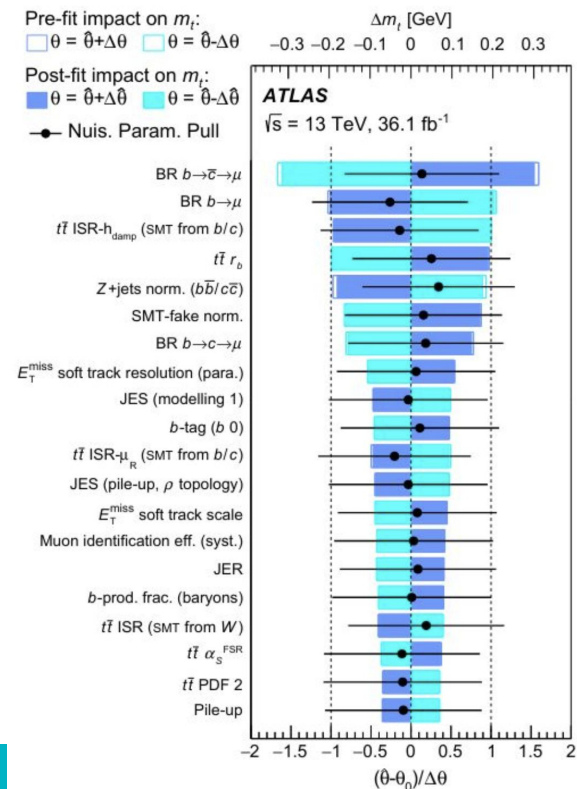


Pulls and rankings: A closer look at Nuisance Parameters

A way to look at the contribution to NLL from Nuisance Parameters (NPs) θ

Shows pre-fit and post-fit *impact* of individual NP on the determination of μ :

- Each NP fixed to ± 1 pre-fit and post-fit sigmas ($\Delta\theta$ and $\Delta\hat{\theta}$ = uncertainty on $\hat{\theta}$)
- Fit re-done with N-1 parameters
- Impact extracted as difference in central value of μ



Post-fit not smaller than pre-fit = no constraint

$0 \pm 1 = \text{No pull}$

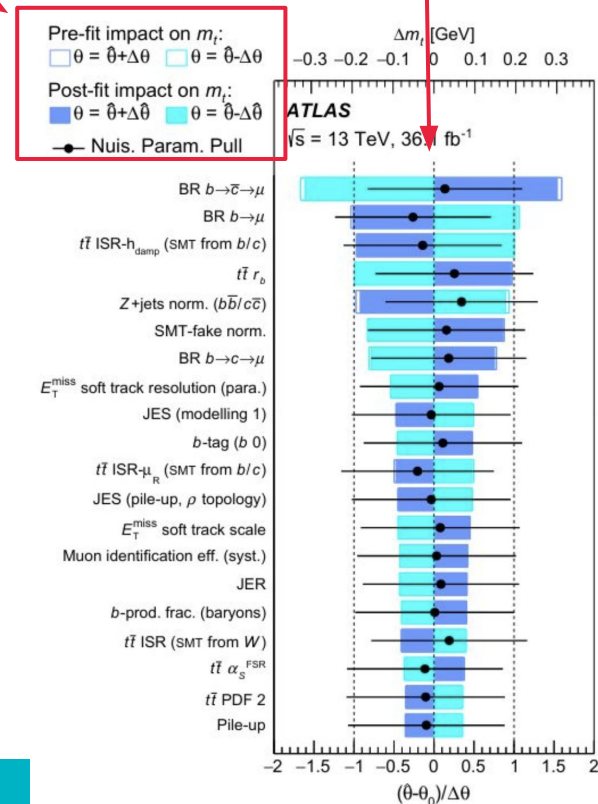
Pulls and rankings: A closer look at Nuisance Parameters

A way to look at the contribution to NLL from Nuisance Parameters (NPs)

Shows pre-fit and post-fit *impact* of individual NP on the determination of μ :

- Each NP fixed to ± 1 pre-fit and post-fit sigmas ($\Delta\theta$ and $\Delta\hat{\theta} = \text{uncertainty on } \hat{\theta}$)
- Fit re-done with N-1 parameters
- Impact extracted as difference in central value of μ

Nomenclature: fitting and parameter estimation



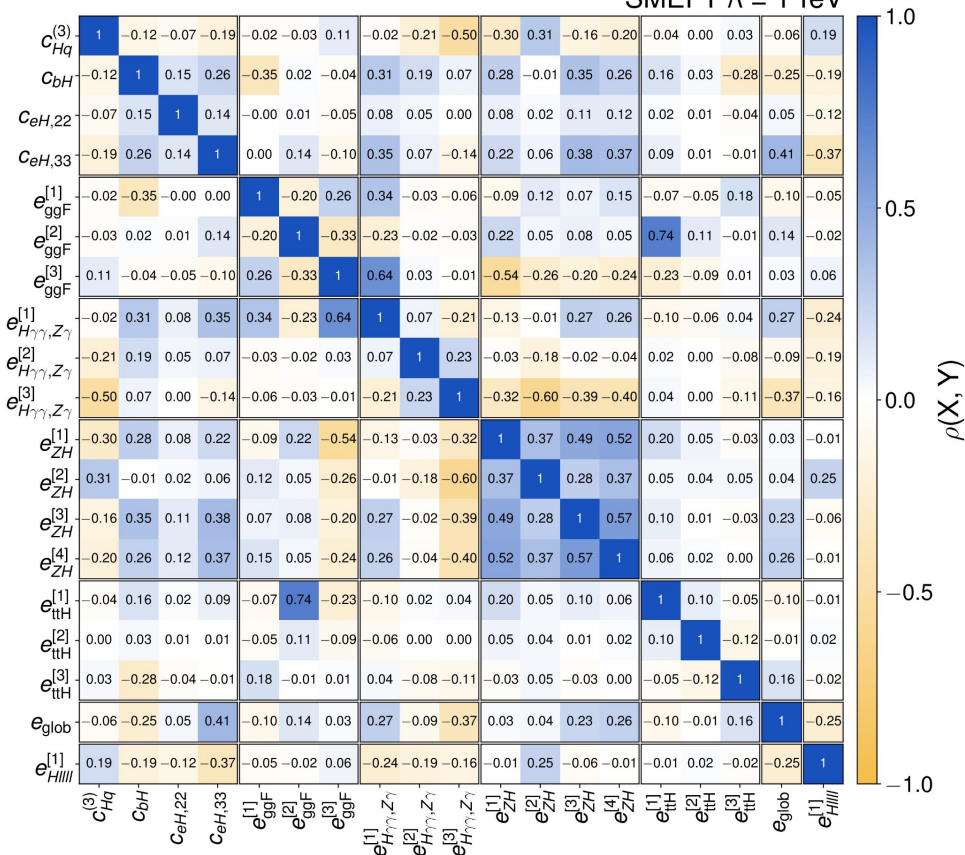
The correlation matrix

Large correlations can make the fit unstable

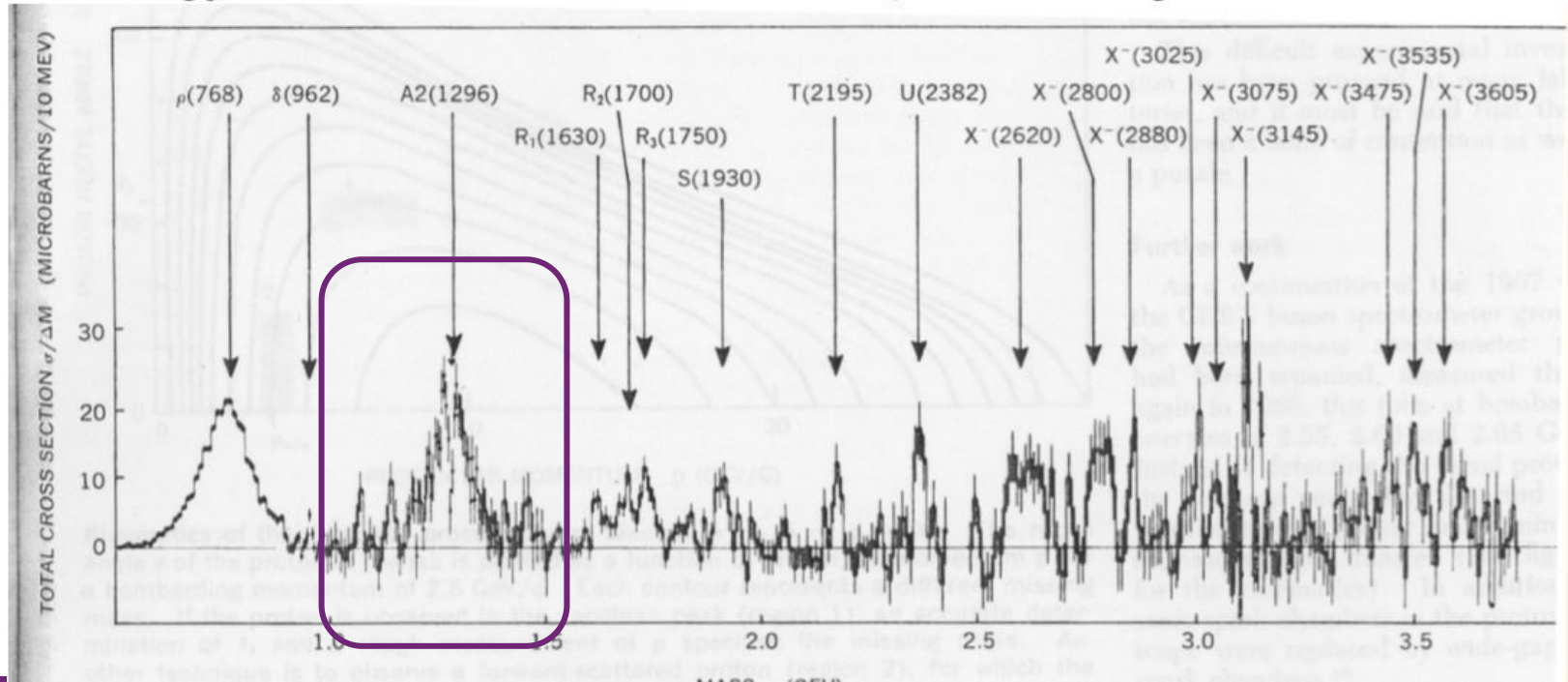
- Not necessarily a problem but worth investigating
- Change of basis can make constraints explicit
- Consider carefully when *reparametrizing*

