

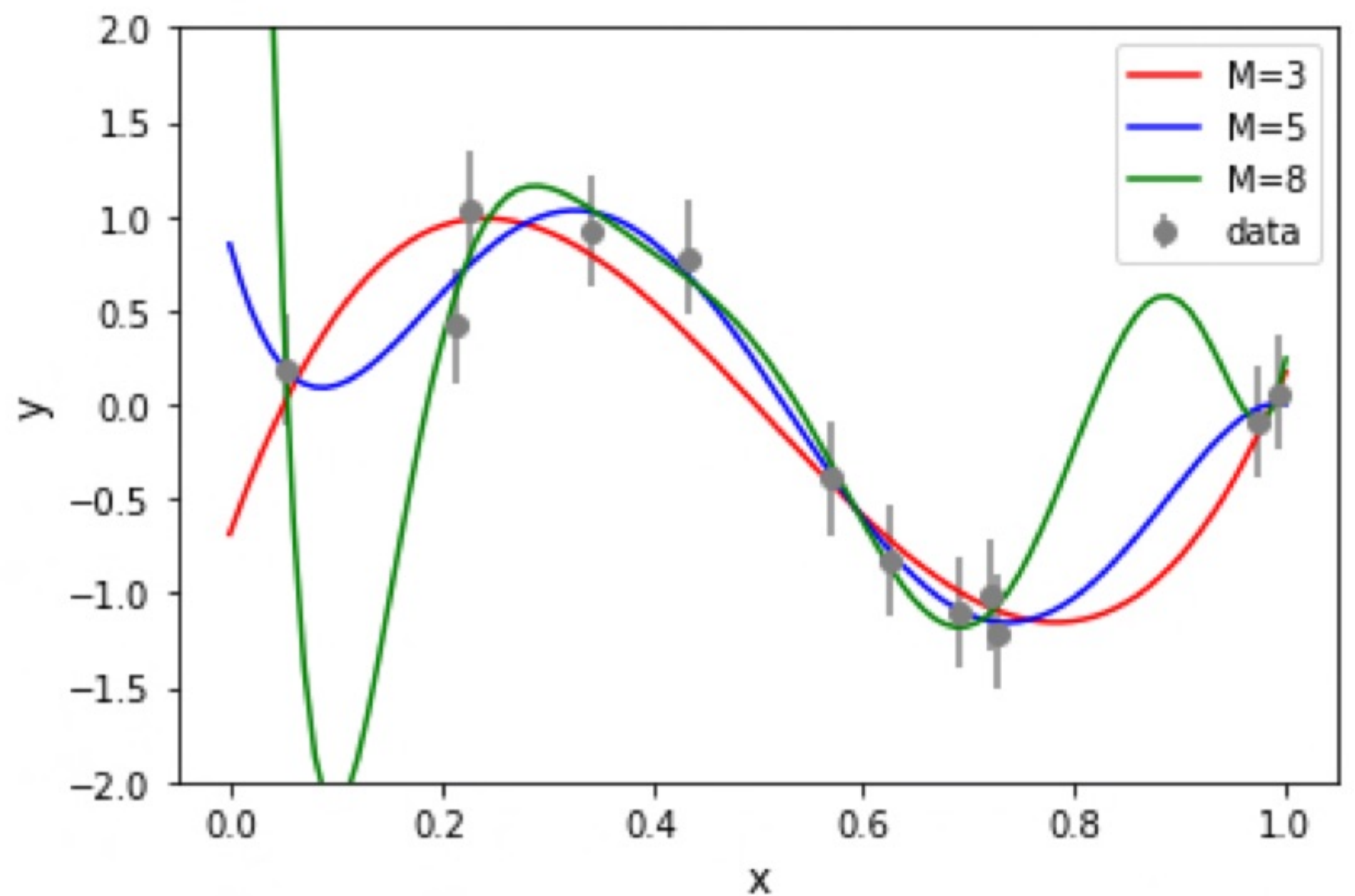
## • Model selection (both freq and Bayesian)

By looking at this plot  
what would be the best model?

