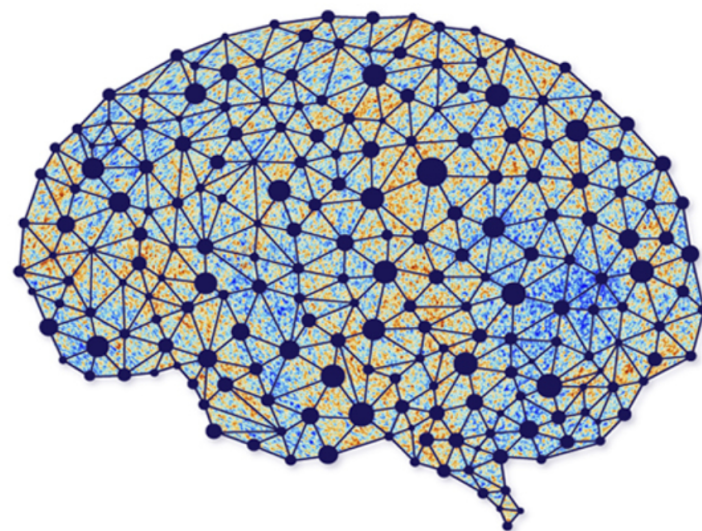


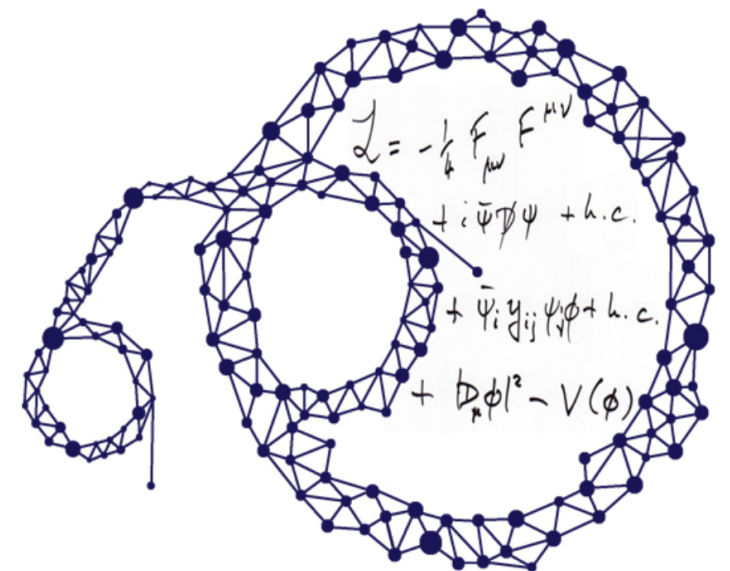
Physics 361 - Machine Learning in Physics

Lecture 3 – Background

Jan. 28th 2025



AI
∩
Universe



Moritz Münchmeyer

Unit 1: Background

1.2 Classical Statistics and Data Analysis Background

Sources:

- Cowan - Statistical data analysis
- also mostly covered in deeplearningbook.org

Estimators

- A (point) estimator makes a prediction for some quantity of interest λ .

E.g. estimator of the mean height $\bar{h}$

We write $\hat{\lambda}$, where "hat" means estimator.

$$\hat{\bar{h}} = \frac{1}{N} \sum_i h_i$$

- Given a dataset $\{x^{obs}\}$ drawn / sampled from a PDF $P(x)$, an estimator is some function

$$\hat{\lambda} = \mathcal{E}[\{x^{obs}\}]$$

↖ some function, "guessed" or derived

- An optimal estimator is
 - unbiased: $\langle \hat{\lambda} \rangle = \lambda$ "correct on average"
 - minimum variance: $\text{Var}[\hat{\lambda}]$ is as small as possible (smallest error)

Some common estimators

notation: $\bar{X}$
"bar" means mean

(sample) mean: $\hat{\bar{X}} = \frac{1}{n} \sum_{i=1}^n x_i^{\text{obs}}$

variance: $\hat{V} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \hat{\bar{X}})^2$

covariance: $\hat{cov} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \hat{\bar{X}})(y_i - \hat{\bar{Y}})$

These estimators have a variance, e.g.

$$V[\hat{\bar{X}}] = \frac{\sigma^2}{n} \quad \text{where } \sigma^2 \text{ is the variance of } X$$

Likelihood, Posterior, Prior

- The Likelihood is the probability of measuring data $\vec{d}$ given a model M with parameters $\vec{\lambda}$.

$\mathcal{L}(\vec{d} | M, \vec{\lambda})$ — often not explicitly written
often used as the Loss function
in machine Learning.