ATLAS

$\sqrt{s} = 13 \text{ TeV}, 139 \text{ fb}^{-1}$
 $m_H = 125.09 \text{ GeV}, |y_H| < 2.5$
 SMEFT $\Lambda = 1 \text{ TeV}$

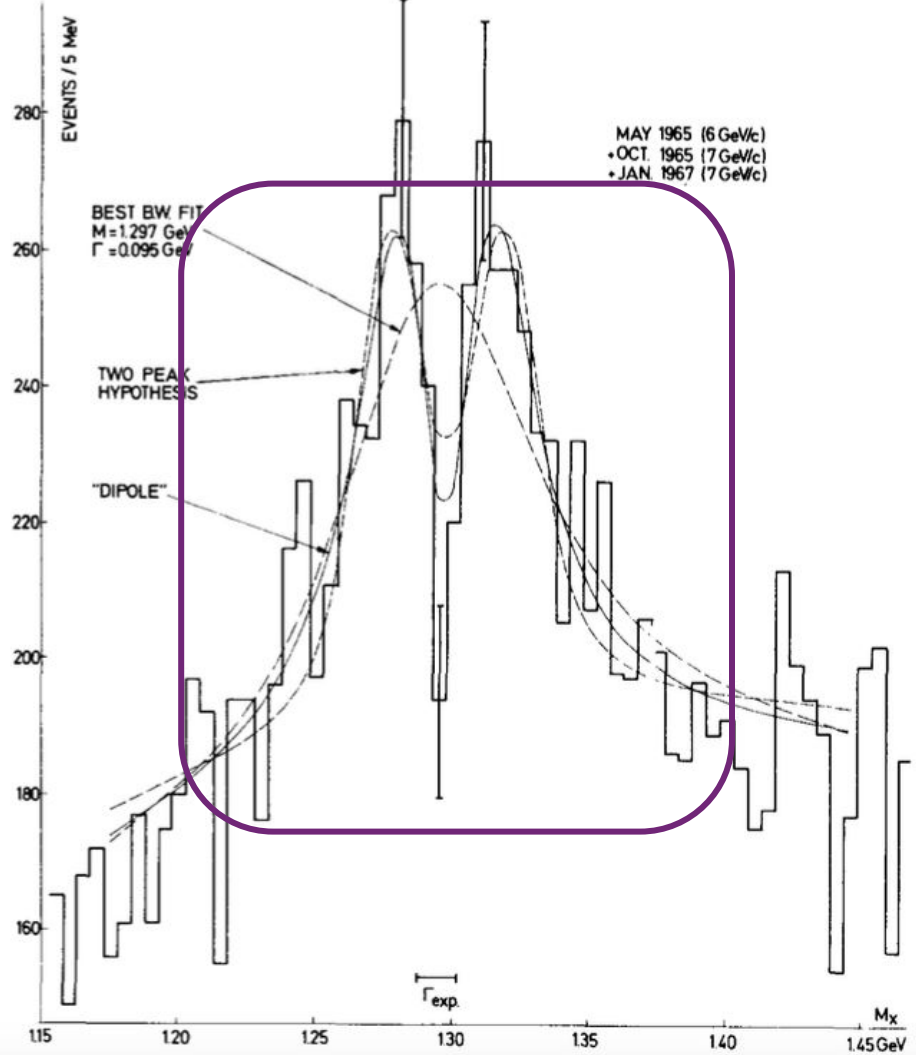


Splitting the A2 meson



Splitting the A2 meson

First reported 1967
More than 6σ effect!

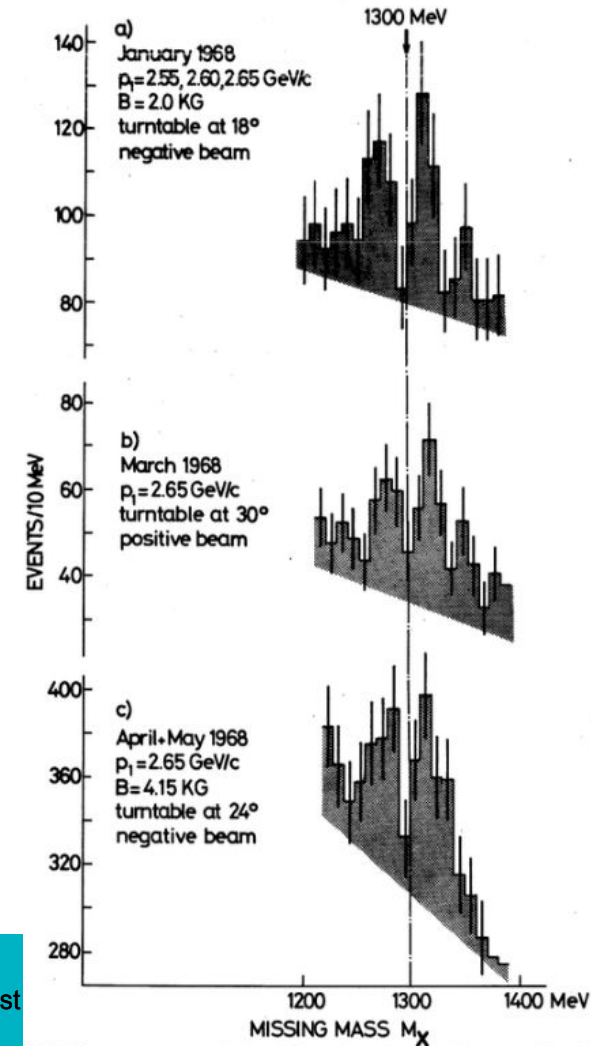


Splitting the A2 meson

First reported 1967
More than 6σ effect!

- Confirmed in 1969
- Went away 1971

1. “Creative” analysis cuts
2. Wanting to see it
 - a. And not publishing when you don't see it

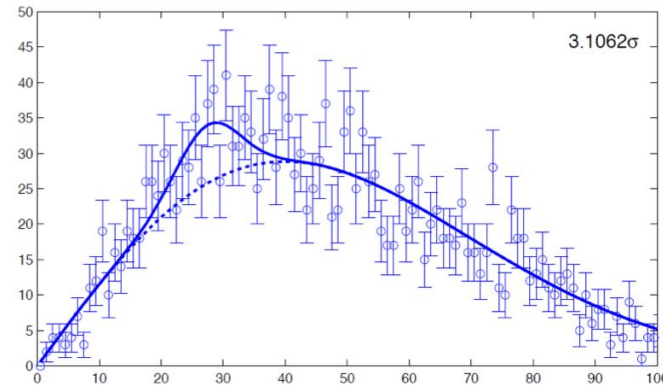
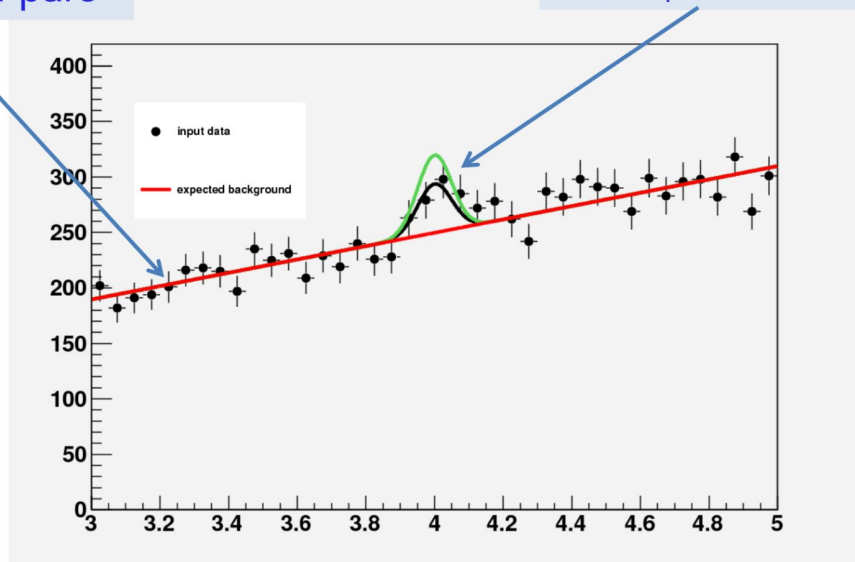


How many parameters do I need?