3<sup>rd</sup>-order polynomial

5<sup>th</sup>-order one

or 8<sup>th</sup>-order one?



Note: model  $m=8$  fits the data better than the rest!

Intuitively model  $M=8$  is not the best though. This is because it seems like if  $M=8$  doesn't look like [the unknown function which generated this data]

↳ this is ultimately what we are after in statistical inference!

So if a new point appears after the fitting is done, it will probably be poorly predicted by  $M=8$

$M=8$  is said to overfit the data

Somehow the whole problem of model selection is to find the model which describes the data well (i.e. not doing underfitting either) while not doing overfitting



\* One of the most common ways to do model comparison in physics is by computing the  $\chi^2/\text{d.o.f.}$

The justification for that is based on what we saw previously about using  $\chi^2$  as a goodness-of-fit:

Assume that observations  $\{y_i\}$  follow a Gaussian distribution with mean

$$\mu_i \equiv \mu(x_i; \vec{\theta})$$

↪ a set of  $m$  parameters

and known std's  $\sigma_i$  } e.g. the errors of the measurements

then the quantity  $\chi^2 = \sum_{i=1}^N \left( \frac{y_i - \mu_i}{\sigma_i} \right)^2$

follows a  $\chi^2$ -distnb with  $N-m$  d.o.f.

So  $\chi^2$  is a measure of how likely it is that data comes from the model whose predictions are the  $\mu_i$ , in units of the stds.

On the other hand, noting that the mean of  $\chi^2$  distnb. is equal to the DOF, then the quantity

rule of thumb /  $\chi^2/\text{DOF}$  :

- if  $\gg 1$  : bad model, underfitting
- if  $\sim 1$  : fit is reasonable, within 1 std
- if  $\ll 1$  : overfitting, or the errors have been overestimated



This rule-of-thumb has several disadvantages:

1) if two models give

$$(\chi^2/\text{DOF})_1 \approx (\chi^2/\text{DOF})_2 \approx 1$$

Then one cannot really favour one model against the other

2) Suppose model #1 gives  $\chi^2 = 15$  and  $\text{DOF} = 10$ , whereas " #2 "  $\chi^2 = 150$  and  $\text{DOF} = 100$ ;