- The Posterior is the probability of parameters $\vec{\lambda}$ given observed data $\vec{d}$.

$$P(\vec{\lambda}, M | \vec{d})$$

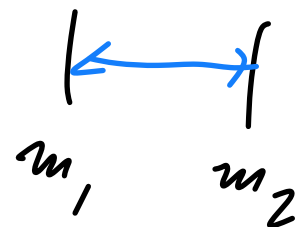
e.g. M : standard model
 $\vec{\lambda}$: Higgs mass
 $\vec{d}$: collision data

- We can use Bayes theorem to get P from L

$$P(\vec{\lambda} | \vec{d}) = \frac{\mathcal{L}(\vec{d} | \vec{\lambda}) P(\vec{\lambda})}{P(\vec{d})}$$

- The prior $P(\vec{\lambda})$ is the probability of $\vec{\lambda}$ before we perform the measurement.

E.g. • flat in some interval



- Gaussian around some prior measurement.

Prior is important if the data is not very constraining.

- Finally we have the evidence

$P(\vec{d})$ is the probability of the data under model M for ANY parameters $\vec{\lambda}$.

$$P(\vec{d}) = \int \mathcal{L}(d|\vec{\lambda}) P(\vec{\lambda}) d\vec{\lambda}$$

Often used in model comparison.

Maximum Likelihood and Maximum A Posteriori

- Combining the concepts of estimators and Likelihoods we define the

Maximum Likelihood estimator

$$\hat{\lambda}_{ML} = \arg \max_{\lambda} \mathcal{L}(\vec{d} | \lambda)$$

Sometimes we can do it analytically:

$$\left. \frac{\partial \mathcal{L}(\vec{d} | \lambda)}{\partial \lambda} \right|_{\lambda = \hat{\lambda}} = 0$$

If not analytically tractable use a numerical optimizer.

- We define the **maximum a posteriori estimator**

$$\hat{\lambda}_{MAP} = \arg \max_{\lambda} \mathcal{P}(\lambda | \vec{d}) \quad (\text{MAP})$$

Gaussian Likelihood

Gaussian Likelihoods are often good approximations of data in physics.

Example: We measure a person's weight w .
To get an uncertainty we measure 100 times.

We assume that the data is described by

$$d_i = w + n_i$$

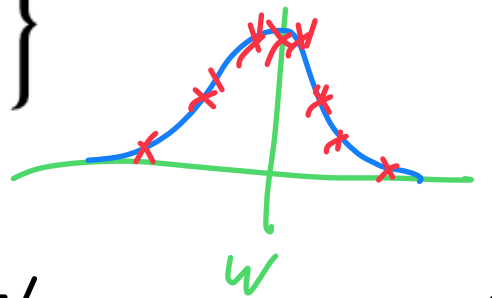
truth

Gaussian noise,
with variance
 σ_w^2

We want to measure the parameters w and σ_w^2 .

For a single observation the Likelihood is:

$$\mathcal{L}(d|w, \sigma_w) \equiv P(d|w, \sigma_w) = \frac{1}{\sqrt{2\pi\sigma_w^2}} \exp\left\{-\frac{(d-w)^2}{2\sigma_w^2}\right\}$$



For m independent measurements $\mathcal{L}$ is the product:

$$\mathcal{L}(\{d_i\}_{i=1}^m | w, \sigma_w) = \frac{1}{(2\pi\sigma_w^2)^{m/2}} \exp\left\{-\frac{\sum_{i=1}^m (d_i - w)^2}{2\sigma_w^2}\right\}$$

We get the posterior from Bayes theorem:

$$P(w, \sigma_w | \{d_i\}) = \frac{\mathcal{L}(\{d_i\} | w, \sigma_w) P(\lambda)}{P(\{d_i\})}$$

If the prior is flat we can work with the maximum Likelihood estimator.

The max. likelihood estimator for w is:

$$\frac{\partial \mathcal{L}}{\partial w} = \frac{\sum_{j=1}^m (d_j - w)}{\sigma_w^2 (2\pi \sigma_w^2)^{m/2}} \exp \left\{ -\frac{\sum_{i=1}^m (d_i - w)^2}{2\sigma_w^2} \right\}$$

const. *≠ const.*

$$\frac{\partial \mathcal{L}}{\partial w} \stackrel{!}{=} 0 \quad \Leftrightarrow \quad \sum_{j=1}^m (d_j - w) = 0$$

solve for w :

$$w = \hat{w} = \frac{1}{m} \sum_{i=1}^m d_i$$

This is the mean,
as expected.