$$q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}_s(m=30) + b)}$$

1) H_1 : add more background pars

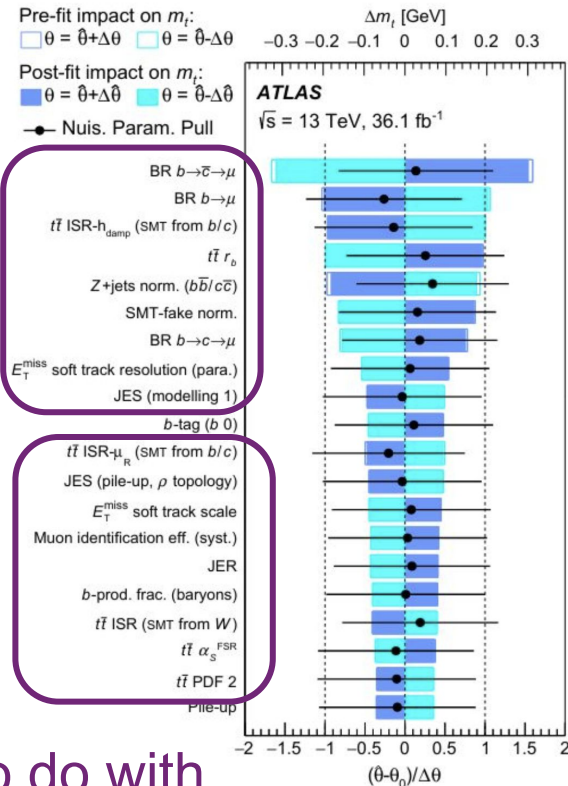
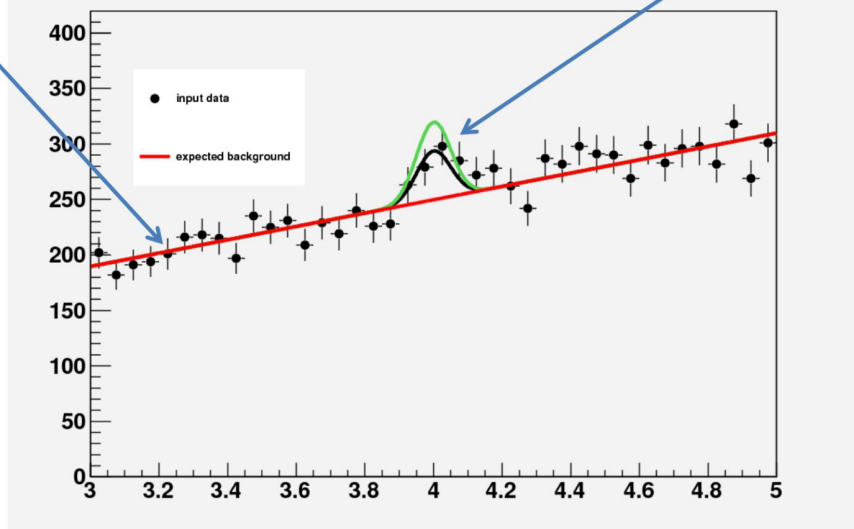
2) H_1 : there is a signal



How many parameters do I need?

1) H_1 : add more background pars

2) H_1 : there is a signal



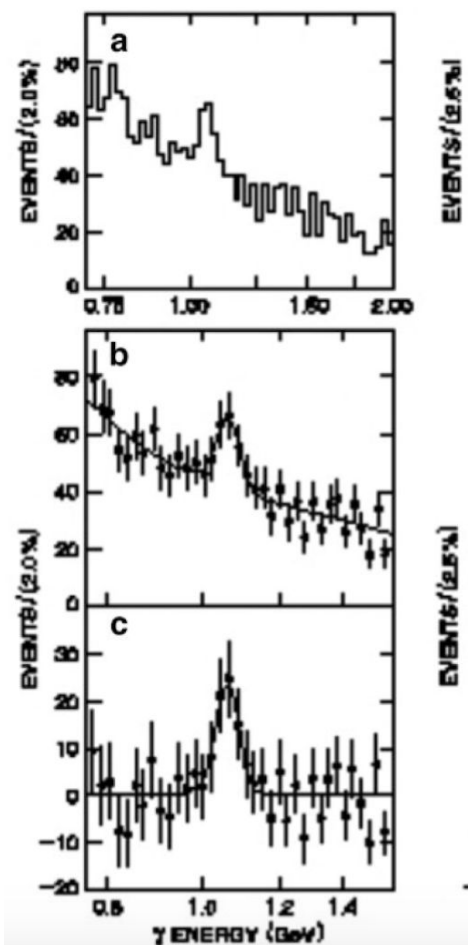
All to do with background estimation

The $\zeta(8.3)$

“Discovered” in 1984 by the Crystal Ball experiment at DESY.

- Measure energy of photons
- Single energy peak seen!!
- Signals $e^+e^- \rightarrow Y \rightarrow \zeta\gamma$
- 4.2 sigma effect
- Plots show (a) raw data , (b) fit, and (c) background-subtracted fit

When more data was taken (in 1985) the peak went away.



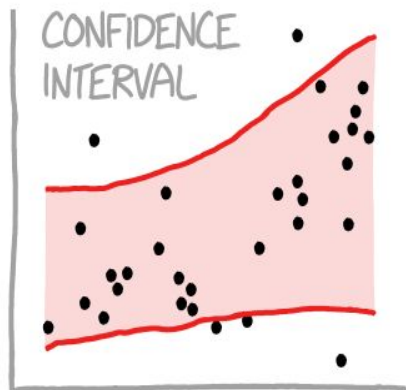
When things go wrong...

“It was easy - I just got a block of marble and chipped away anything that didn't look like David.” - Michaelangelo Buonarotti(attrib.)

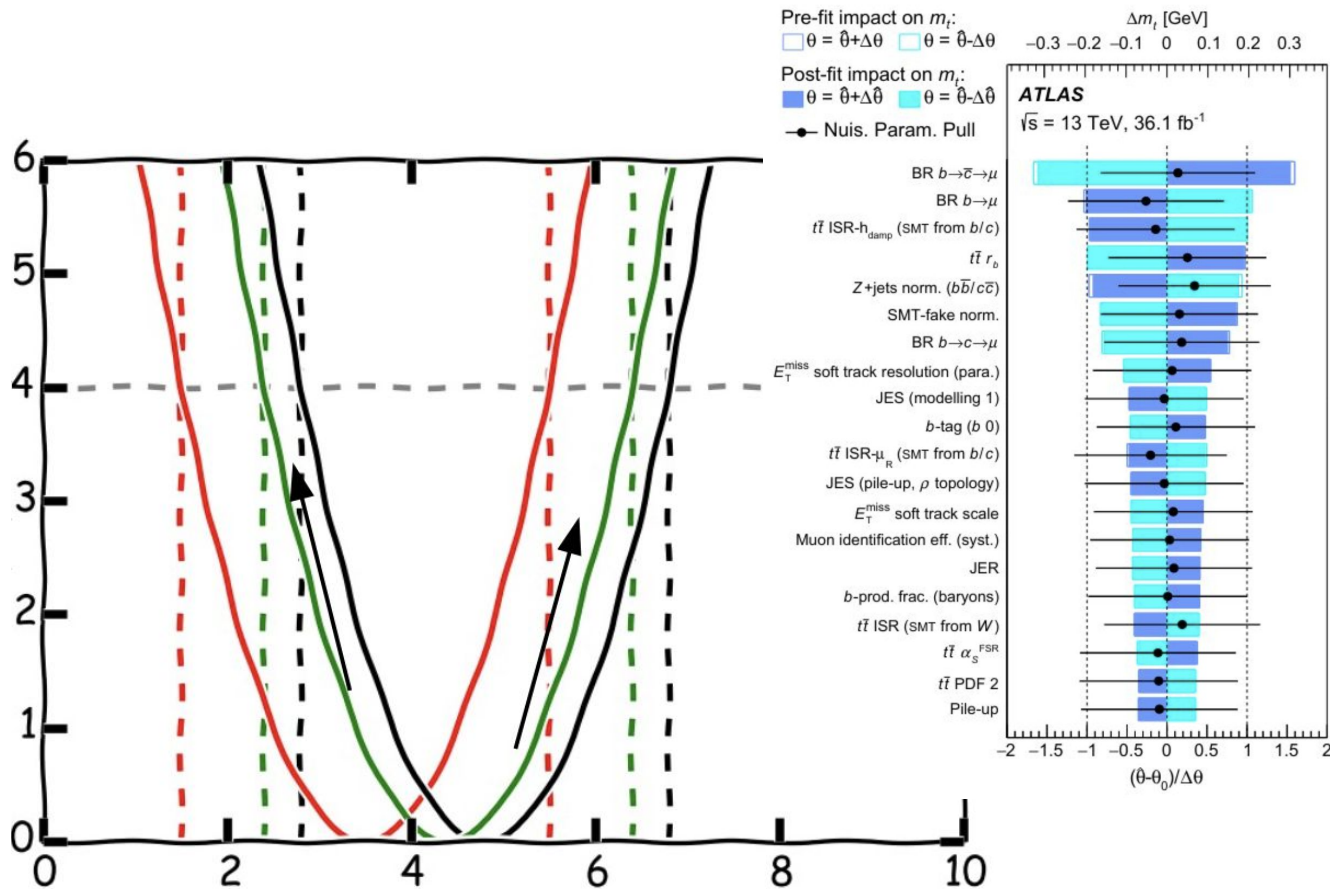
- Maybe good way of creating sculpture - but very bad way of doing physics
- To resist temptation: For searches, devise cuts before looking at the data. Use Monte Carlo simulations, and/or data in ‘sidebands’. Only when cuts are optimised do you ‘open the box’.
- For measurements, apply secret offset until selection frozen.