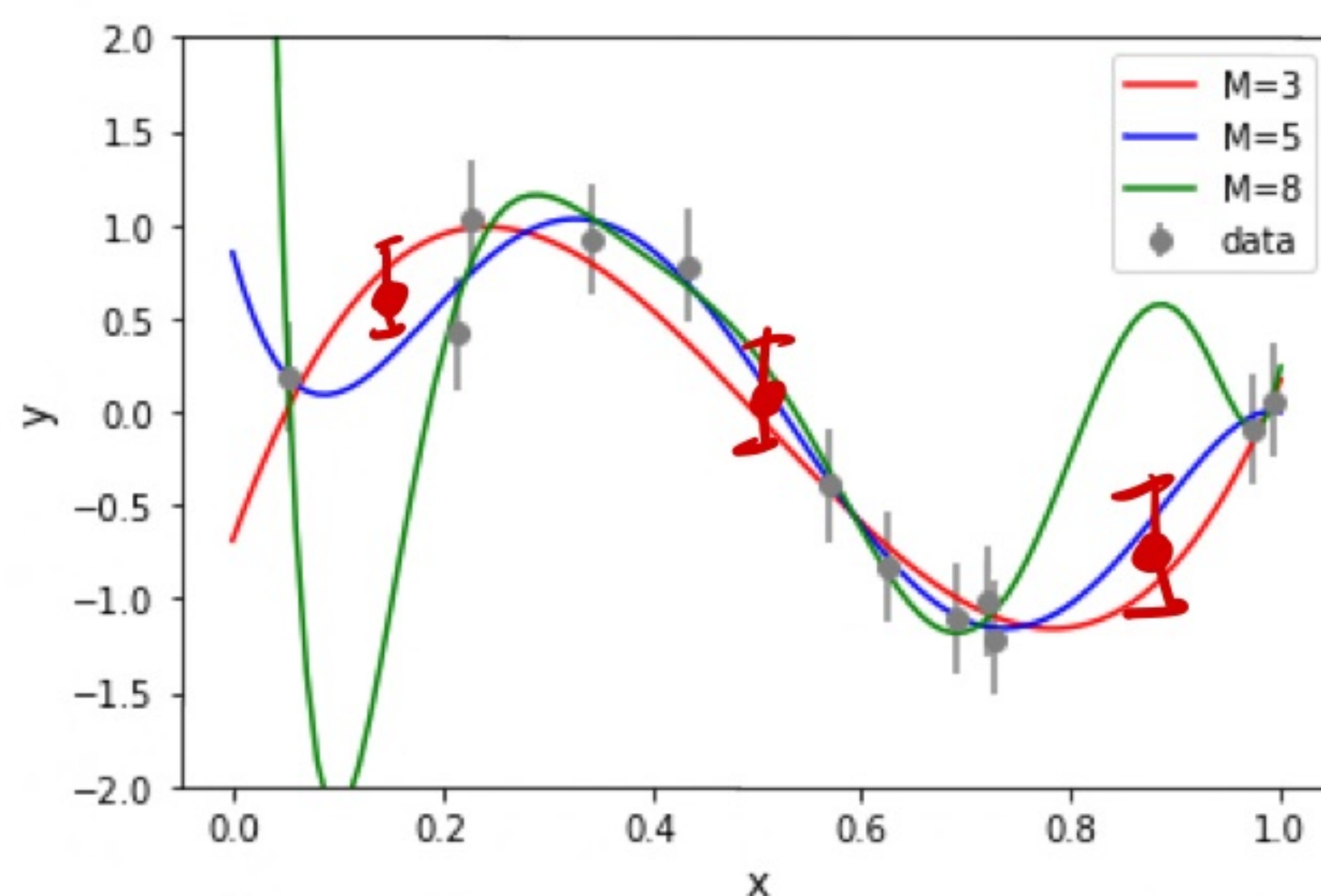
both have the same  $\chi^2/\text{DOF}$ , however

$$(\text{p-value})_1 = 0.13, (\text{p-value})_2 = 9 \times 10^{-4}$$

3) If data is not normally distributed, then the conclusions may be incorrect.

\* Another frequentist method, probably the most widely used by the ML community, is to fit the model with part of the data ("training data" in the ML jargon) and then evaluate how good the fitted model performs in the rest of the data ("test data")