We could also do $\frac{\partial \mathcal{L}}{\partial \sigma_w}$ to find $\hat{\sigma}_w$.

While here the result is "obvious", this is a general procedure to find estimators given a model (= Likelihood).

Sampling the posterior: MCMC

- We have Learned how to get the posterior for parameters $\hat{\lambda}$ given data $\vec{d}$.

- It is given by
$$P(\lambda | d) = \frac{P(d | \lambda) P(\lambda)}{P(d)}$$

Often we only need the unnormalized posterior

$$P(\lambda | d) \propto P(d | \lambda) P(\lambda)$$

The evidence does not depend on $\vec{\lambda}$ and can thus be ignored when finding λ .

- There is still a problem to deal with: computational difficulty.

Example: Assume $P(\lambda | \vec{d})$ depends on 20 parameters.

How do we work with such a high dimensional posterior?

If we evaluate it on a grid with 100 points per dimension we'd need $(100)^{20}$ evaluations.

This cannot be done.

Instead one works with such posteriors by sampling from them using a method called Markov Chain Monte Carlo (MCMC).

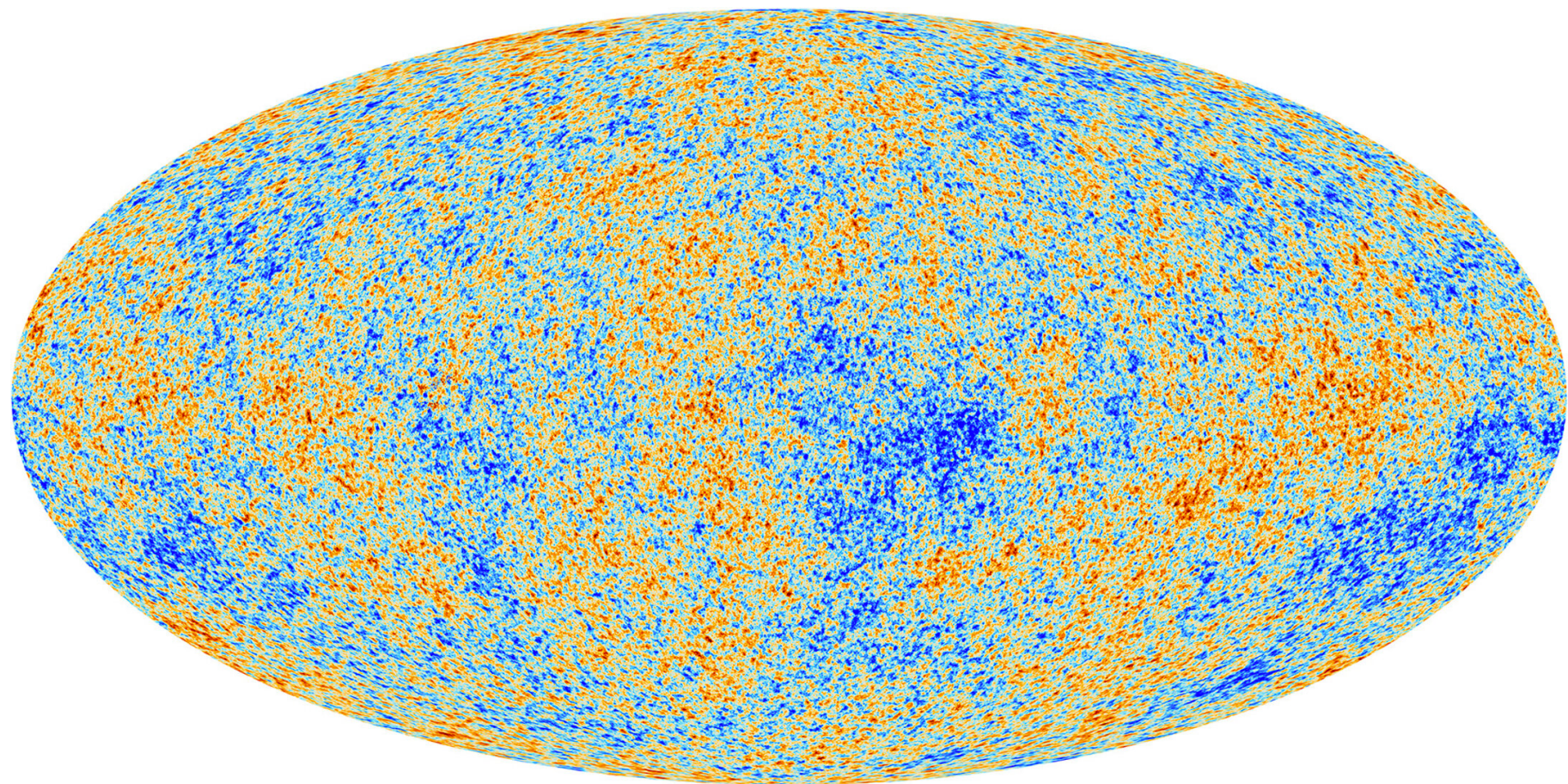
Perhaps covered later in this course.