Partial summary

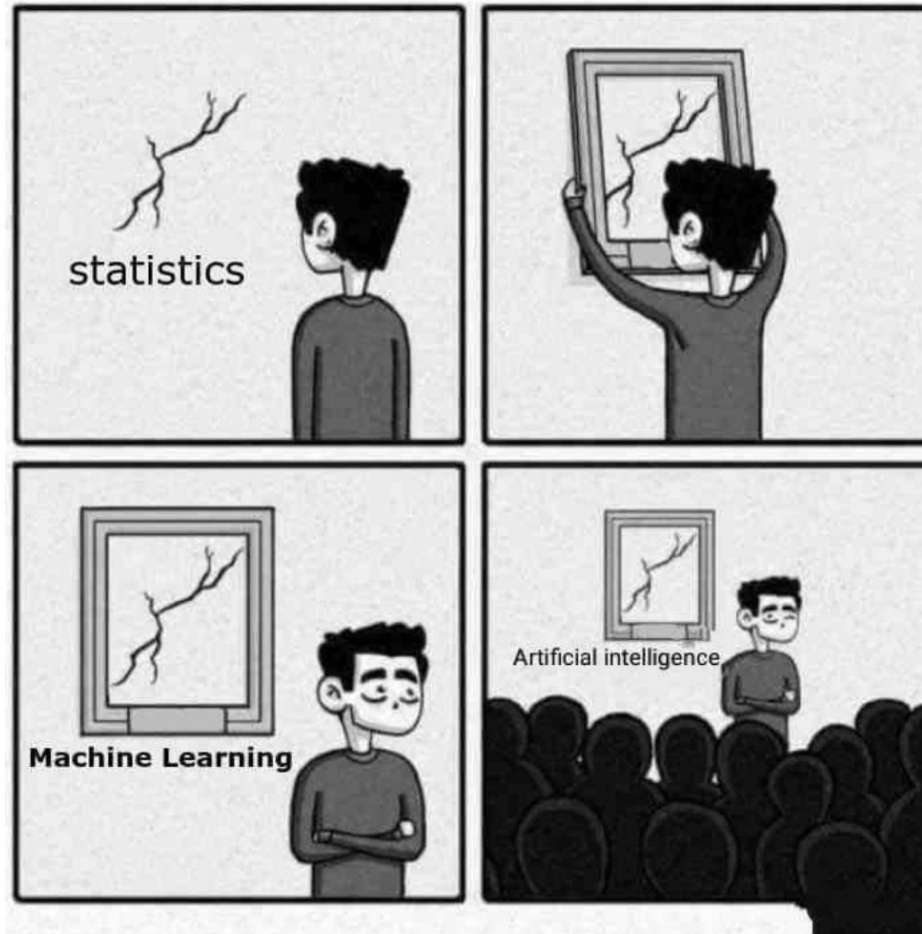
Not too bad...



"LISTEN, SCIENCE IS HARD.
BUT I'M A SERIOUS
PERSON DOING MY BEST."



Stats Meets ML?



What changes for ML?

Ask yourself: Why is my ML doing better?

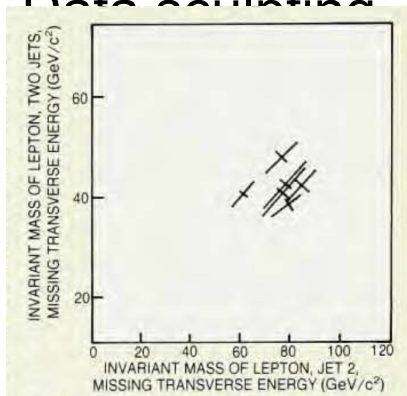
- Multidimensional
 - Does it have to do with correlated parameters?
 - Event-by-event
- Faster
 - Does it parallelise?
 - Simplify
- Wrong assumptions
 - Is there something we assume to be true that it doesn't learn?

General types of ML used in physics

- Supervised Learning
 - Classification
 - Regression
- Unsupervised Learning
 - Clustering
 - Dimensionality reduction and compression
 - Outlier detection
- Weak and Semi-supervised Learning
 - Weak: Label smoothing or correction, Loss reweighting or Multiple-instance learning (MIL)
 - Semi: e.g. Pseudo-labeling with confidence thresholds

What could go wrong?

- Learning MC issues
- Finding the right answer, but no uncertainties
- Data simulation



General types of ML used in physics: special cases

- Physics-informed ML
 - Grounding data-driven models in domain knowledge
 - Goal: accurate, robust, and interpretable outcomes
- Uncertainty aware ML
 - Helps look at intervals rather than best fit values



"LISTEN, SCIENCE IS HARD.
BUT I'M A SERIOUS
PERSON DOING MY BEST."

Interpretability: A warm fuzzy feeling.... Or more than that

When things go wrong... ML edition

“It was easy - I just got a block of marble and chipped away anything that didn't look like David.” - Michaelangelo Buonarotti(attrib.)

- Often the way ML works! You just want to make sure you find a David not Goliath
- To resist temptation: For searches, devise cuts before looking at the data. Use Monte Carlo simulations, and/or data in ‘sidebands’. Only when cuts are optimised do you ‘open the box’.
 - ◆ Much harder to do for ML!!
- For measurements, apply secret offset until selection frozen.

Overfitting and techniques to prevent it

- **Regularization**: Adds a penalty term to the loss function to discourage large parameter values.
- **Dropout**: During training, randomly sets a fraction of neurons to zero at each iteration, which prevents co-adaptation of neurons and acts as an ensemble method.
- **Early Stopping**: Monitors performance on a validation set and halts training when the validation loss stops improving, reducing the risk of overfitting to training data.
- **Data Augmentation**: Increases the effective size of the training set by applying physically valid transformations

The difficult part

- To resist temptation: For searches, devise cuts before looking at the data. Use Monte Carlo simulations, and/or data in ‘sidebands’. Only when cuts are optimised do you ‘open the box’.

Possible checks:

- Injection test
- Smeared or adjusted input data
- Testing on MC from a different generator
- Stay blind

Summary



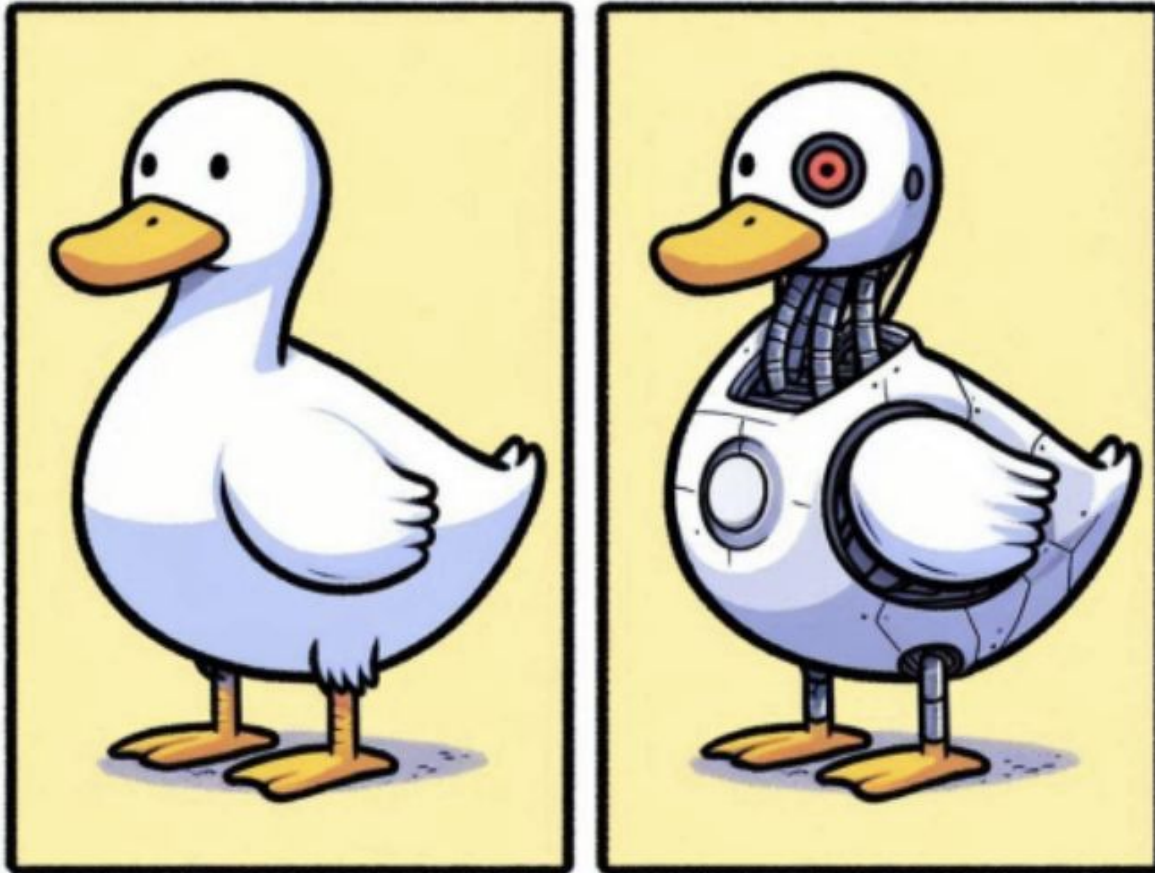
"LISTEN, SCIENCE IS HARD.
BUT I'M A SERIOUS
PERSON DOING MY BEST."

- Using a high level of significance (5σ)
 - ◆ History tells us this is necessary
 - ◆ High output of results
- Us "blind" analysis
- Focus in intervals
 - ◆ Rather than best estimated (fitted) values
- Independent validations
 - ◆ Rumours are hard to get rid off...
- Use common sense and healthy scepticism
- To learn more:
<https://indico.cern.ch/event/1407421>

If it looks like a duck, swims like a duck, and quacks like a duck, then it probably is a duck.

[If it looks like data, it's a sufficiently good simulator?]

