

# Reconstruction in high-energy Cherenkov neutrino telescopes

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free
- Lecture 2: Likelihood methods in detail: **Poisson process likelihoods**
  - > obtain response-PDFs (Analytic / B-Splines)
  - > Tricks for systematics circumvention (1st hit likelihood)
  - > Parametrization: Stochastics
  - > Uncertainties

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free
- Lecture 2: Likelihood methods in detail: **Poisson process likelihoods**
  - > obtain response-PDFs (Analytic / B-Splines)
  - > Tricks for systematics circumvention (1st hit likelihood)
  - > Parametrization: Stochastics
  - > Uncertainties
- Lecture 3: All things **neural networks**
  - > Neural network-based likelihoods
  - > Inverse likelihood-free methods
  - > Encodings and symmetries (Transformers etc.)
  - > Putting it all together

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free
- Lecture 2: Likelihood methods in detail: **Poisson process likelihoods**
  - > obtain response-PDFs (Analytic / B-Splines)
  - > Tricks for systematics circumvention (1st hit likelihood)
  - > Parametrization: Stochastics
  - > Uncertainties
- Lecture 3: All things **neural networks**
  - > Neural network-based likelihoods
  - > Inverse likelihood-free methods
  - > Encodings and symmetries (Transformers etc.)
  - > Putting it all together

Classical statistical tools



Neural networks

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free
- Lecture 2: Likelihood methods in detail: **Poisson process likelihoods**
  - > obtain response-PDFs (Analytic / B-Splines)
  - > Tricks for systematics circumvention (1st hit likelihood)
  - > Parametrization: Stochastics
  - > Uncertainties
- Lecture 3: All things **neural networks**
  - > Neural network-based likelihoods
  - > Inverse likelihood-free methods
  - > Encodings and symmetries (Transformers etc.)
  - > Putting it all together

Classical statistical tools

Information theory

Neural networks

# Overview

- Lecture 1: Introduction and **high-level overview**
  - > Bayes theorem
  - > Frequentist vs Bayesian
  - > Information-theoretic grounding (KL-divergence of the “joint”)
  - > Theme of these lectures: likelihood vs likelihood-free
- Lecture 2: Likelihood methods in detail: **Poisson process likelihoods**
  - > obtain response-PDFs (Analytic / B-Splines)
  - > Tricks for systematics circumvention (1st hit likelihood)
  - > Parametrization: Stochastics
  - > Uncertainties
- Lecture 3: All things **neural networks**
  - > Neural network-based likelihoods
  - > Inverse likelihood-free methods
  - > Encodings and symmetries (Transformers etc.)
  - > Putting it all together

Classical statistical tools

**Revolution:**  
**Likelihood-**  
**Free inference**

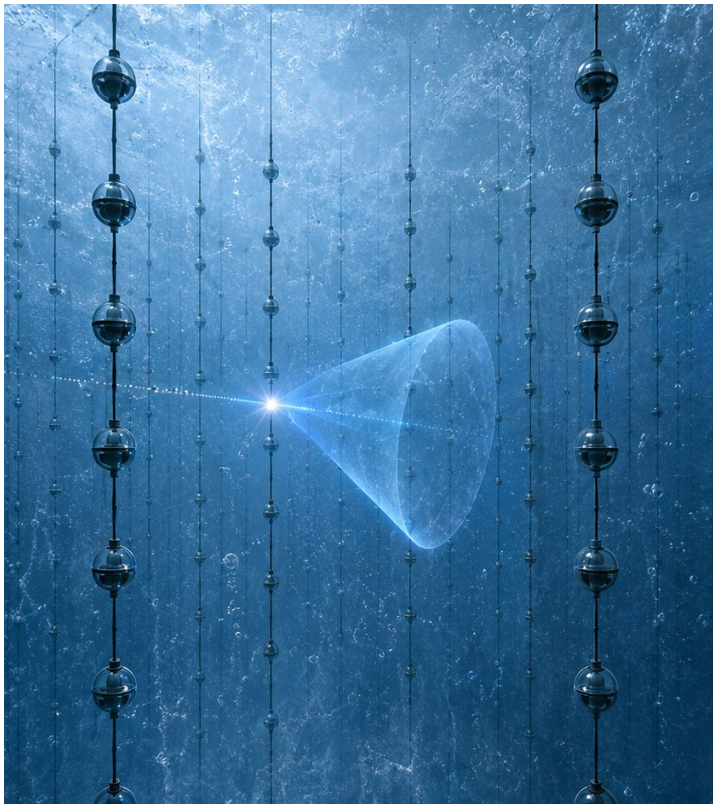
**Information theory**

**Neural networks**

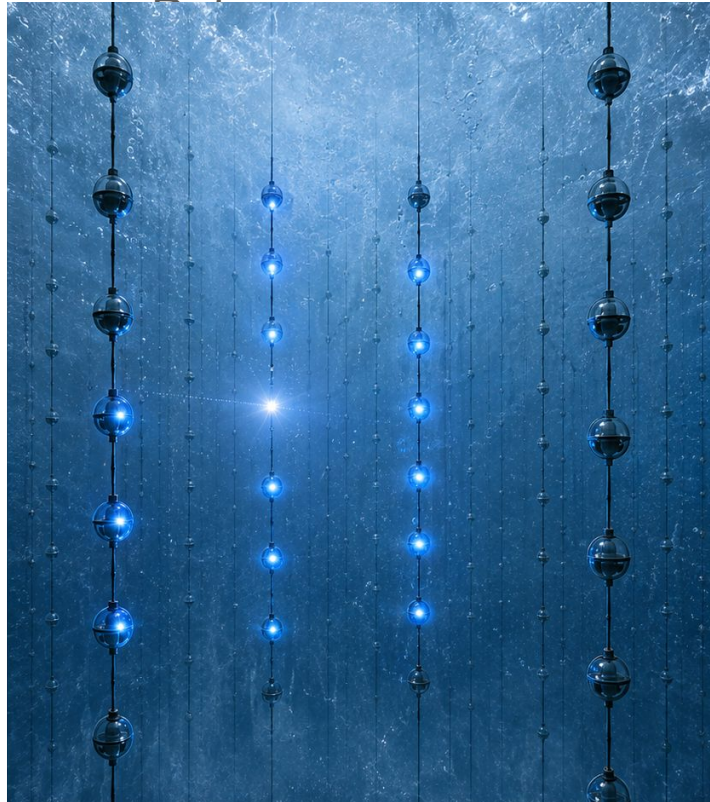
# Learning objectives lecture 1:

- Bayes' Theorem and the “joint” distribution
- Frequentist / Bayesian parameter estimation
- Information theory: KL divergence and its interpretation
- likelihood-based vs likelihood-free methods

Interaction happens...

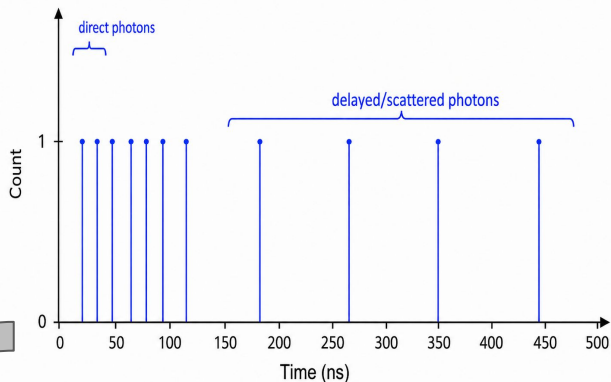
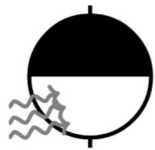


What we actually see later



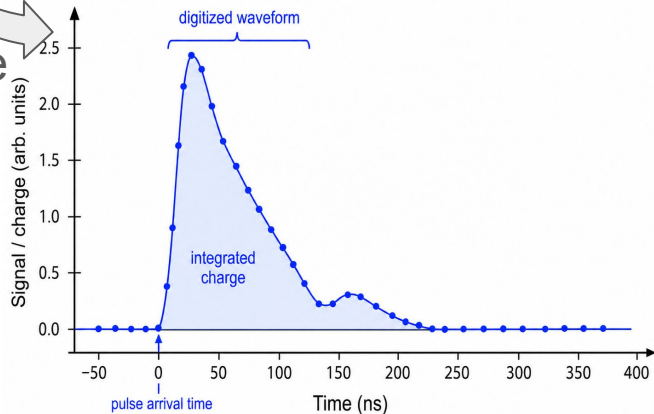
## Physical photon hit times

Photons hit DOM



## Schematic charge response

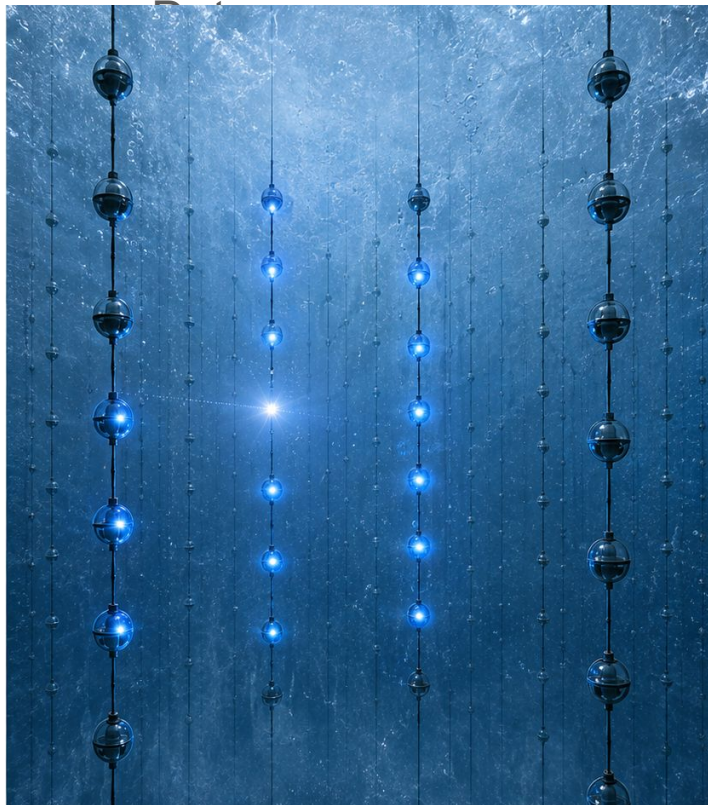
PMT  
Response



Data "x":