Unit 1: Background

1.3 An example of data analysis in physics that uses these concepts

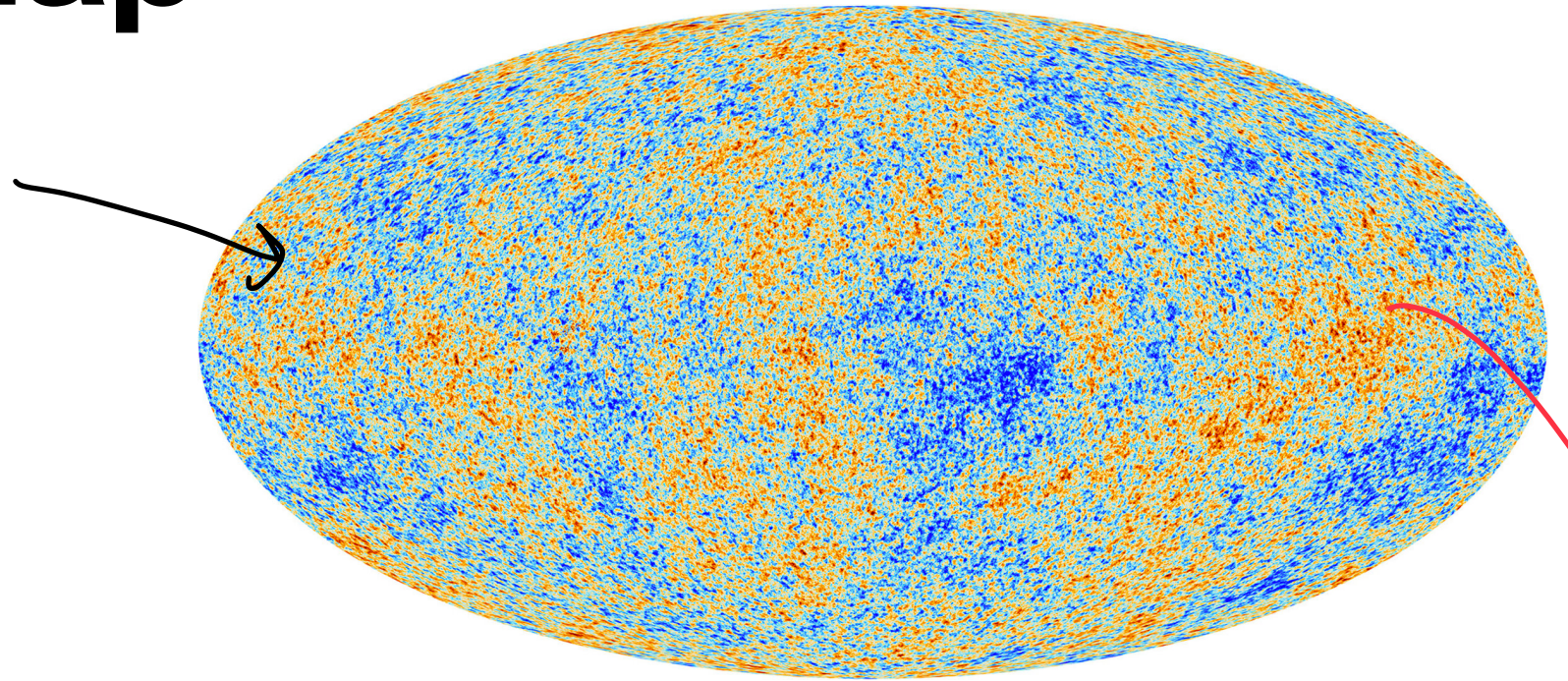
A physics application of these concepts: The CMB power spectrum

- I want to show you an application of these ideas to real physics. The CMB power spectrum analysis is one of the jewels of physics, telling us much of what we know about the history of the universe.
- Of course this is a complicated topic and I can only give you a brief idea.
- The Cosmic Microwave Background is a radiation that permeates the universe. It has a temperature of about 3K and has been measure extremely precisely.