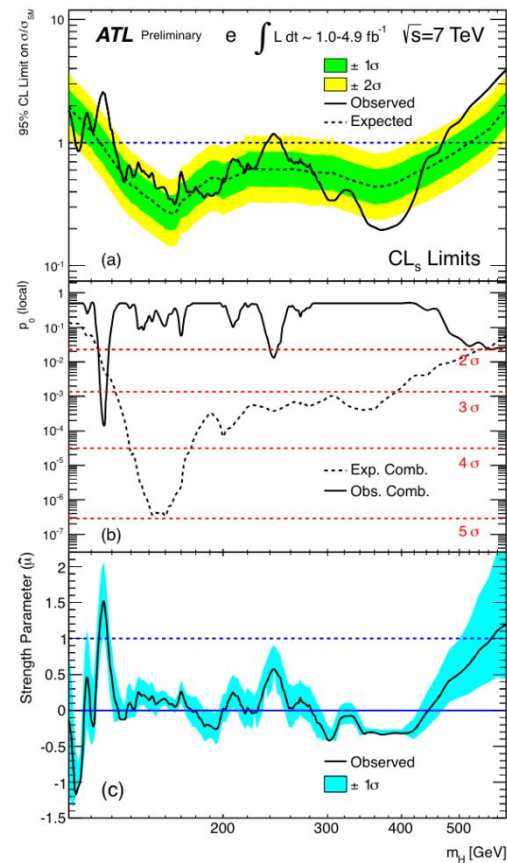
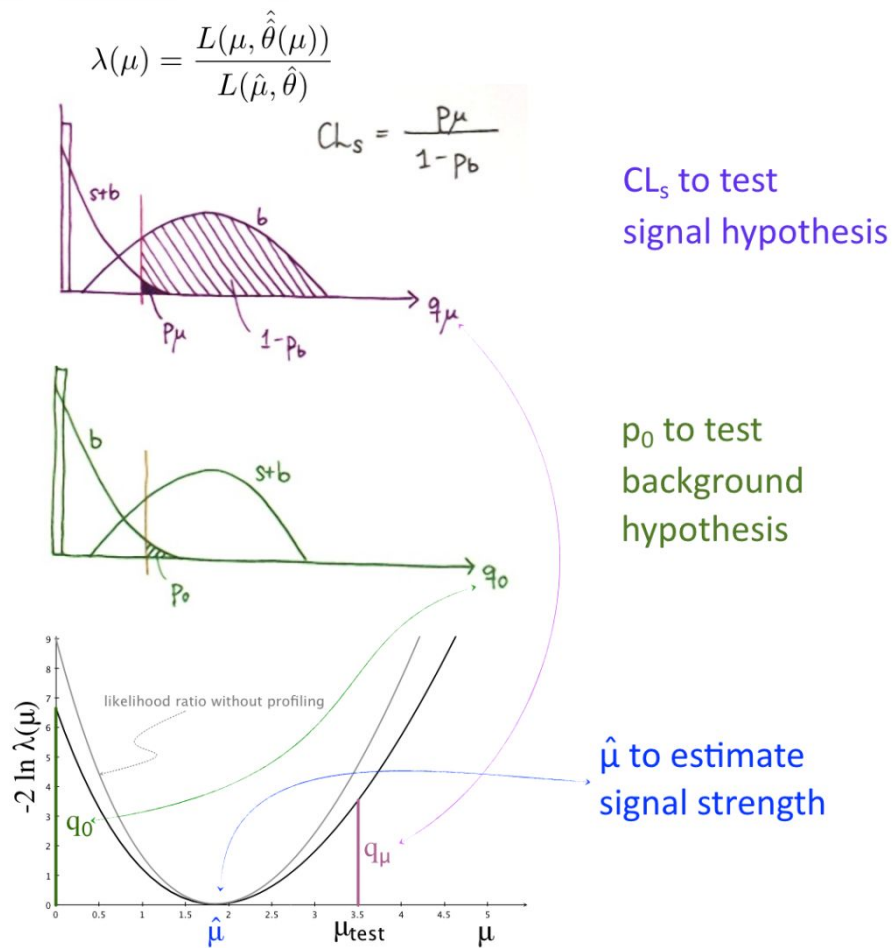
Thanks!



[DALL-E 3 take on the topic]

Back up

Thumbnail of the LHC statistical procedures



Parameter dependence

Say we want to model either $p(x|\theta)$ or $r(x|\theta) = \frac{p(x|\theta)}{p_{\text{ref}}(x)}$.

A few approaches that change structure of the model

- **Point-by-Point:** model $p(x|\theta_i)$ for a set of points $\{\theta_i\}$
 - Not explicitly parametrized in θ , no structure
- **Parametrized Network:** NN models both x -dependence and θ -dependence
 - Most flexible, but doesn't exploit any physics knowledge
- **Fixed Interpolation:** multiple NNs model x -dependence, but form of θ -dependence is fixed & defined by physicist
 - e.g. this is possible for EFT coefficients (exact)
 - This is what HistFactory etc. do for nuisance parameters but this makes assumptions

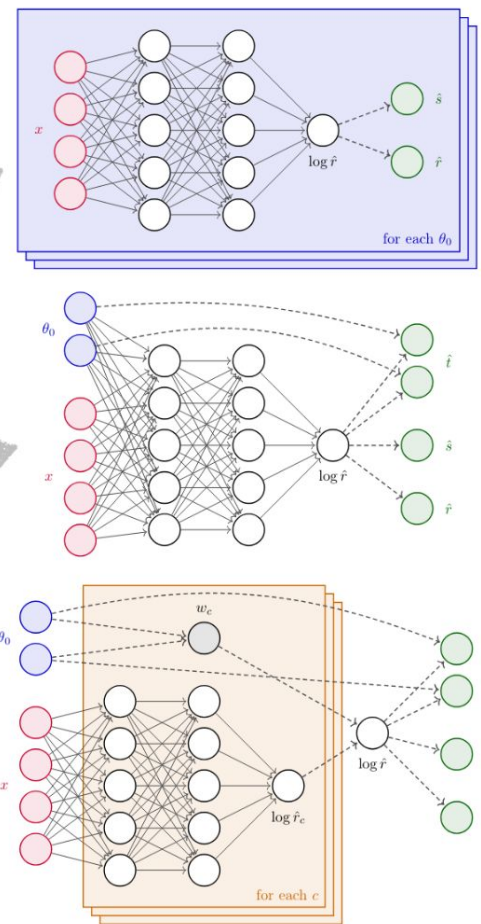
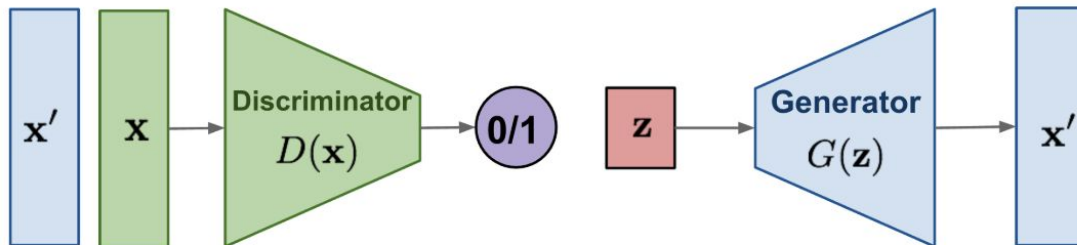
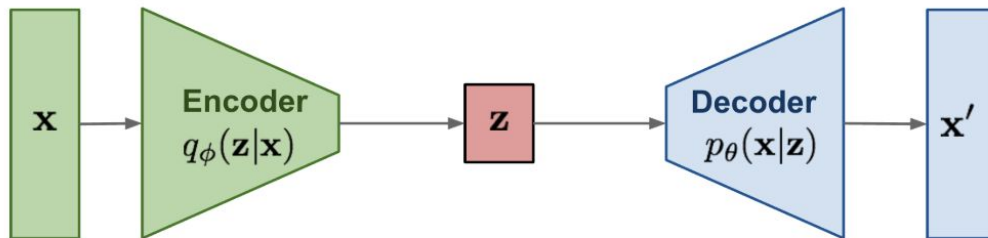


Figure 8: Schematic neural network architectures for point-by-point (top), agnostic parameterized (middle), and morphing-aware parameterized (bottom) estimators. Solid lines denote dependencies with learnable weights, dashed lines show fixed functional dependencies.

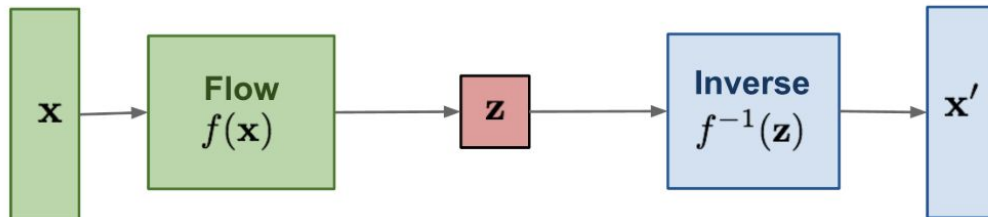
GAN: Adversarial training



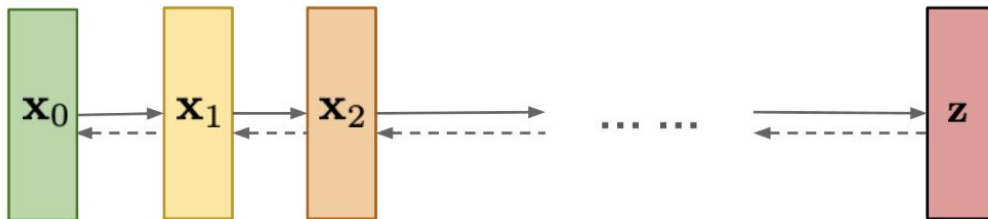
VAE: maximize variational lower bound



Flow-based models:
Invertible transform of distributions



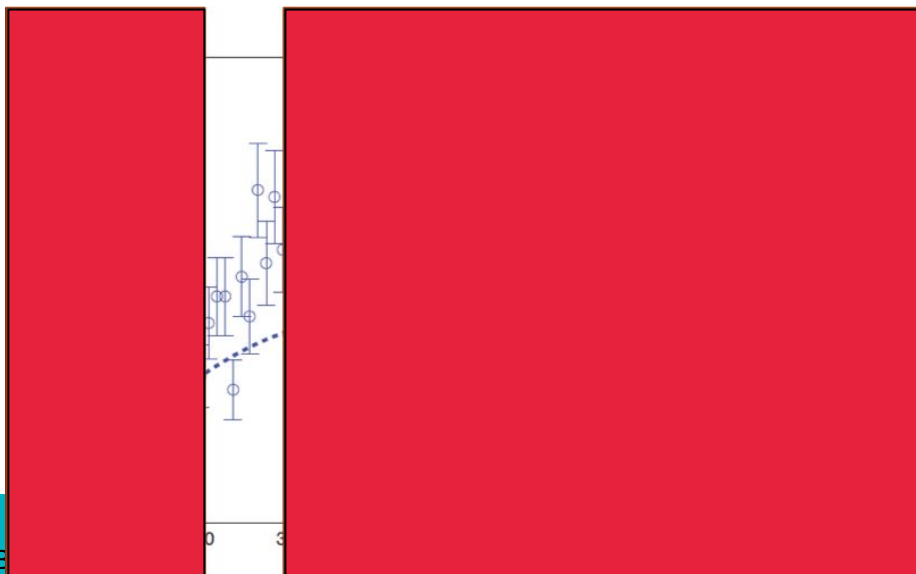
Diffusion models:
Gradually add Gaussian noise and then reverse



Look elsewhere effect: Option 1 Fixed Mass

Scan mass range in predefined steps.