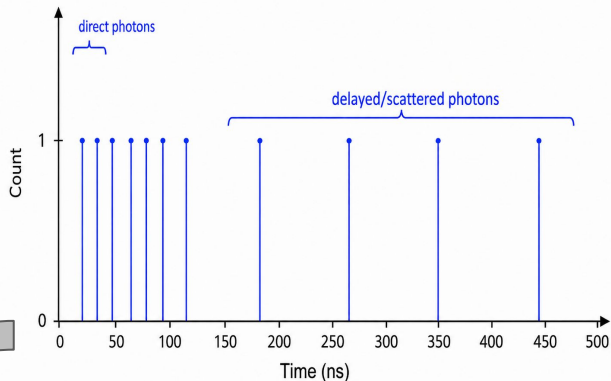
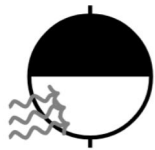
Total charge, charge vs time, leading edge

## What we actually see later

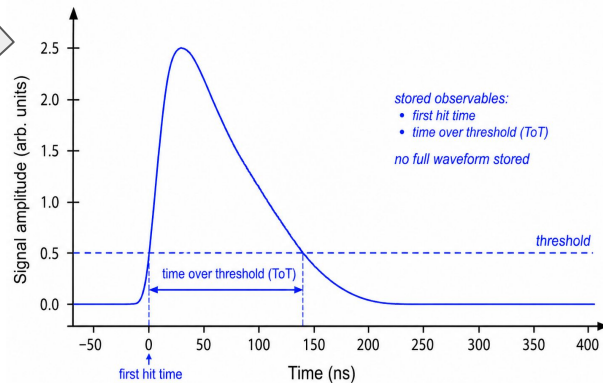


## Physical photon hit times

Photons hit DOM



## Schematic charge response (KM3NeT)

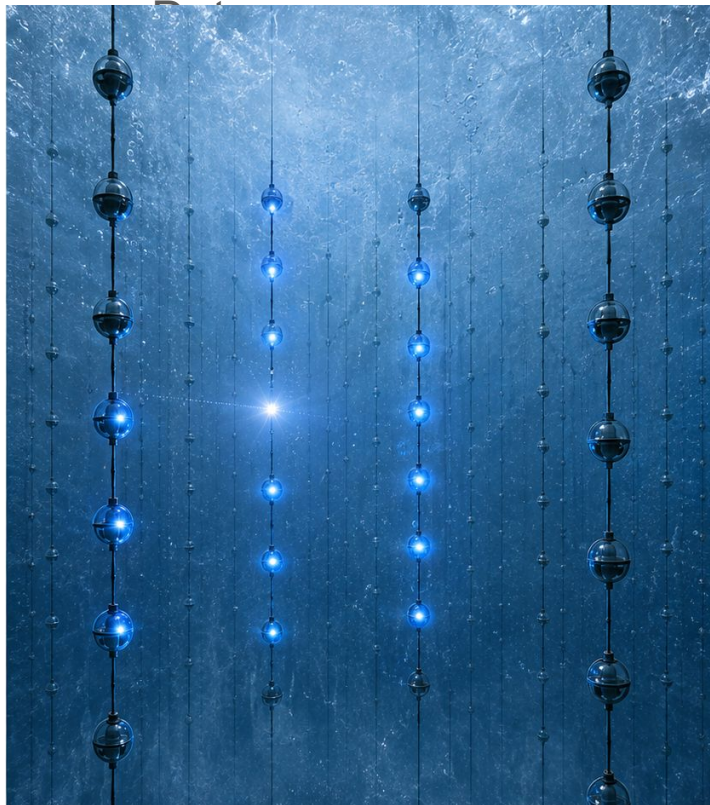


PMT  
Response

Data "x":

leading edge, ToT

## What we actually see later

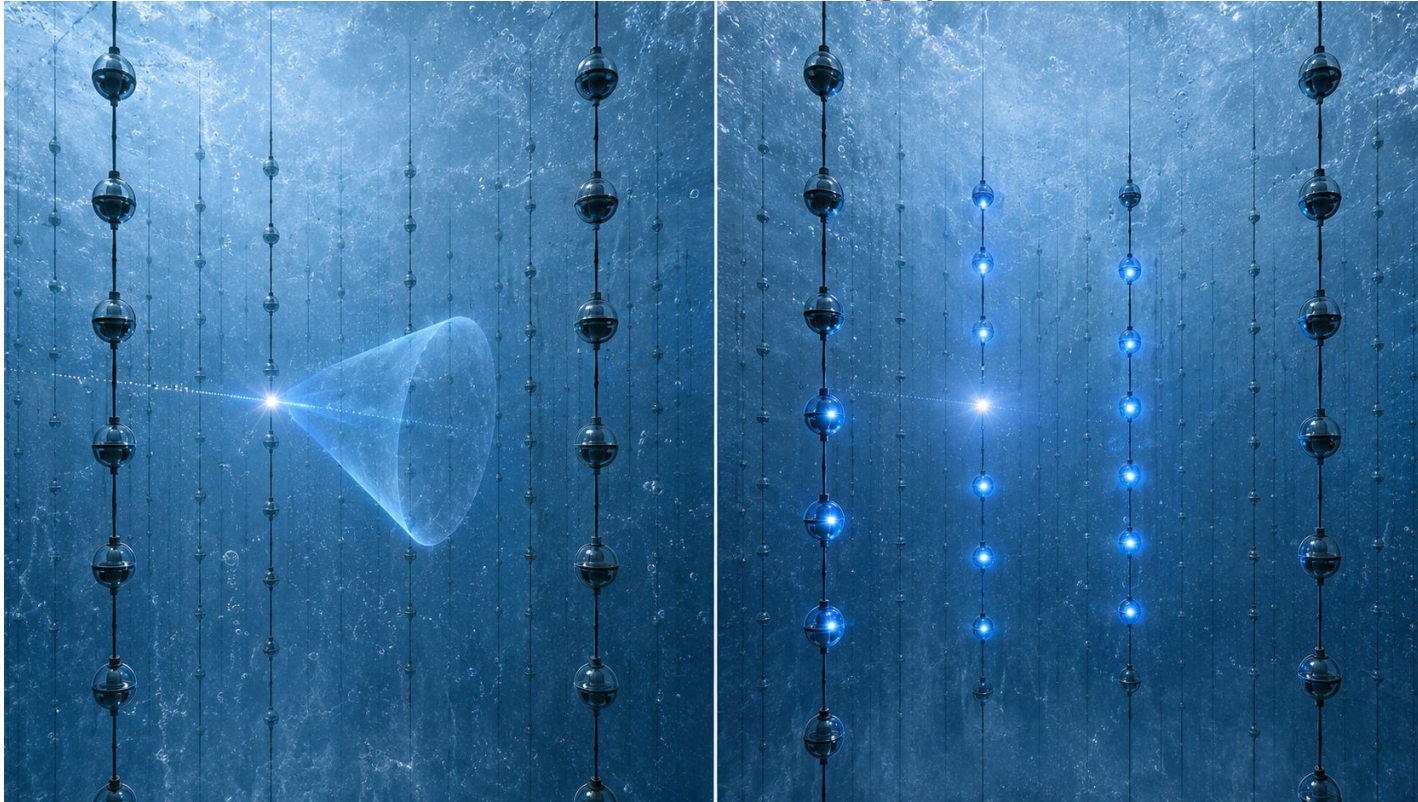


## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

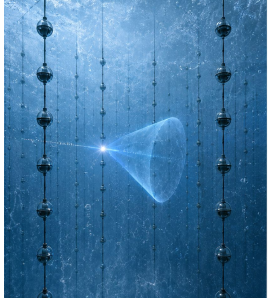
## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules



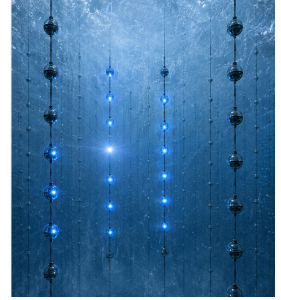
## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$



## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules



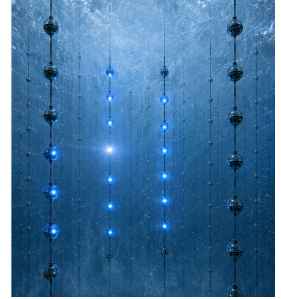
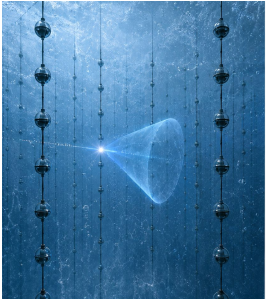
## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

We know direction, energy etc. influence the data  
-> correlations exist!  
There is a non-trivial joint distribution  $p(\theta, x)$  !



## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “ $x$ ”

$x = \{x_1, x_2, \dots, x_N\}$  for  $N$  modules

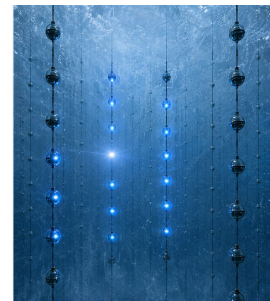
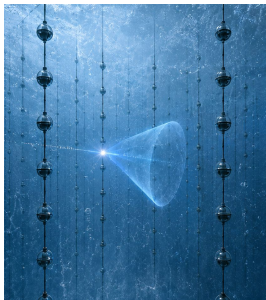
We know direction, energy etc. influence the data  
→ correlations exist!

There is a non-trivial joint distribution  $p(\theta, x)$  !

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$



Law of conditional probabilities



## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

We know direction, energy etc. influence the data

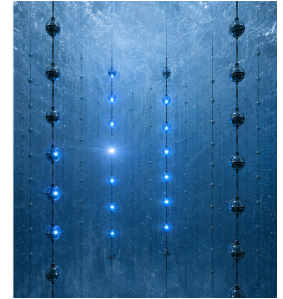
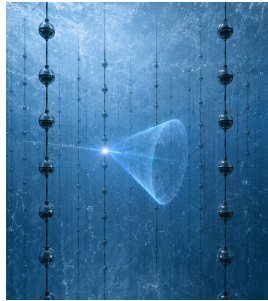
-> correlations exist!

There is a non-trivial joint distribution  $p(\theta, x)$  !

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$

Bayes' Theorem  $\longrightarrow$

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



## Event parameters $\theta$ :

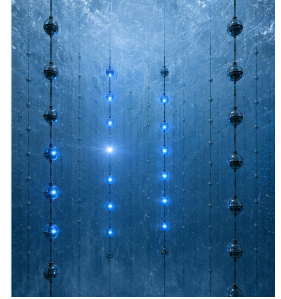
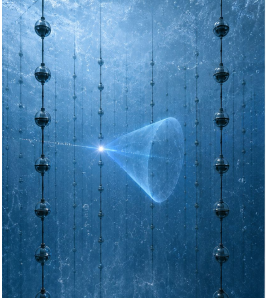
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



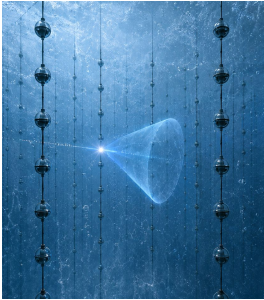
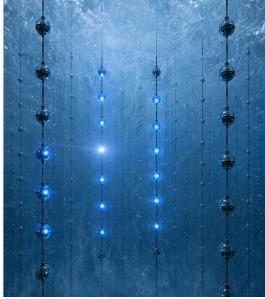
## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “ $x$ ”

$x = \{x_1, x_2, \dots, x_N\}$  for  $N$  modules

## Bayes' Theorem


$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$


Posterior  
(anti-causal / inverse)

## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

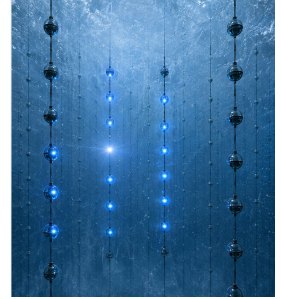
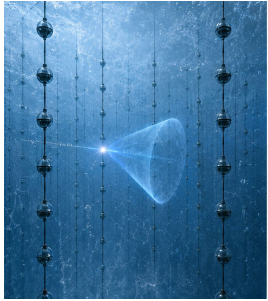
$x = \{x_1, x_2, \dots, x_N\}$  for N modules

Bayes' Theorem

Prior (apriori parameter assumptions)

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Posterior  
(anti-causal / inverse)



## Event parameters $\theta$ :

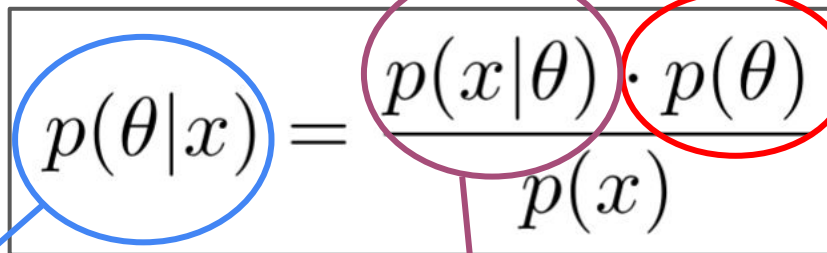
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

Bayes' Theorem

Prior (apriori parameter assumptions)



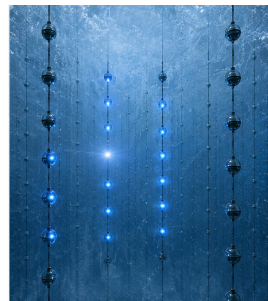
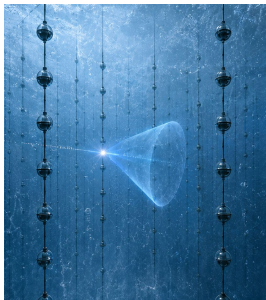
The diagram shows the equation  $p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$  enclosed in a black rectangular box. A blue circle highlights the term  $p(\theta|x)$ , with a blue line extending from it to a box labeled 'Posterior (anti-causal / inverse)'. A purple circle highlights the term  $p(x|\theta)$ , with a purple line extending from it to a box containing 'Response function / Data generation function / Forward model (Causal direction)' and 'Likelihood  $p(x|\theta) \equiv \mathcal{L}(\theta)$ '. A red circle highlights the term  $p(\theta)$ , with a red line extending from it to a box labeled 'Prior (apriori parameter assumptions)'.

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Posterior  
(anti-causal / inverse)

Response function /  
Data generation function /  
Forward model  
(Causal direction)

**Likelihood**  $p(x|\theta) \equiv \mathcal{L}(\theta)$



## Event parameters $\theta$ :

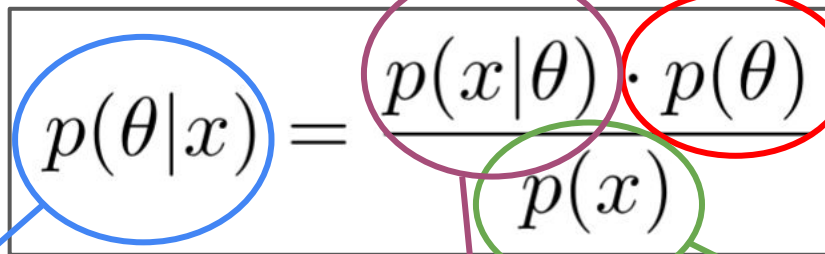
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “ $x$ ”

$x = \{x_1, x_2, \dots, x_N\}$  for  $N$  modules

## Bayes' Theorem

Prior (apriori parameter assumptions)



The diagram shows the equation  $p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$  enclosed in a black box. Annotations include: a blue circle around  $p(\theta|x)$  with a line pointing to a blue box labeled 'Posterior (anti-causal / inverse)'; a red circle around  $p(\theta)$  with a line pointing to a red box labeled 'Prior (apriori parameter assumptions)'; a green circle around  $p(x)$  with a line pointing to a green box labeled 'Marginal likelihood / Evidence (constant w.r.t.  $\theta$ )'; and a purple circle around  $p(x|\theta)$  with a line pointing to a purple box labeled 'Response function / Data generation function / Forward model (Causal direction)'.