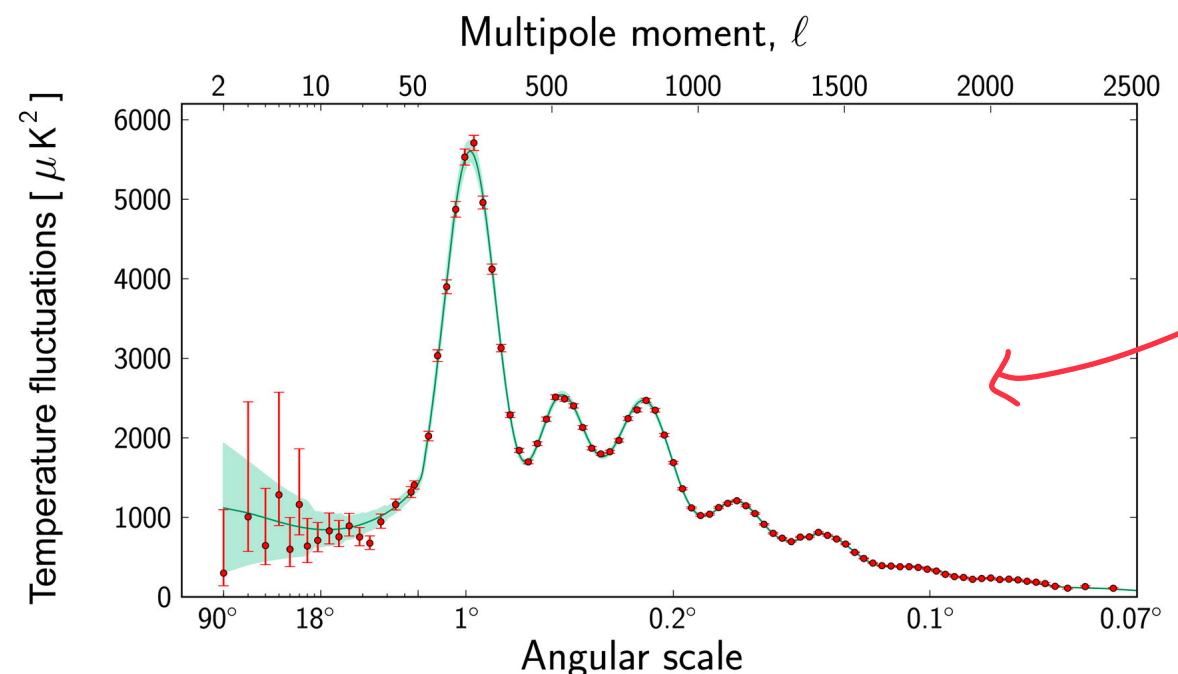
The CMB map

CMB temperature
field $T^{\text{CMB}}(\hat{n})$



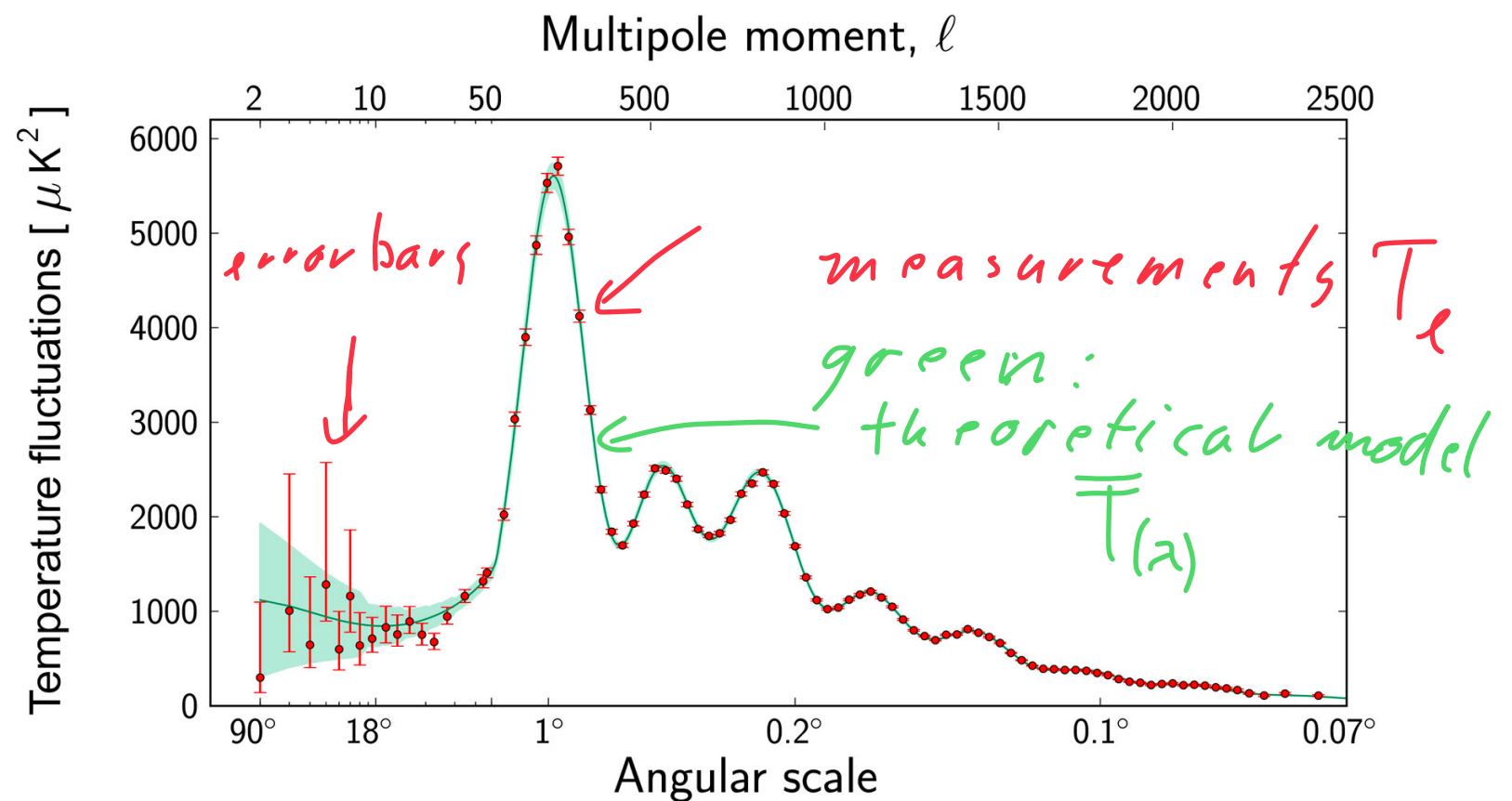
Rather than analyzing this map directly, one first compresses the data into an optimal summary statistic:

CMB power spectrum



The dataset

- Our measurements will be the power spectrum as a function of scale. It can be extracted from the previous map by taking its Fourier transform and squaring the mode amplitudes.
- Our data points are:



data set: $\{T_\ell\}$ with some error bars

↑
Fluctuation amplitudes of CMB temperature.

The likelihood

- We will make a Gaussian likelihood.
- Our theory model is that the mean curve (green) is a known function of cosmological parameters, which we want to measure.

Theoretical model (green curve):

$\bar{T}_e(\vec{\lambda})$ depends on cosmological parameters, e.g.:

Ω_c : Dark matter density

Ω_b : baryon density

It can be calculated from the Laws of nature (e.g. General Relativity)