clearly the  $M=8$  performs the worst regarding the test data points



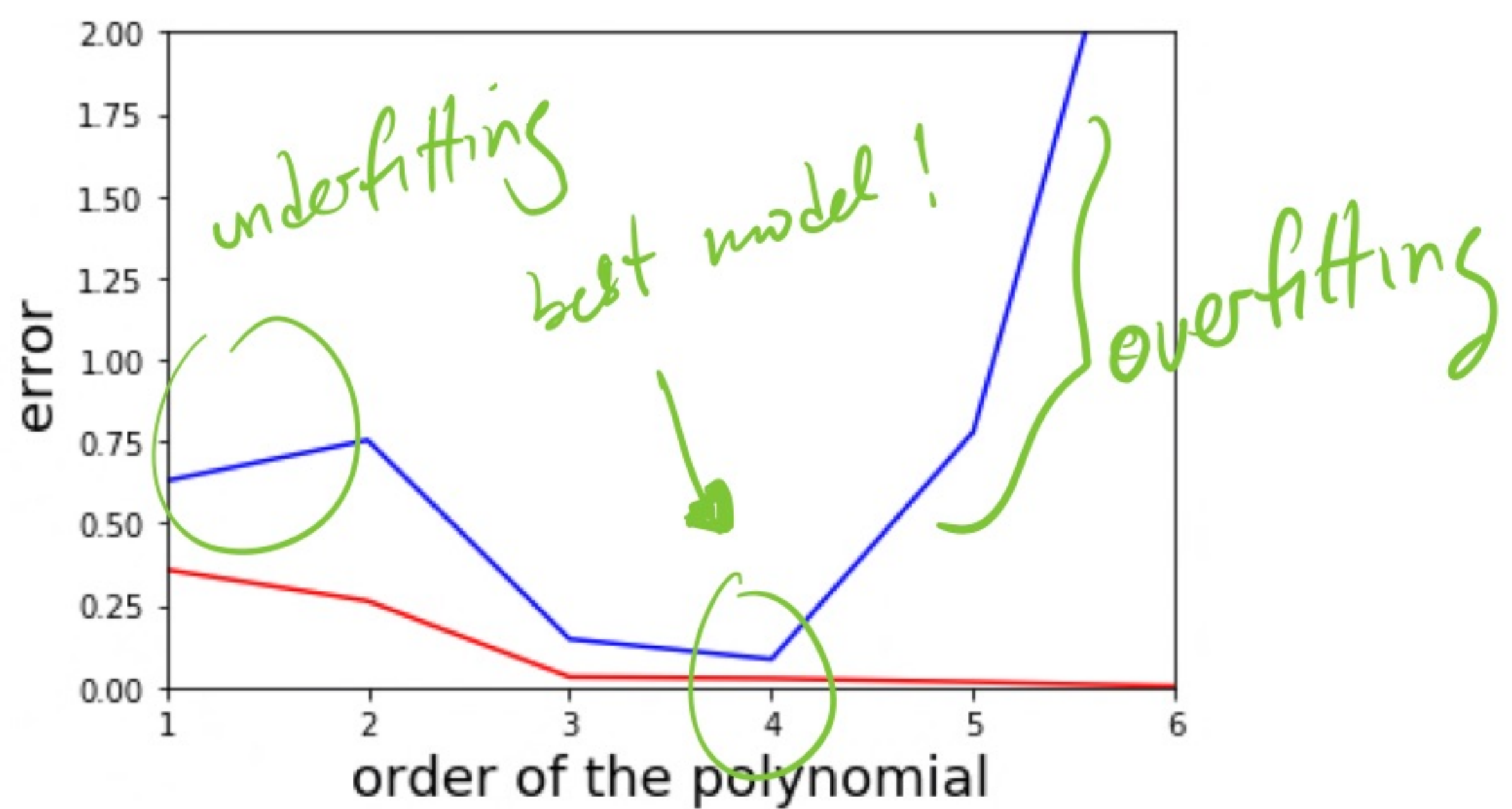


This is also an heuristic method, but it is much more powerful than the  $\chi^2/\text{DOF}$  in many ways; e.g.:

- By computing the discrepancy between data and predictions, using the likelihood  $p(y_i | \vec{\theta})$

one can choose the model which gives the lowest discrepancy (the largest likelihood) on the test data

- The method is not restricted to normally distributed data!



- Drawback: maybe not convenient for small datasets: one ideally would like to use enough "training" data for the fit to be good, but it may be insufficient in this case

- Note that when a model underfits, the bias (diff. between prediction and data) is large, while



the variance is small (predictions don't change much from point to point)

On the contrary, if a model overfits, the bias is small but the variance is high.

⇒ So a good model has a "Bias-Variance" trade-off.

Actually, it can be shown that the MSE

$$MSE = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2$$

is a linear combination of bias and variance:

Let  $y(x) = f(x) + \epsilon$  {  $f(x)$  is the true underlying function we are trying to model with  $\hat{f}(x)$  }

$\epsilon \sim N(0, \sigma^2)$

$$\text{Var}[\hat{f}] = \mathbb{E}[\hat{f}^2] - (\mathbb{E}[\hat{f}])^2$$
$$\text{Bias}[\hat{f}] = \mathbb{E}[\hat{f}] - f$$