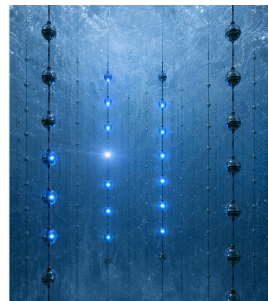
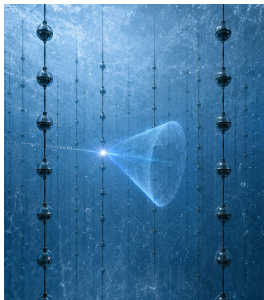
$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Posterior  
(anti-causal / inverse)

Response function /  
Data generation function /  
Forward model  
(Causal direction)

**Likelihood**  $p(x|\theta) \equiv \mathcal{L}(\theta)$

Marginal likelihood / Evidence  
(constant w.r.t.  $\theta$ )

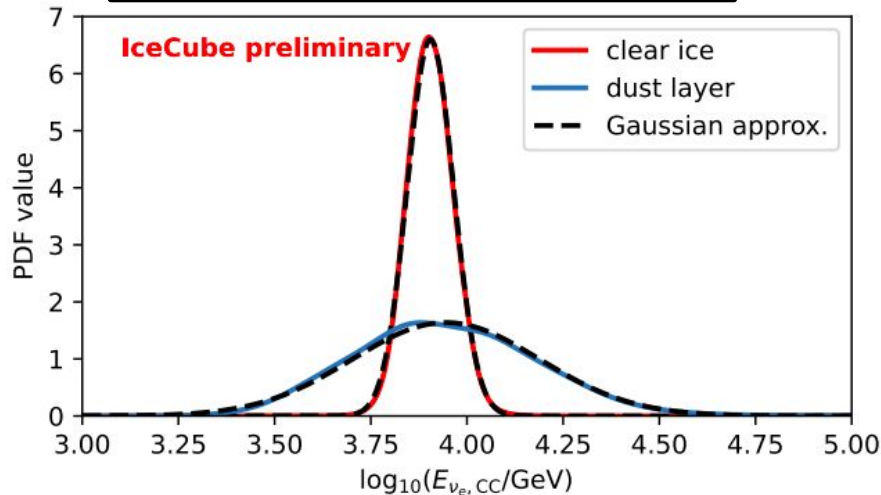


All reconstruction methods:

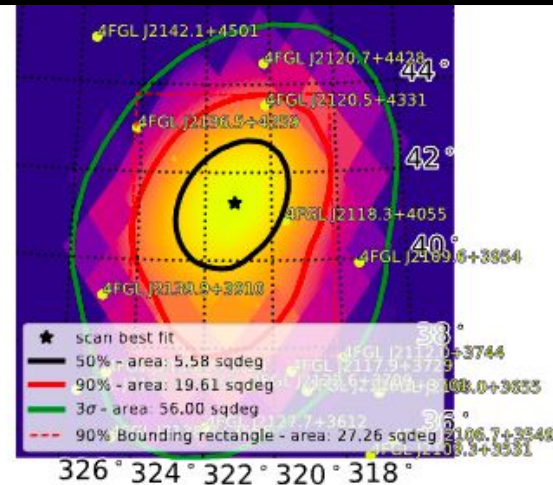
Utilize different parts of “Bayes’ Theorem” to construct confidence / credible intervals

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Posterior over “log-energy”



Confidence Region over direction (likelihood-scan)



All reconstruction methods:

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Utilize different parts of “Bayes’ Theorem” to construct confidence / credible intervals

### Interactive Question 1:

Which part of Bayes’ Theorem do we need for both Frequentist and Bayesian parameter estimation?

Why?

Hint: There are 4 possible parts!

$p(x)$

$p(x|\theta)$

$p(\theta|x)$

$p(\theta)$

# Learning objectives lecture 1:

- Bayes' Theorem and the “joint” distribution
- **Frequentist / Bayesian parameter estimation**
- Information theory: KL divergence and its interpretation
- likelihood-based vs likelihood-free methods

Event parameters  $\theta$ :

Bayes' Theorem

Data "x"

$\mathbf{x}=\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$   
for N modules

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

In Practice:

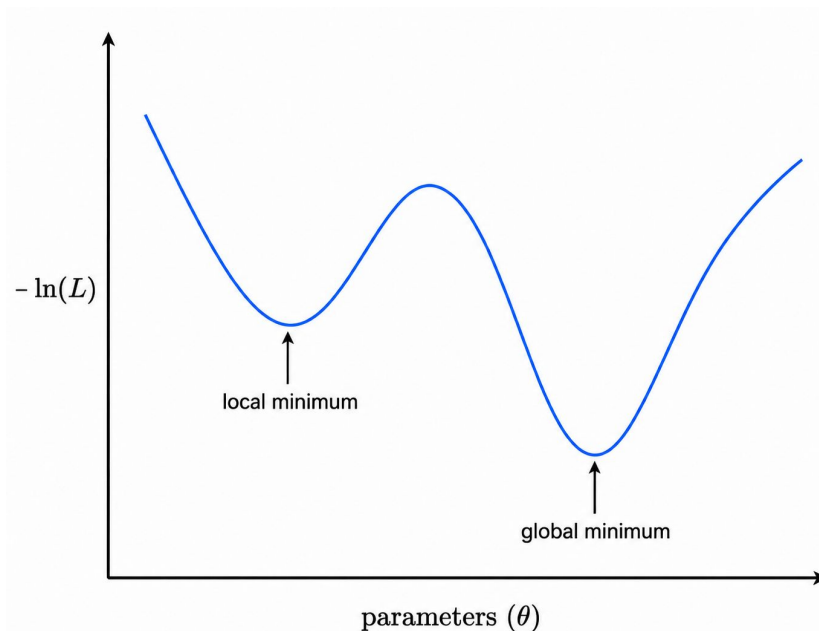
Maximum of  $\log(L)$

-> **Minimum of  $-\log(L)$**

Maximum Likelihood  
(Fisher, Pearson, Wilks', 20s / 30s)

Response function /  
Data generation function /  
Forward model  
(Causal direction)

**Likelihood**  $p(x|\theta) \equiv \mathcal{L} \equiv L$



Event parameters  $\theta$ :

Bayes' Theorem

Data "x"

$\mathbf{x}=\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$   
for N modules

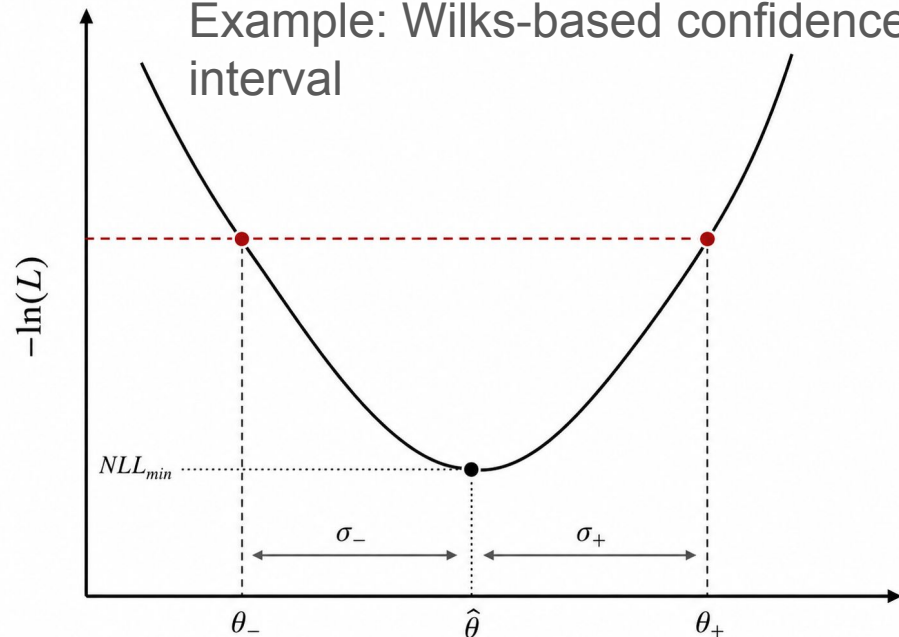
$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Maximum Likelihood  
(Fisher, Pearson, Wilks', 20s / 30s)

Response function /  
Data generation function /  
Forward model  
(Causal direction)

**Likelihood**  $p(x|\theta) \equiv \mathcal{L}(\theta)$

Example: Wilks-based confidence interval



Event parameters  $\theta$ :

Bayes' Theorem

Data "x"

$\mathbf{x}=\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$   
for N modules

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

Maximum Likelihood  
(Fisher, Pearson, Wilks', 20s / 30s)

Response function /  
Data generation function /  
Forward model  
(Causal direction)

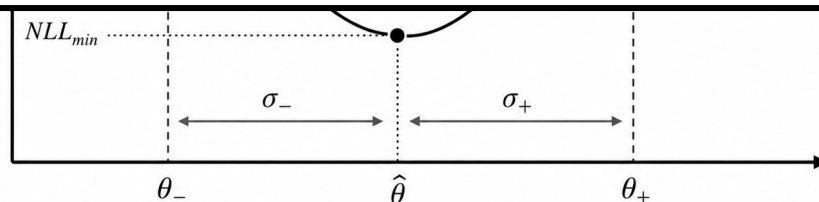
**Likelihood**  $p(x|\theta) \equiv \mathcal{L}(\theta)$

Example: Wilks-based confidence interval

95% confidence interval != 95% probability to contain true value

If a value is true, it ends up inside in 95% of repeated realizations -> good for exclusion