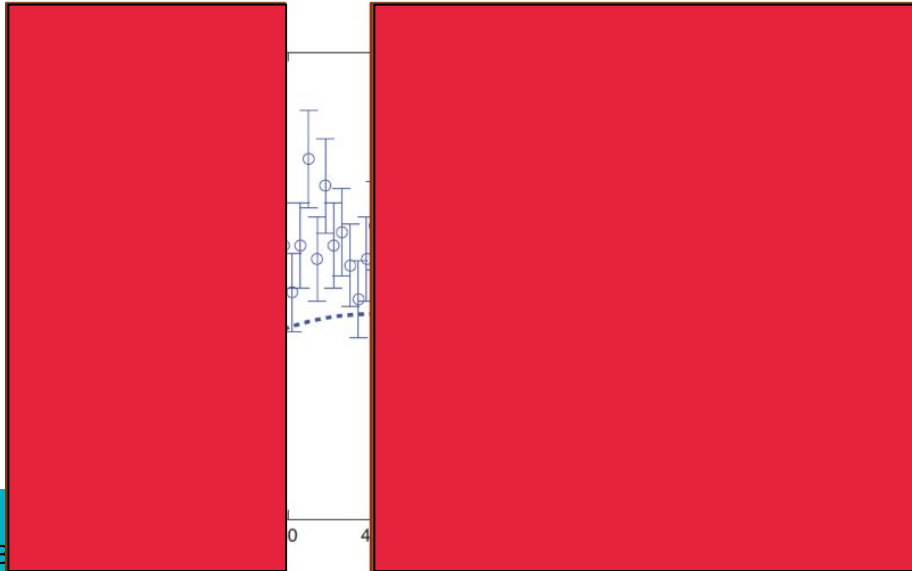
Perform a fixed mass analysis at each point $q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}_s(m=30) + b)}$



Look elsewhere effect: Option 1 Fixed Mass

Scan mass range in predefined steps.

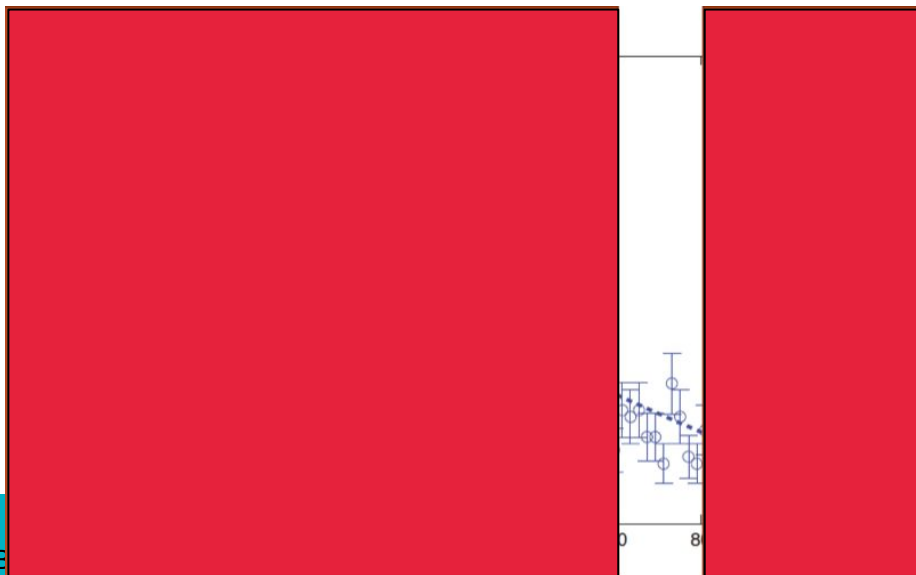
Perform a fixed mass analysis at each point $q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}_s(m=30) + b)}$



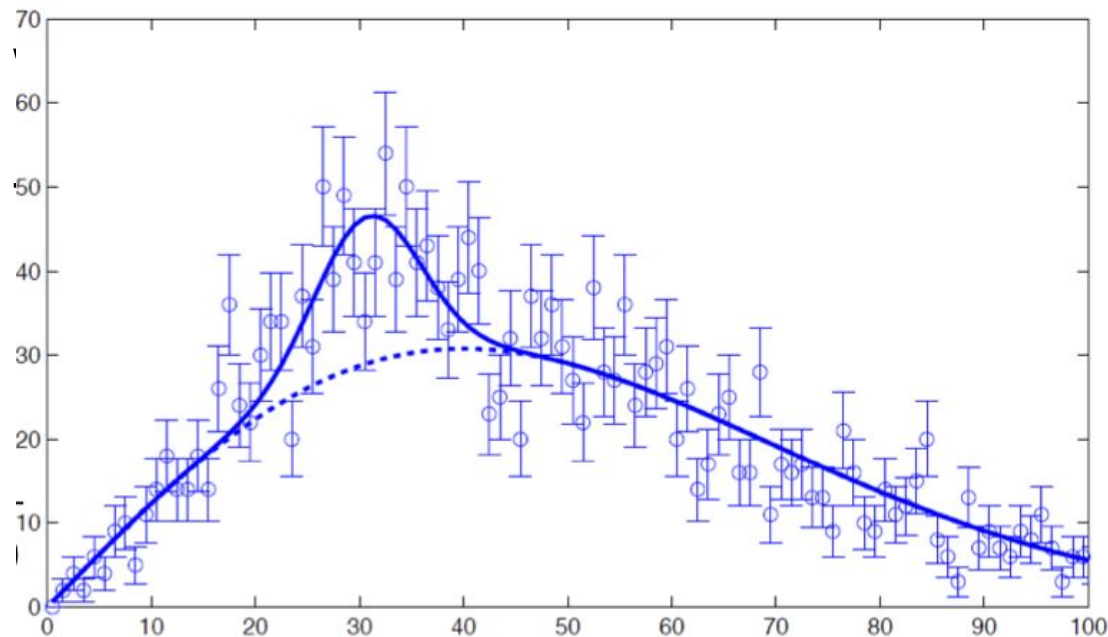
Look elsewhere effect: Option 1 Fixed Mass

Scan mass range in predefined steps.

Perform a fixed mass analysis at each point $q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}_s(m=30) + b)}$



Is there a signal?



$$q_{fix,obs} = -2 \ln \frac{L(b)}{L(\hat{\mu}s(m=30) + b)}$$

Is there a signal?

For GOF tests with binned data:

→ compare observed event numbers n_i with expectation values f_i

Since no H1 specified there are many different GOF tests possible

$$\chi^2 = \sum_i (f_i - n_i)^2 / \sigma^2$$

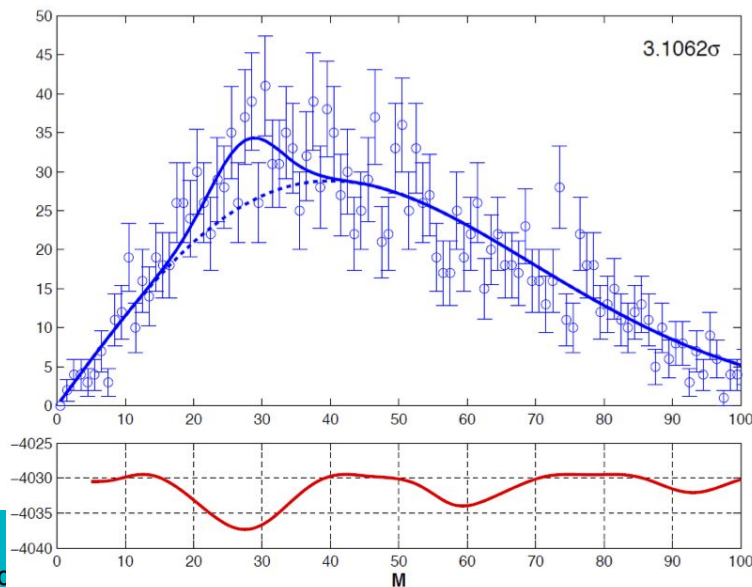
→ Basically does what you do by eye; Minimise distance from hypothesis to the data points

Look elsewhere effect: Option 2 Floating mass

Leave the mass floating

Use a modified test statistic

$$q_{float,obs}(\hat{\mu}, \hat{m}) = -2 \ln \frac{L(b)}{L(\hat{\mu}s(\hat{m}) + b)}$$