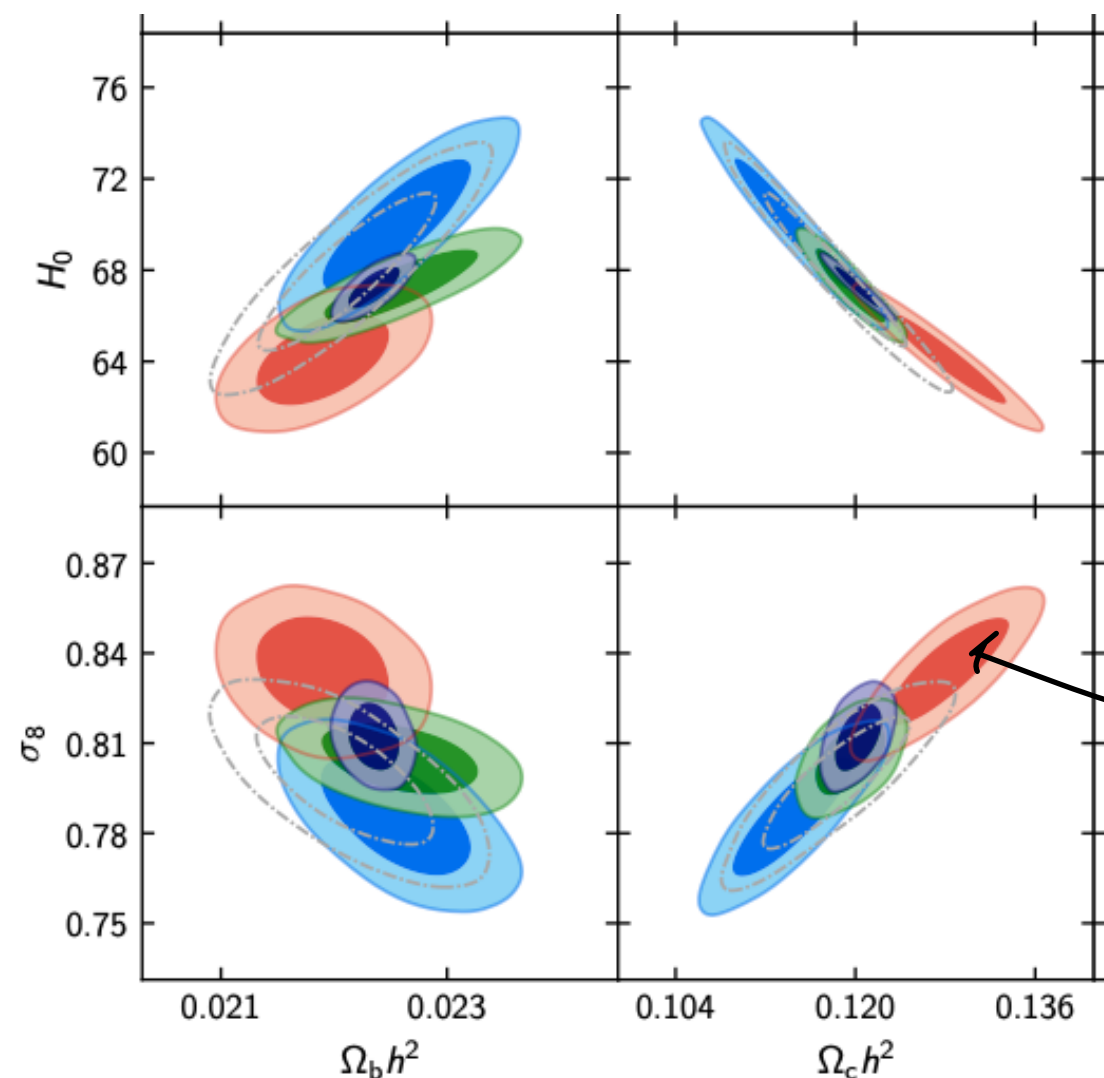
$\Rightarrow$ Gaussian Likelihood

$$\ln \mathcal{L}(\{T_e\} | \lambda) \propto \frac{1}{2} \sum_{\ell=1}^n \frac{(T_{e,\ell} - \bar{T}_e(\vec{\lambda}))^2}{2\sigma_{\ell}^2}$$

$\Rightarrow$ Maximize Likelihood to find theory parameters $\vec{\lambda}$ that fit the data best.

Sampling the posterior

- Since we have a likelihood, we can now use Bayes theorem and get the posterior.
- Then, one can sample from the posterior, to get the likely values of the desired cosmological parameters.
- The result looks are Monte Carlo plots like this:



H_0 : Hubble parameter
 Ω_c : dark matter
 Ω_b : baryons

one and two
sigma regions
of the posterior
 $P(\mathbf{z} | \text{data})$

Unit 1: Background

1.4 Information theory background

What is information theory?

- Quantify how much information is in a given signal.

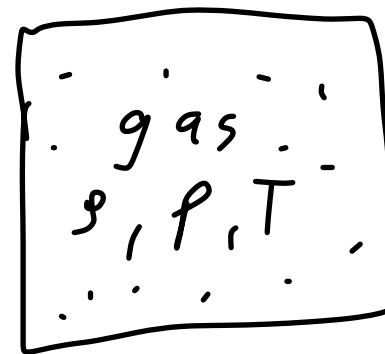
→ Optimal compression

- The information in a signal depends on the probability of the signal.
 - A certain signal/event contains no information.
 - A highly unlikely event contains a lot of information.
- In machine Learning, information theory is used to characterize PDFs and their similarity.