$-\ln(L)$



## Event parameters $\theta$ :

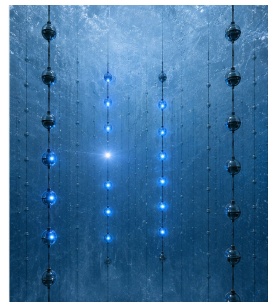
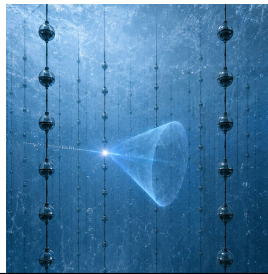
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

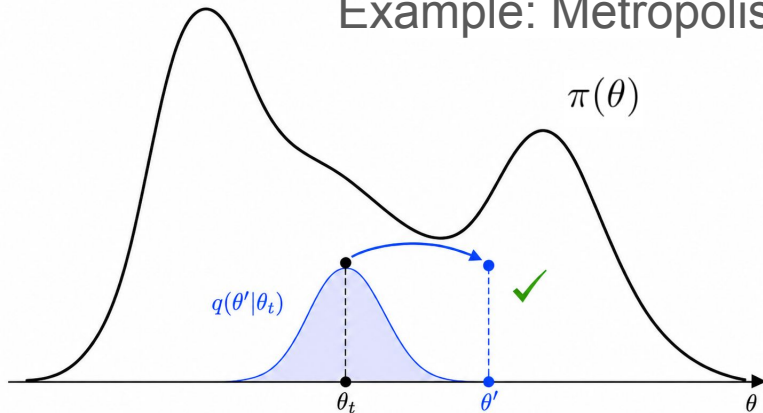


Markov-Chain Monte Carlo  
(Bayesian, '50s-today)

Likelihood and Prior

$$p(x|\theta) \cdot p(\theta) \equiv \pi(\theta)$$

## Example: Metropolis



$$\alpha = \min\left(1, \frac{\pi(\theta')}{\pi(\theta_t)}\right)$$

## Event parameters $\theta$ :

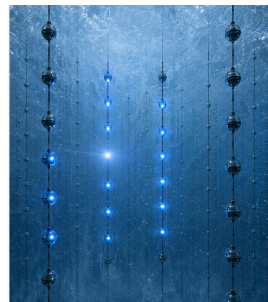
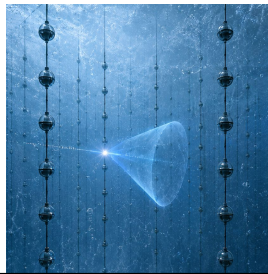
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

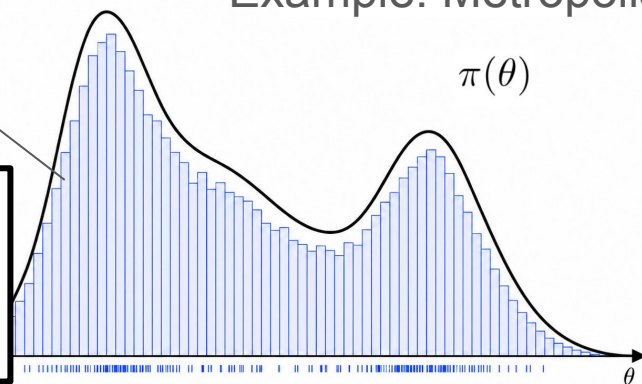


Markov-Chain Monte Carlo  
(Bayesian, '50s-today)

Likelihood and Prior

$$p(x|\theta) \cdot p(\theta) \equiv \pi(\theta)$$

Samples  
From  
Posterior



Example: Metropolis

$$\alpha = \min\left(1, \frac{\pi(\theta')}{\pi(\theta_t)}\right)$$

## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{\int p(x|\theta) \cdot p(\theta) d\theta}$$

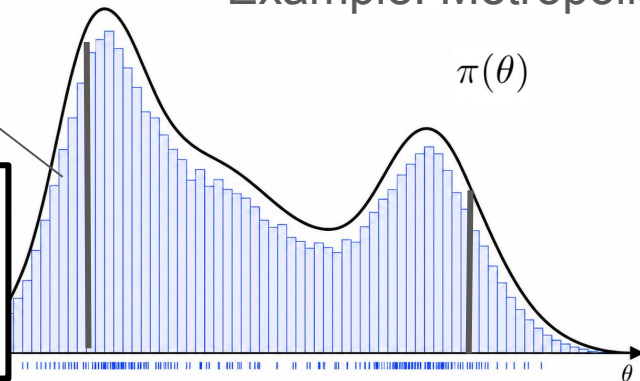
95% Credible interval = 95% probability to contain true value under Prior

Markov-Chain Monte Carlo  
(Bayesian, '50s-today)

Likelihood and Prior

$$p(x|\theta) \cdot p(\theta) \equiv \pi(\theta)$$

Samples  
From  
Posterior



$$\alpha = \min\left(1, \frac{\pi(\theta')}{\pi(\theta_t)}\right)$$

## Event parameters $\theta$ :

$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{Z}$$

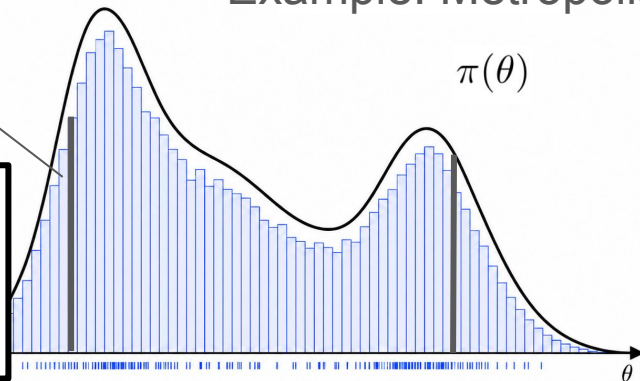
95% Credible interval = 95% probability to contain true value under Prior

Markov-Chain Monte Carlo  
(Bayesian, '50s-today)

Likelihood and Prior

$$p(x|\theta) \cdot p(\theta) \equiv \pi(\theta)$$

Samples  
From  
Posterior



$$\alpha = \min\left(1, \frac{\pi(\theta')}{\pi(\theta_t)}\right)$$

## Event parameters $\theta$ :

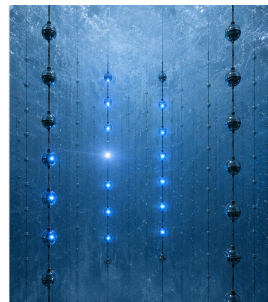
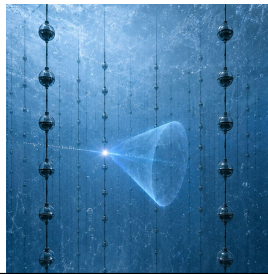
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

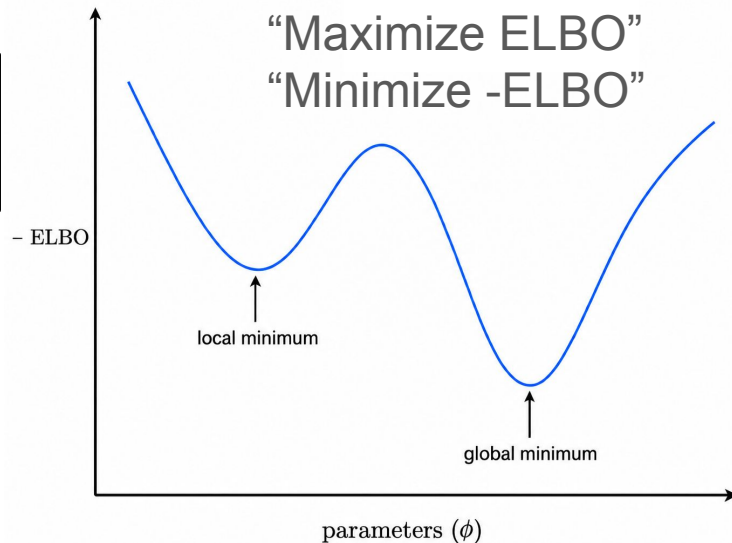
## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



Variational Inference -> Maximize ELBO  
(Bayesian, '90s-today)

$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$



## Event parameters $\theta$ :

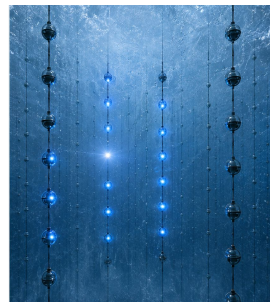
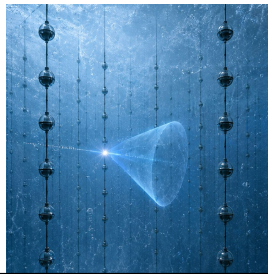
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