Look elsewhere effect: Option 3 Trials factor

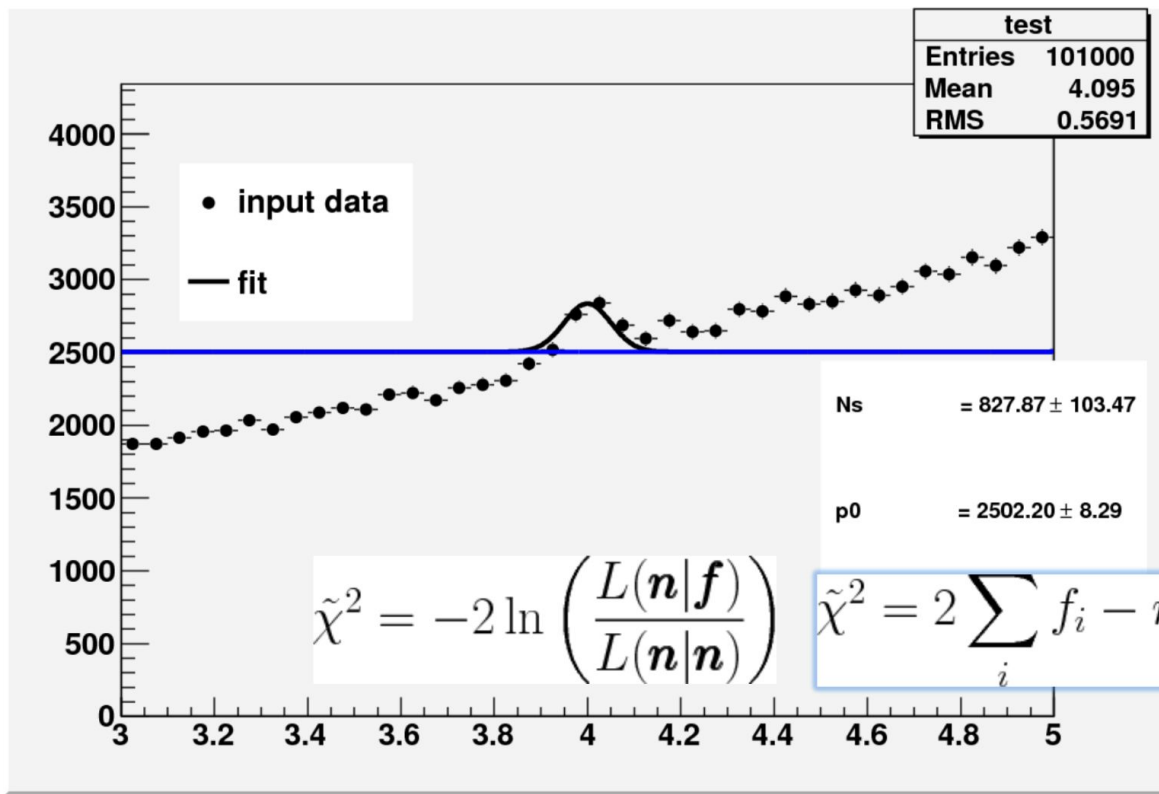
Define your significance of a local excess of events by taking into account the probability of observing an excess anywhere in the range. Quantify in terms of a trial factor

- Ratio between the probability of observing the excess at some fixed mass point to the probability of observing it anywhere in the range.
- The trial factor grows linearly with the (fixed mass) significance

$$trial\# = \frac{\int_{q_{obs}} f(q_{float} | H_0) dt_{float}}{\int_{q_{obs}} f(q_{fix} | H_0) dt_{fix}} = \frac{p_{float}}{p_{fix}} > 1$$

Recommended

How many parameters do I need? – Example

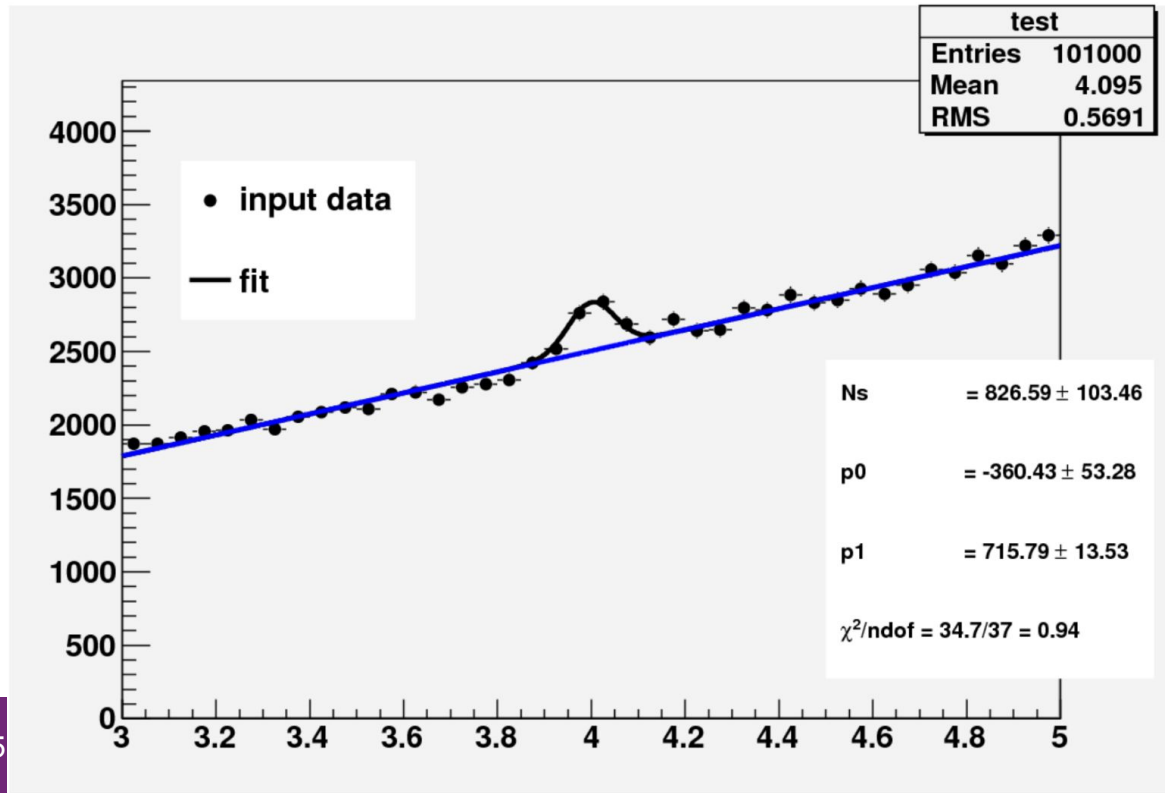


Fit function:
gauss+p0

$$\chi^2 = \sum_i (f_i - n_i)^2 / \sigma^2$$

$$\tilde{\chi}^2 = 2880$$

How many parameters do I need? – Example

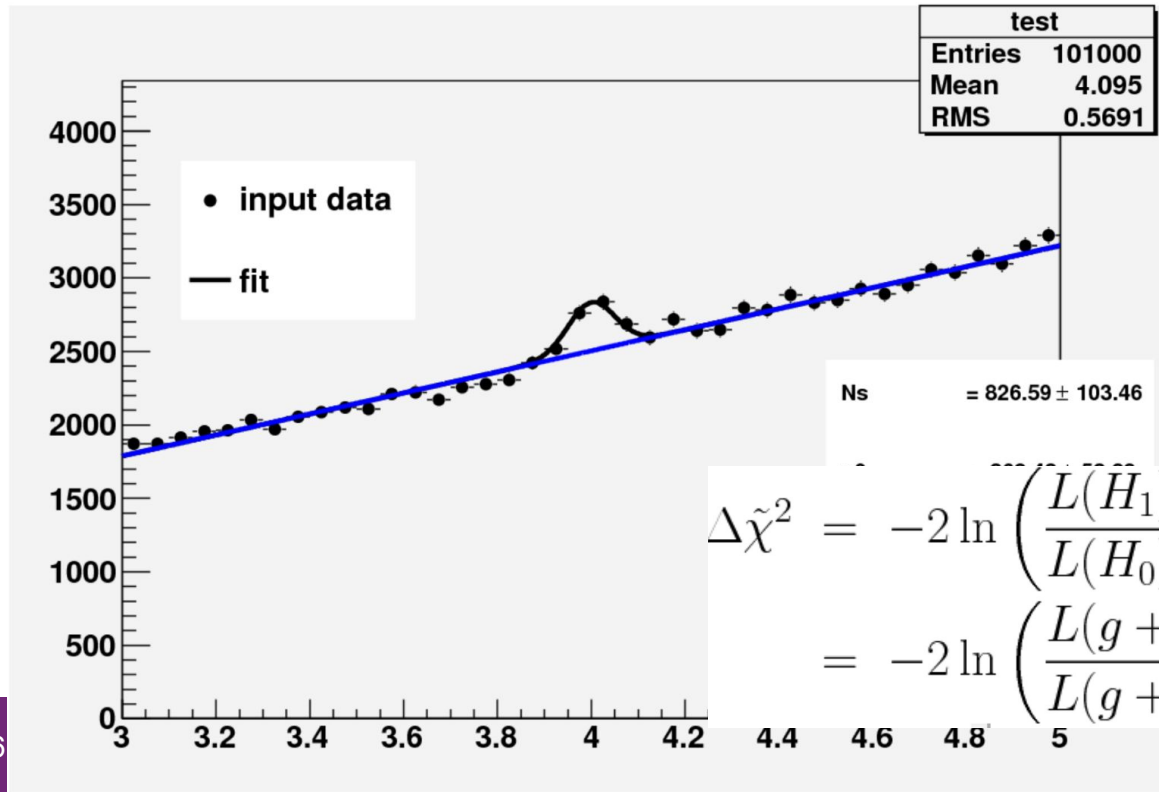


Fit function:
gauss+p1

$$\tilde{\chi}^2 = 34.7$$

Should we stop here?