Entropy

- In statistical physics (thermodynamics) the **entropy** is given by

$$S = (K_B) \log \Omega$$



↑ here we set it to 1

Ω : number of microstates (equally likely)

- We can re-write this as

$$S = -\log p \quad \text{where} \quad p = \frac{1}{\Omega}$$

is the probability of each m. state

- The general definition of entropy

Shannon entropy

$$S = - \sum_i p_i \log p_i$$

- $\log$ here is base e $\rightarrow$ entropy is in "nats"
(base 2 $\rightarrow$ " in bits)

Properties of the entropy

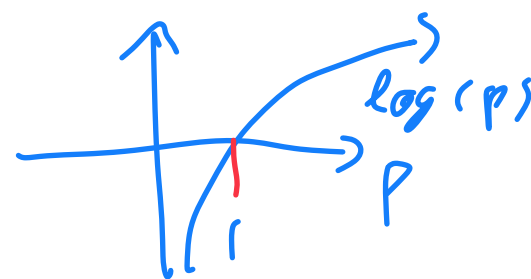
- discrete PDF

$$S = - \sum_i p_i \log(p_i)$$

$-\log(p_i)$ is the „self-information“ of event i

$$= - \mathbb{E}_{x \sim P(x)} [\log(P(x))]$$

A unlikely event has a large self info.



self information $I(x) = -\log(P(x))$

A certain event has self inform. 0

maximum possible entropy is

- continuous PDF

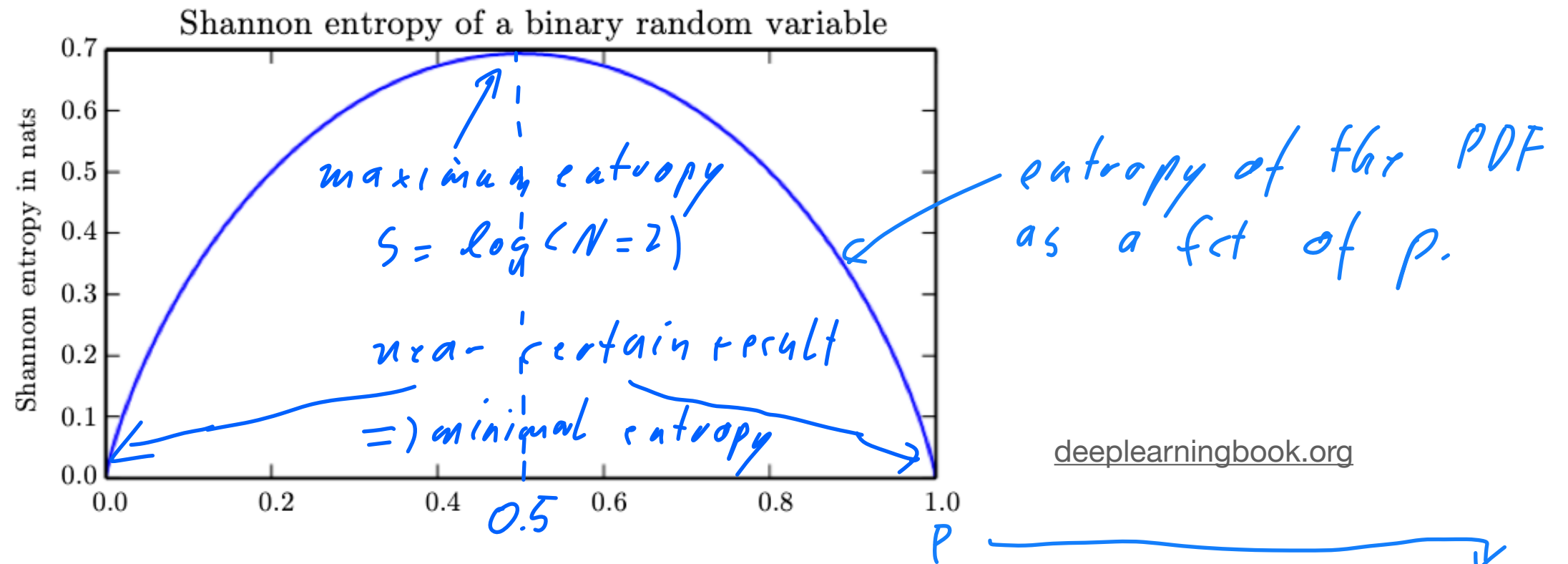
$$S = - \int dx p(x) \log(p(x))$$

$$= - \mathbb{E}_{x \sim P(x)} [\log(P(x))]$$

$$S = \log(N)$$

More uniform distributions have a higher entropy.
(spread out)

Example: Shannon entropy of a binary variable



Example: binary random variable $X: \begin{cases} 1 : p \\ 0 : 1-p \end{cases}$

calculate the entropy

$$S = - \sum_i p_i \log(p_i)$$

$$= - (1-p) \log(1-p) - p \log(p)$$

Course logistics

- **Reading for this lecture:**
 - **For example:** [Deeplearningbook.org](https://deeplearningbook.org) chapter 3 and 5.
- **Problem set:** First problem set this week.