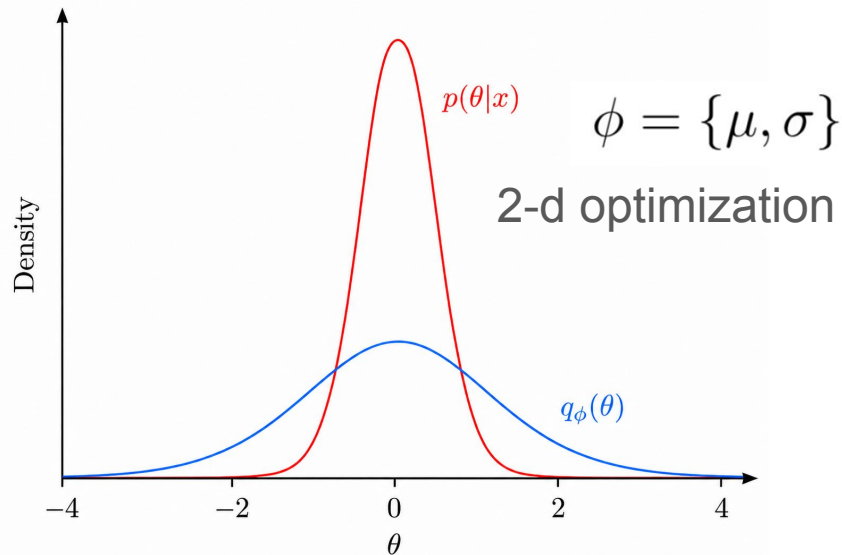
## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



Variational Inference -> Maximize ELBO  
(Bayesian, '90s-today)

$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$



## Event parameters $\theta$ :

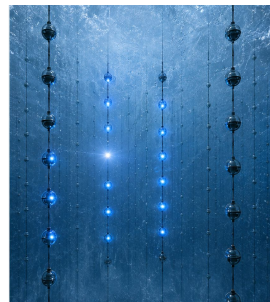
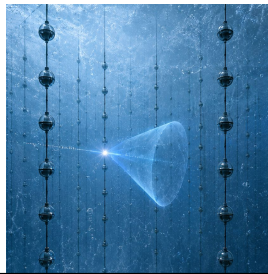
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

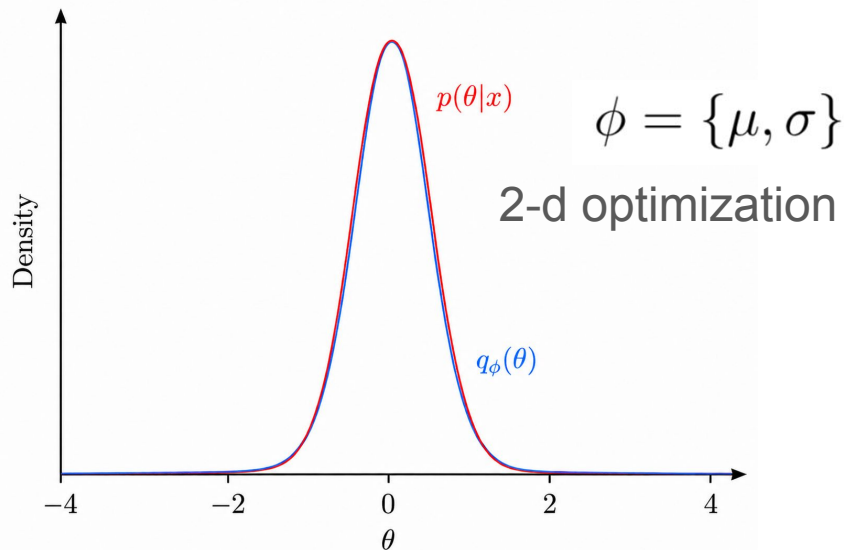
## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



Variational Inference -> Maximize ELBO  
(Bayesian, '90s-today)

$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$



## Event parameters $\theta$ :

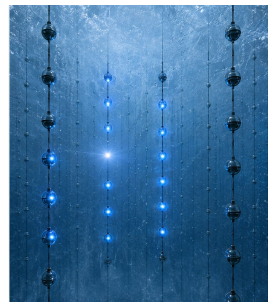
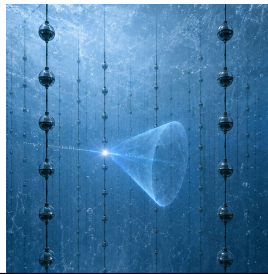
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

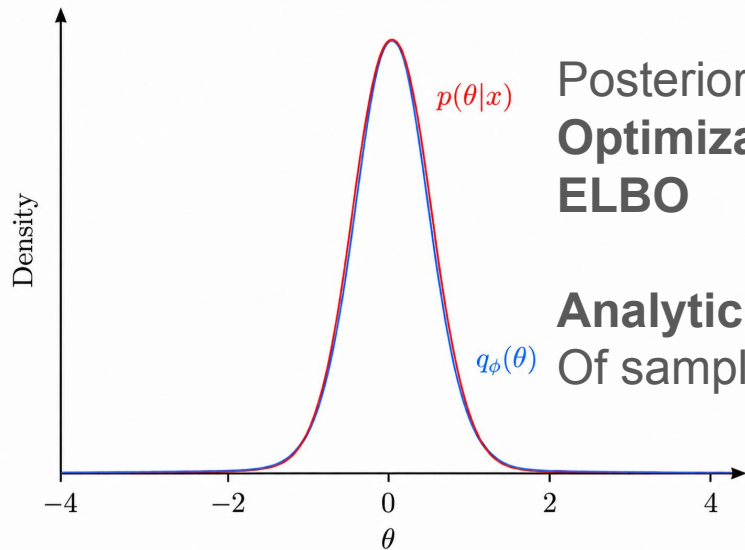
## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



Variational Inference -> Maximize ELBO  
(Bayesian, '90s-today)

$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$



Posterior via  
**Optimization of  
ELBO**

**Analytic** instead  
Of samples

## Event parameters $\theta$ :

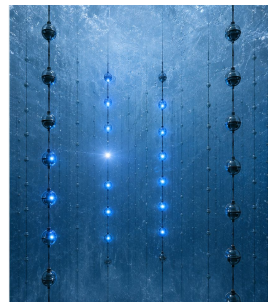
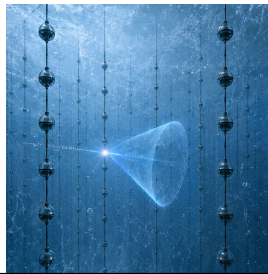
$\theta = \{\text{Direction, Energy, Position, Type, anything we record in MC, ...}\}$

## Data “x”

$x = \{x_1, x_2, \dots, x_N\}$  for N modules

## Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$



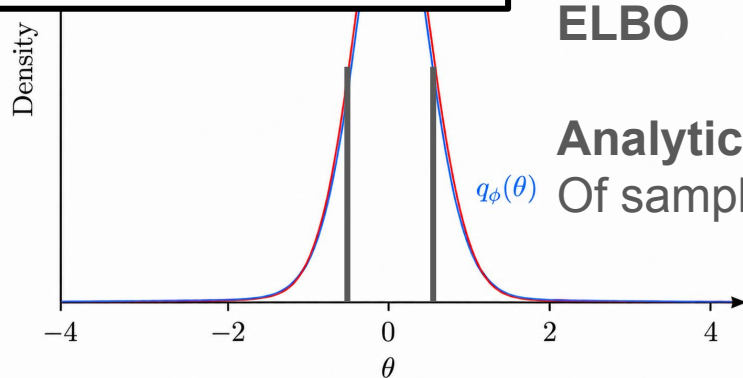
Variational Inference -> M (Bayesian, '90s-today)

95% Credible interval = 95% probability to contain true value under Prior

Posterior via **Optimization of ELBO**

**Analytic** instead Of samples

$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$



Observation: In both Frequentist (Max. likelihood) and traditional Bayesian methods (MCMC / variational inference) we require the likelihood function!

## **Interactive Question 2:**

**What are the pros / cons of relying on the likelihood function?**

**Try to find at least 1 pro / con**

Observation: In both Frequentist (Max. likelihood) and traditional Bayesian methods (MCMC / variational inference) we require the likelihood function!

**Question: What are the pros / cons of relying on the likelihood function?**

**Discuss! (5 min)**

Pros:

- Well defined and mature statistical tools use it
- Frequentist guarantees (asymptotically optimal, saturating Cramer-Rao bound, etc.)
- Established methods use it: MCMC / variational inference to get posterior in traditional Bayesian analysis

Observation: In both Frequentist (Max. likelihood) and traditional Bayesian methods (MCMC / variational inference) we require the likelihood function!