{  $\mathbb{E}$  is taken wrt different datasets one could have access to }

Actually:

$$MSE = \mathbb{E}[(y - \hat{f}(x))^2]$$

↪ new point, unseen before ("test" point)

$$= \mathbb{E}[(f + \epsilon - \hat{f})^2]$$

$$= \mathbb{E}[(f + \epsilon - \hat{f} + \mathbb{E}[\hat{f}] - \mathbb{E}[\hat{f}])^2]$$



$$\begin{aligned}
 &= \mathbb{E} \left[ (f - \mathbb{E}[\hat{f}])^2 + \mathbb{E}[\varepsilon]^2 + (\mathbb{E}[\hat{f}] - \hat{f})^2 \right. \\
 &\quad + 2(f - \mathbb{E}[\hat{f}])\varepsilon + 2(\mathbb{E}[\hat{f}] - \hat{f})\varepsilon \\
 &\quad \left. + 2(f - \mathbb{E}[\hat{f}])(\mathbb{E}[\hat{f}] - \hat{f}) \right] \\
 &= (f - \mathbb{E}[\hat{f}])^2 + \mathbb{E}[\varepsilon]^2 + \mathbb{E}[(\mathbb{E}[\hat{f}] - \hat{f})^2] \\
 &\quad + 2 \left\{ (f - \mathbb{E}[\hat{f}]) + \mathbb{E}[\mathbb{E}[\hat{f}] - \hat{f}] \right\} \mathbb{E}[\varepsilon] \rightarrow 0 \\
 &\quad + 2(f - \mathbb{E}[\hat{f}]) \underbrace{\mathbb{E}[\mathbb{E}[\hat{f}] - \hat{f}]}_{= \mathbb{E}[\hat{f}] - \mathbb{E}[\hat{f}] = 0}
 \end{aligned}$$

$$= \underbrace{(f - \mathbb{E}[\hat{f}])^2}_{\text{Bias}^2} + \underbrace{\mathbb{E}[\epsilon]^2}_{\text{Var}[\epsilon] + (\mathbb{E}[\epsilon])^2} + \underbrace{\mathbb{E}[(\mathbb{E}[\hat{f}] - \hat{f})^2]}_{\text{Var}}$$

$$= \text{Bias}^2[\hat{f}] + \text{Var}[\hat{f}] + \sigma^2 \quad \text{□□}$$

## \* Bayesian Model Comparison

Expanded version of the Bayes theorem:

$$p(\vec{\theta} \mid \text{model}, \text{data}) = \frac{p(\text{data} \mid \vec{\theta}, \text{model}) p(\vec{\theta} \mid \text{model})}{p(\text{data} \mid \text{model})}$$

↙  
"model evidence"



Suppose 2 models:  $m_1$  and  $m_2$

then  $p(m_i | \text{data}) \propto \underbrace{p(\text{data} | m_i)}_{\text{evidence}} \underbrace{p(m_i)}_{\text{prior of the model itself (not of its params!)}}$

So if both models have equal prior,

$\Rightarrow$  the most probable model is obtained by computing the Bayes factors:

$$B = p(\text{data} | m_1) / p(\text{data} | m_2)$$

There is a criterion for  $B$ : The Jeffreys scale

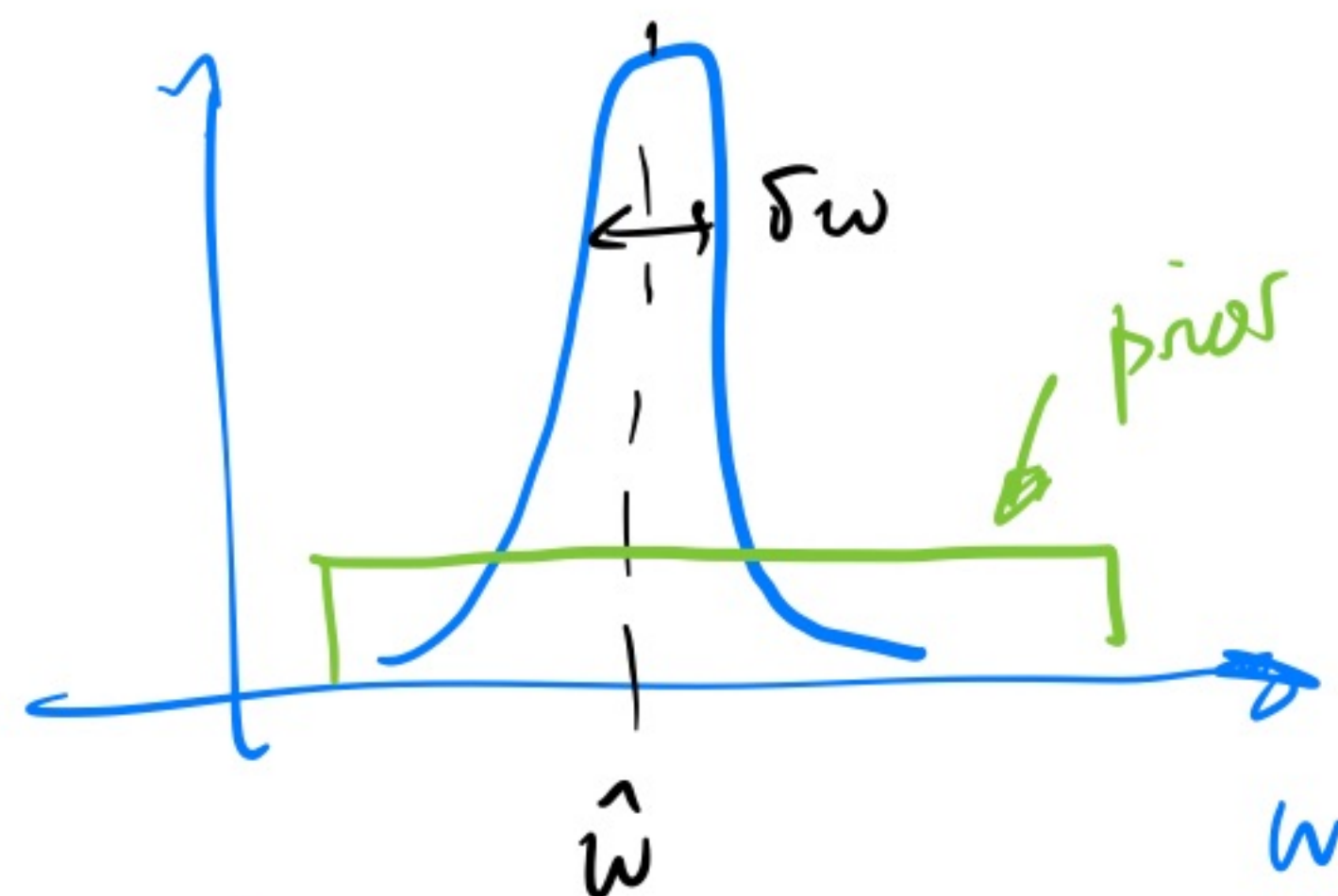
| $B$     | strength of evidence    |
|---------|-------------------------|
| $< 1$   | negative                |
| 1-3     | barely worth mentioning |
| 3-10    | substantial             |
| 10-30   | strong                  |
| 30-100  | very strong             |
| $> 100$ | decisive                |

- In order to get insight about the model evidence, let's suppose:
- a model with 1 single parameter  $w$
  - likelihood is sharply peaked around  $\hat{w}$  having width  $\delta w$
  - flat prior:  $p(w) = \frac{1}{\Delta w_{\text{prior}}}$



then

$$p(D|m) \approx \frac{p(D|\hat{w}) \delta w}{\Delta w_{\text{prior}}}$$



penalisation for models having large prior uncertainties

- Now assume a model with  $K$  parameters  $\vec{w} = \{w_k\}$

Suppose similar ratio  $\frac{\delta w_k}{\Delta w_{\text{prior}}^k} = \frac{\delta w}{\Delta w_{\text{prior}}} \quad \forall k$

then

$$p(D|m) \approx p(D|\hat{\vec{w}}) \left( \frac{\delta w}{\Delta w_{\text{prior}}} \right)^K$$

penalisation for models having too many parameters, (likelihood  $p(D|\hat{\vec{w}})$  improves though, so optimal model should have an intermediate complexity)

Drawback: Computing the evidence is typically intractable!

(a whole topic on approximate Bayesian inference)