How many parameters do I need? – Example



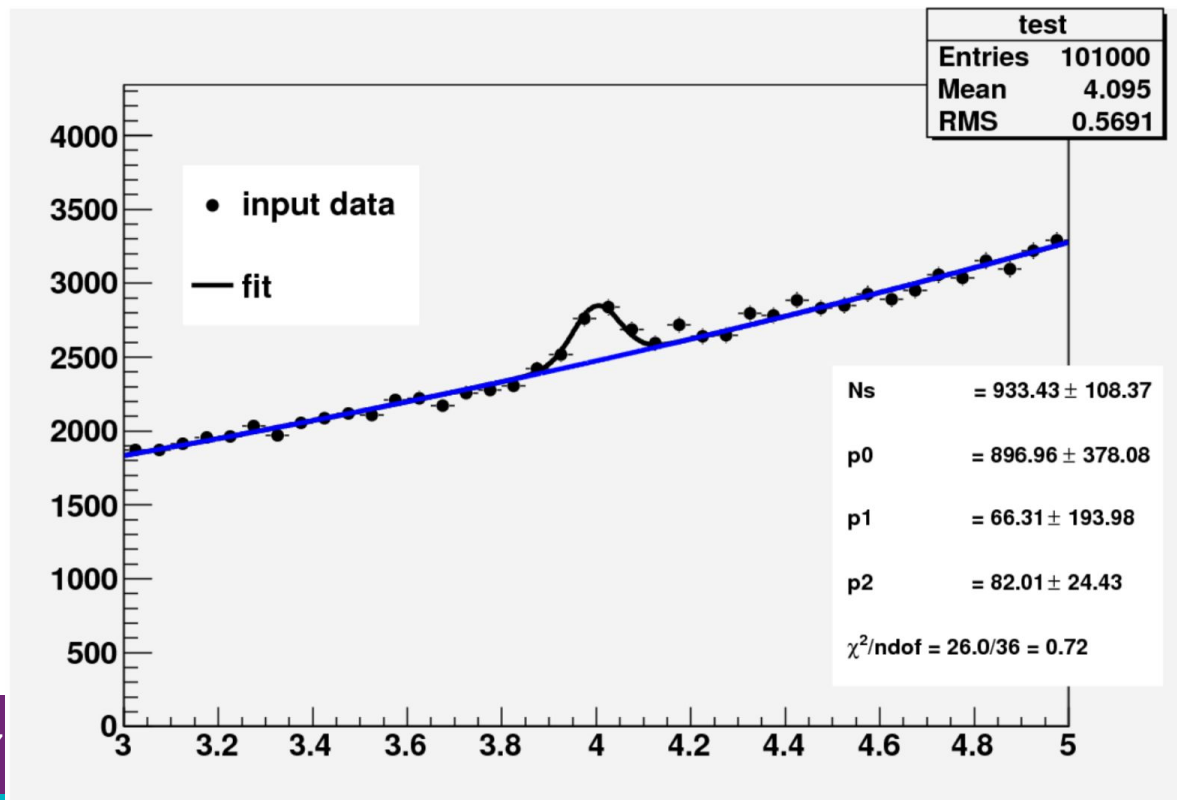
Fit function:
gauss+p1

$$\tilde{\chi}^2 = 34.7$$

Should we stop here?

$$\begin{aligned}
 \Delta\tilde{\chi}^2 &= -2 \ln \left(\frac{L(H_1)}{L(H_0)} \right) \\
 &= -2 \ln \left(\frac{L(g + p_1)}{L(g + p_0)} \right) = -2845.3
 \end{aligned}$$

How many parameters do I need? – Example



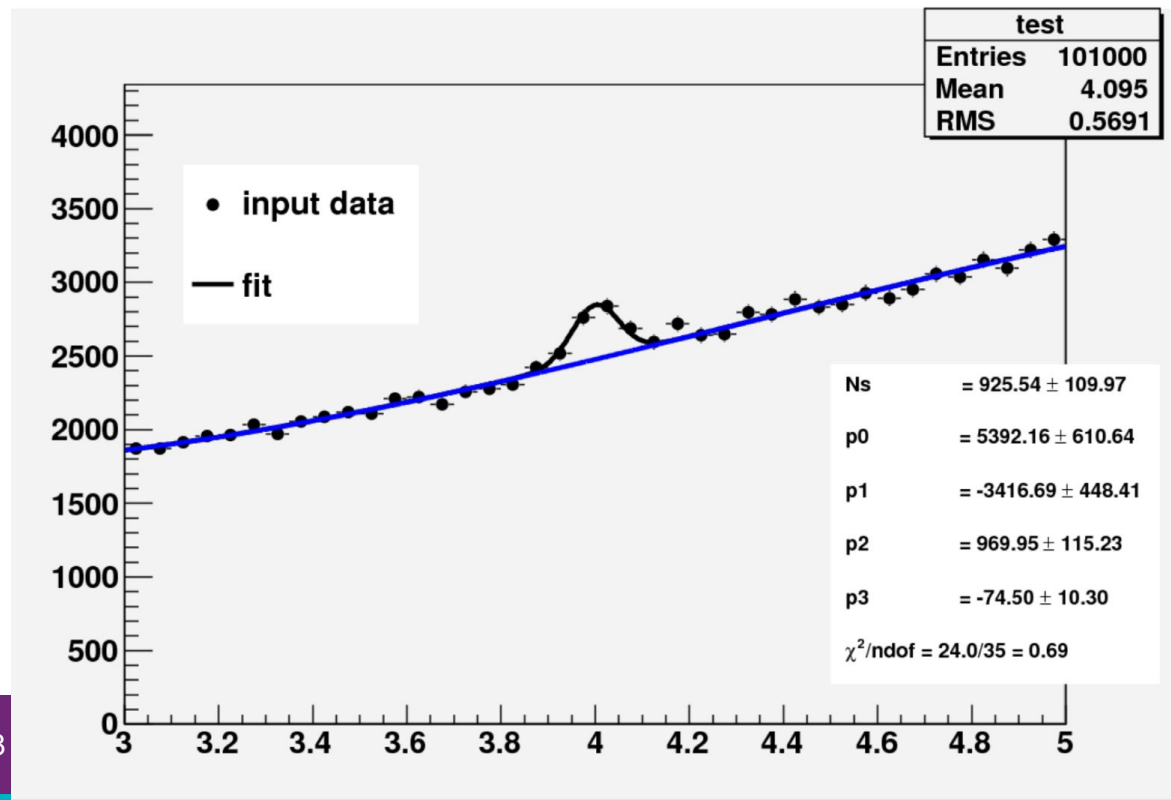
Fit function:
gauss+p2

$$\tilde{\chi}^2 = 26.0$$

$$\Delta\tilde{\chi}^2 = -8.7$$

Should we stop here?

How many parameters do I need? – Example



Fit function:
gauss+p3

$$\tilde{\chi}^2 = 24.0$$

$$\Delta\tilde{\chi}^2 = -2.0$$

Please stop!

How many parameters do I need? – Example

If H_0 correct then according to **Wilks' theorem**: $-\Delta\chi^2$ should follow a χ^2 function with $\text{ndf}=1$ (in asymptotic regime of large n)

g+p0

$$\tilde{\chi}^2 = 2880$$

Note: p-values for χ^2 : `TMath::Prob( $\chi^2$  obs, ndf)`

g+p1

$$\tilde{\chi}^2 = 34.7 \quad \Delta\tilde{\chi}^2 = -2845.3$$

g+p2

$$\tilde{\chi}^2 = 26.0 \quad \Delta\tilde{\chi}^2 = -8.7$$

P value = 0.003

Favoured over g+p1

g+p3

$$\tilde{\chi}^2 = 24.0 \quad \Delta\tilde{\chi}^2 = -2.0$$

P value = 0.15

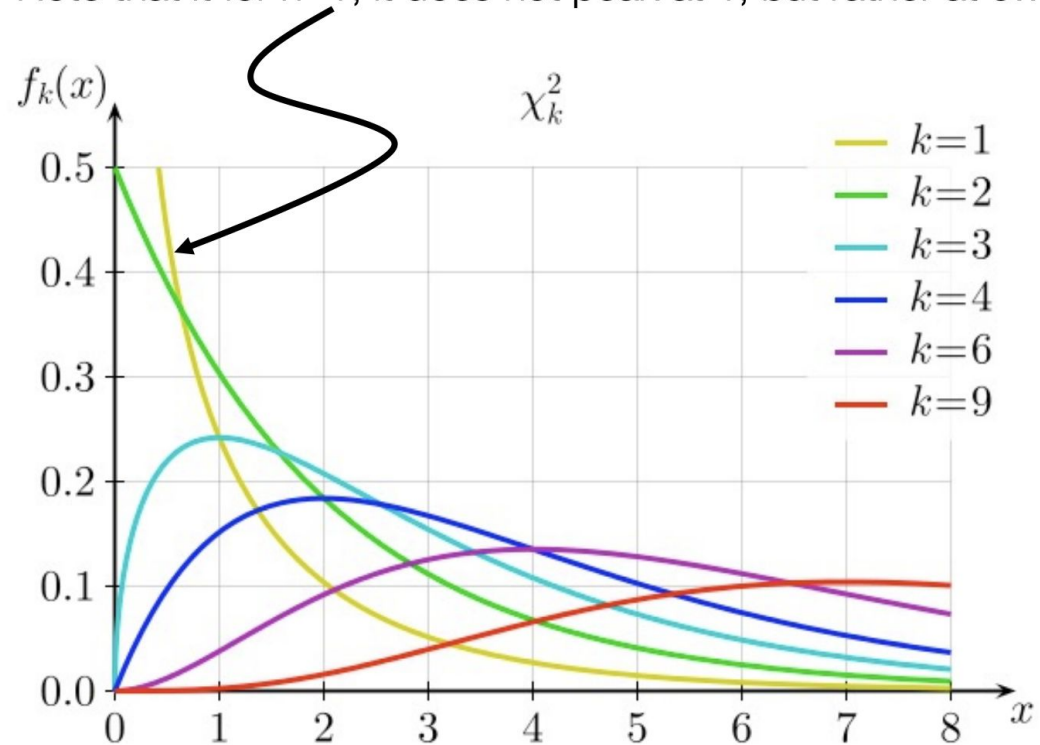
Not favoured over g+p2

What does a χ^2 distribution look like for $n=1$?

χ^2

Remember:

- Note that for $n=1$, it does not peak at 1, but rather at 0...



Improving on χ^2

Likelihood ratio is an improved χ^2 – S. Baker & R.D. Cousins, NIM 221 (1984) 437

- Still a single number
- “Optimal estimator” e.g. parameter estimation
- Requires a second hypothesis
- Allows taking uncertainties into account systematically!!
 - ◆ What if there’s a correlation?
 - ◆ What if there’s a systematic uncertainty

Alternatives to Likelihoods and χ^2

Kolmogorov-Smirnov test

F_c : Cumulative distribution function

F_e : Empirical distribution function

$$q_{GoF,KS} = \sup |F_c(x) - F_e(x)|$$

Anderson-Darling test

$$q_{GoF,AD} = n \cdot \int dF_e(x) \frac{(F_c(x) - F_e(x))^2}{F_e(x) \cdot (1 - F_e(x))}$$