**Question: What are the pros / cons of relying on the likelihood function?**

**Discuss! (5 min)**

Pros:

- Well defined and mature statistical tools use it
- Frequentist guarantees (asymptotically optimal, saturating Cramer-Rao bound, etc.)
- Established methods use it: MCMC / variational inference to get posterior in traditional Bayesian analysis

Cons:

- We need to model it first! Requires likelihood fit for the likelihood function itself!  
-> Numerically challenging! (**next slides, lecture 2**)
- High-dimensional physics hypothesis have many local optima!  
-> Complexity, part 2!  
(**lecture 2**)

- 1) How to cleanly define “likelihood fit of likelihood?”**
- 2) How could we possibly do likelihood-free modelling?**

**-> Information-theoretic viewpoint**



# Learning objectives lecture 1:

- Bayes' Theorem and the “joint” distribution
- Frequentist / Bayesian parameter estimation
- **Information theory: KL divergence and its interpretation**
- likelihood-based vs likelihood-free methods

# Information theory

Established in the 40s by Claude Shannon  
(probabilities  $\leftrightarrow$  coding theory)

(40s - Shannon):

Shannon-Entropy: (discrete)

$$H(p) = - \sum_{x \in \mathcal{X}} p(x) \log p(x)$$

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \sum_{x \in \mathcal{X}} p(x) \log \frac{p(x)}{q(x)}$$



# Information theory

Established in the 40s by Claude Shannon  
(probabilities  $\leftrightarrow$  coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

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

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$



# Information theory

Established in the 40s by Claude Shannon  
(probabilities <-> coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

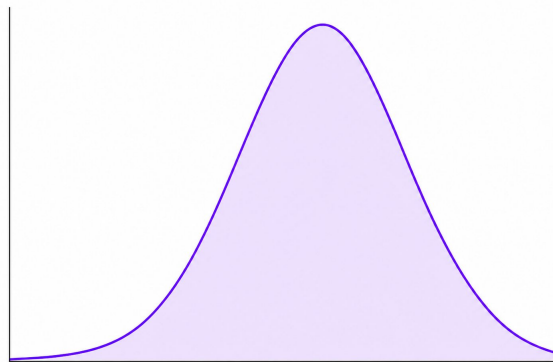
$$H(p) = - \int p(x) \log p(x) dx$$

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

Higher entropy



# Information theory

Established in the 40s by Claude Shannon  
(probabilities <-> coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

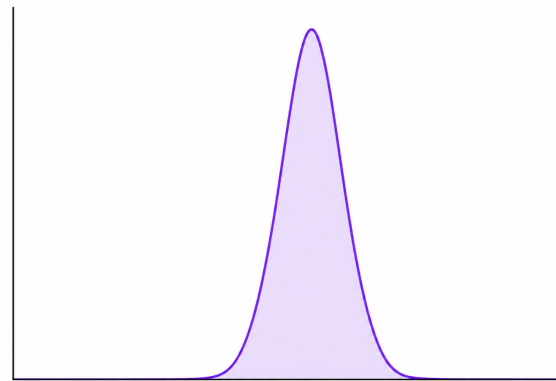
$$H(p) = - \int p(x) \log p(x) dx$$

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

Lower entropy



# Information theory

Established in the 40s by Claude Shannon  
(probabilities <-> coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

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

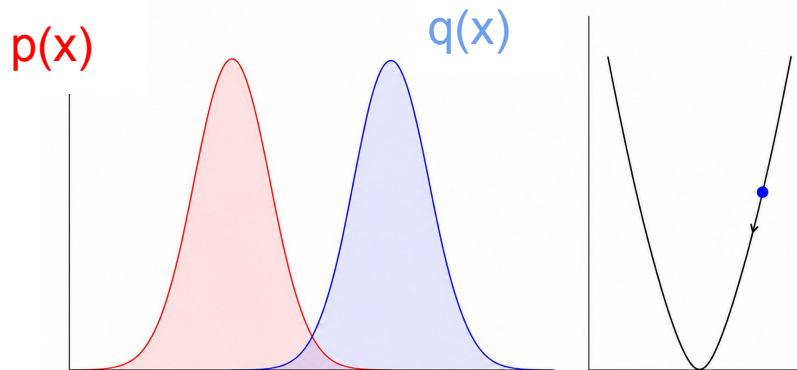
(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

Notation:  $q(x)$  is parametrized  
Approximates  $p(x) / P(x)$

Higher KL



# Information theory

Established in the 40s by Claude Shannon  
(probabilities <-> coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

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

(50s - Kullback / Leibler):

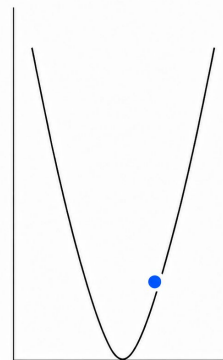
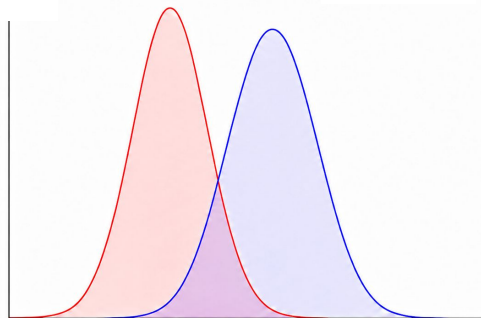
Relative Entropy / KL-Divergence:

$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

Lower KL (0 = distributions the same)

$p(x)$

$q(x)$



# Information theory

Established in the 40s by Claude Shann  
(probabilities <-> coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

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

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

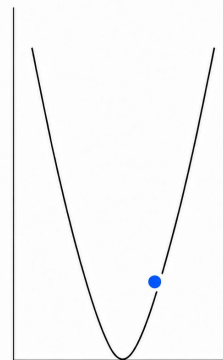
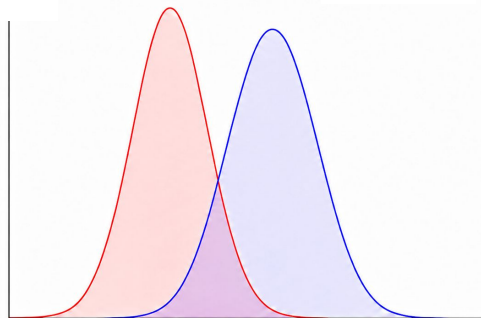
$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

$$- \int p(x) \log q(x) dx - \left( - \int p(x) \log p(x) dx \right)$$

Lower KL (0 = distributions the same)

$p(x)$

$q(x)$



# Information theory

Established in the 40s by Claude Shannon  
(probabilities  $\leftrightarrow$  coding theory)

(40s - Shannon):

Differential Entropy: (continuous)

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

(50s - Kullback / Leibler):

Relative Entropy / KL-Divergence:

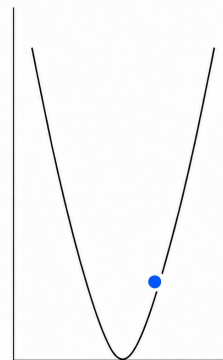
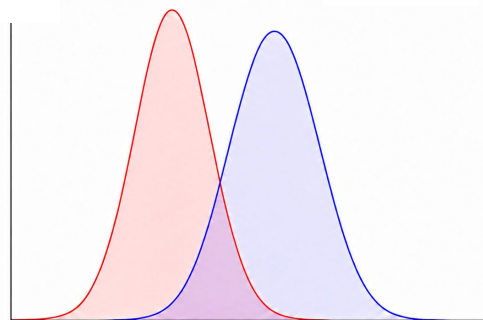
$$D_{\text{KL}}(p \parallel q) = \int p(x) \log \frac{p(x)}{q(x)} dx$$

$$\underbrace{- \int p(x) \log q(x) dx}_{\text{cross-entropy}} - \underbrace{\left( - \int p(x) \log p(x) dx \right)}_{\text{entropy} \rightarrow \text{const w.r.t } q(x)}$$

Lower KL (0 = distributions the same)

$p(x)$

$q(x)$

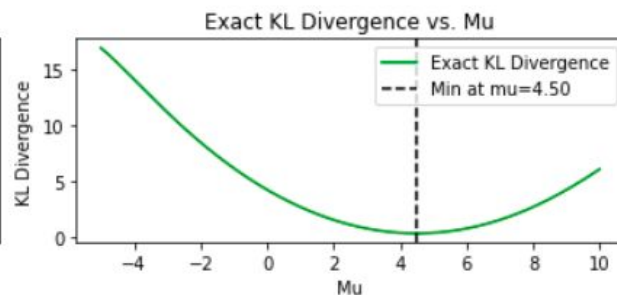
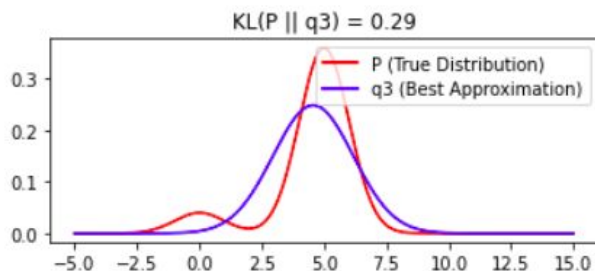
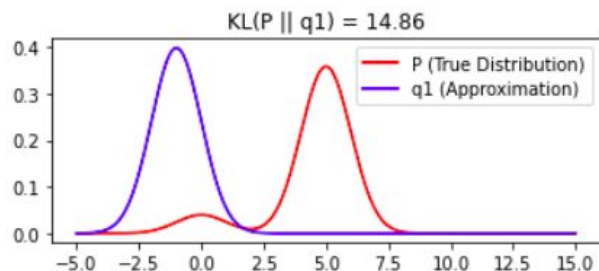


# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x)$

Approximation:  $q(x)$

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$



# PDF-approximation matching via “forward” KL divergence

In practice .. **We do not know  $P(x)$  -> we have samples from  $P(x)$**

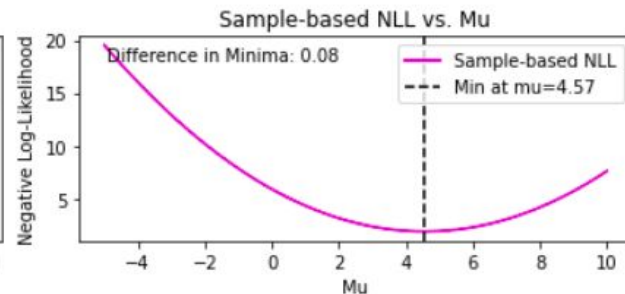
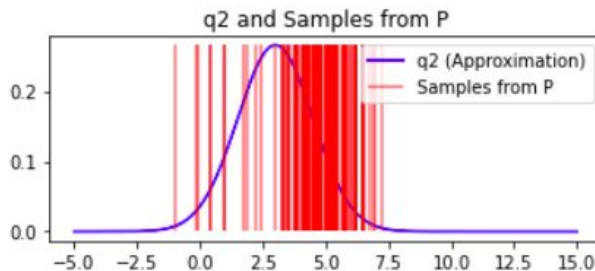
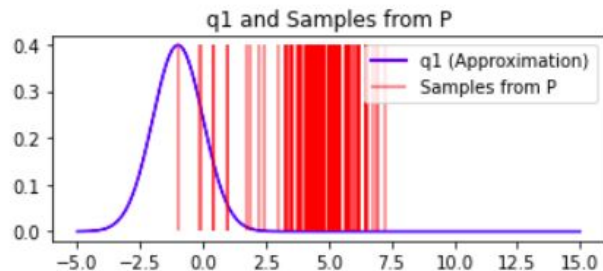
True distribution:  $P(x)$

Approximation:  $q(x)$

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i)}{q_{\theta}(x_i)} \right)$$

$$\longrightarrow \underset{\theta}{\operatorname{argmin}} \left( \frac{1}{N} \sum_i \ln \left( \frac{P(x_i)}{q_{\theta}(x_i)} \right) \right) \quad \underset{\theta}{\operatorname{argmin}} \left( \frac{1}{N} \sum_i -\ln (q_{\theta}(x_i)) \right)$$



### Interactive Question 3:

Why do we not use the so-called “reverse” KL divergence?

Hint: See the sample-based approximation?

Forward KL (previous slides):

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$

Reverse KL:

$$D_{\text{KL}}(q|P)[\theta] = \int q_{\theta}(x) \ln \left( \frac{q_{\theta}(x)}{P(x)} \right) dx$$

$$D_{\text{KL}}(q|P)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{q_{\theta}(x_i[\theta])}{P(x_i[\theta])} \right)$$

**P(x): True distribution (unknown)**    **q(x):**  
**Approximation**

Sample-based  
approximation  $\longrightarrow$

# Question:

**Why do we typically not use the “reverse” KL-divergence here?**

Forward KL (previous slides):

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$

1) Here we do not have access to  $P(x)$   
Evaluate as sample-based expectation value  
(certain assumptions on  $q$ ):

$$D_{\text{KL}}(q|P)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{q_{\theta}(x_i[\theta])}{P(x_i[\theta])} \right)$$

Reverse KL:

$$D_{\text{KL}}(q|P)[\theta] = \int q_{\theta}(x) \ln \left( \frac{q_{\theta}(x)}{P(x)} \right) dx$$



Without  $P$  we cannot  
calculate the gradient

# Question:

**Why do we typically not use the “reverse” KL-divergence here?**

1) Here we do not have access to  $P(x)$

2) The qualitative approximation properties are different, and typically worse!

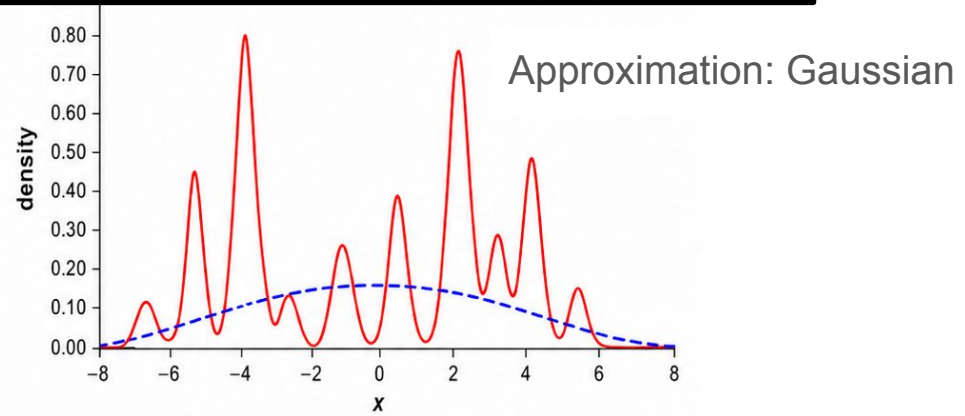
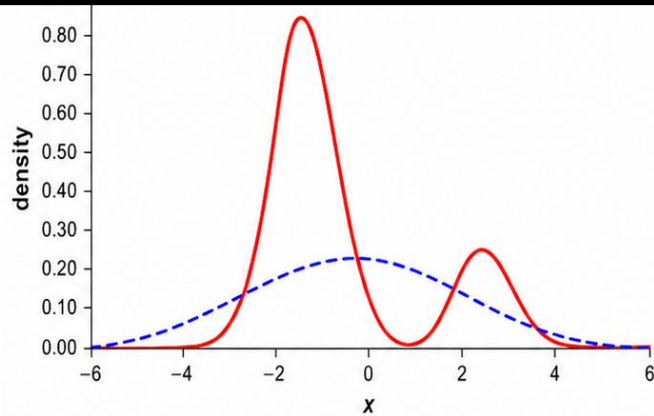
Forward KL (previous slides):

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$

Reverse KL:

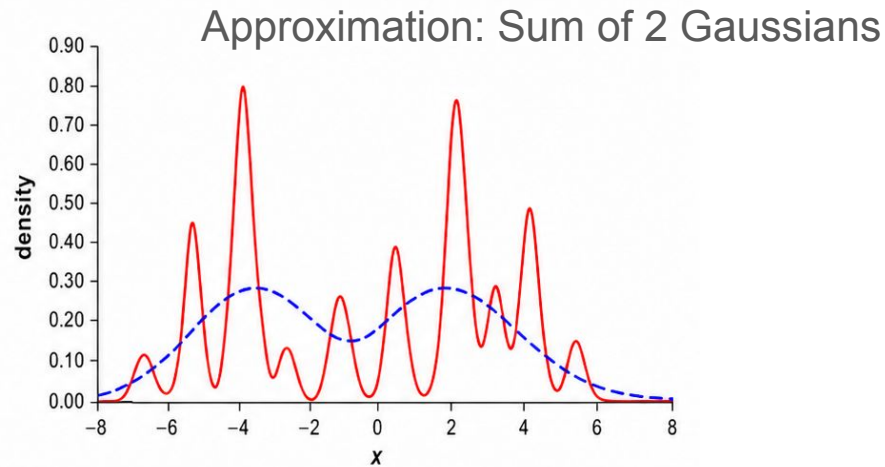
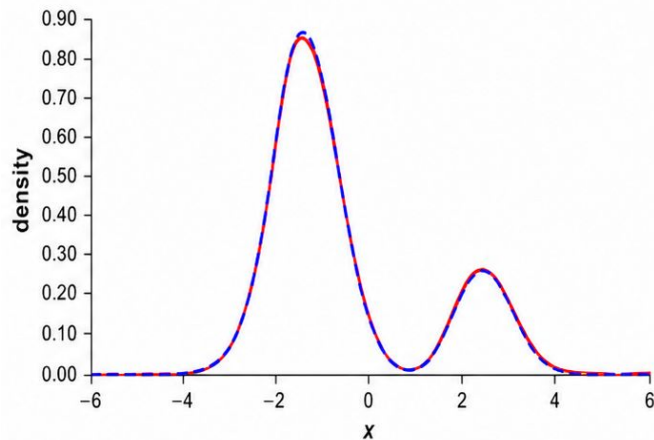
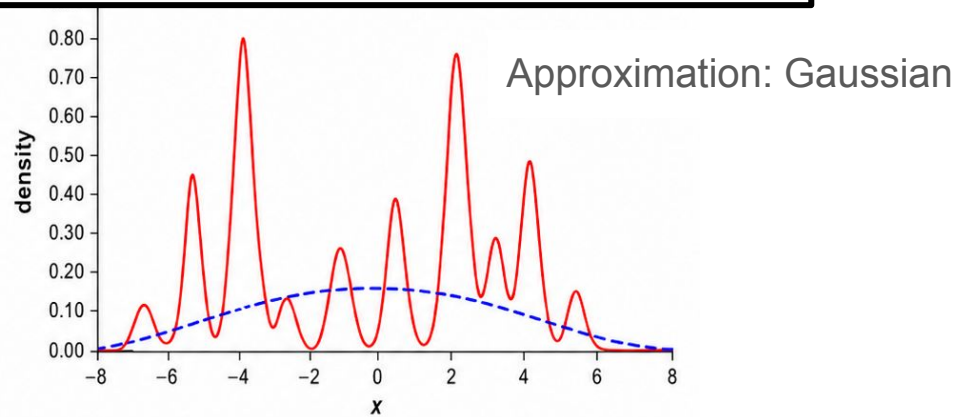
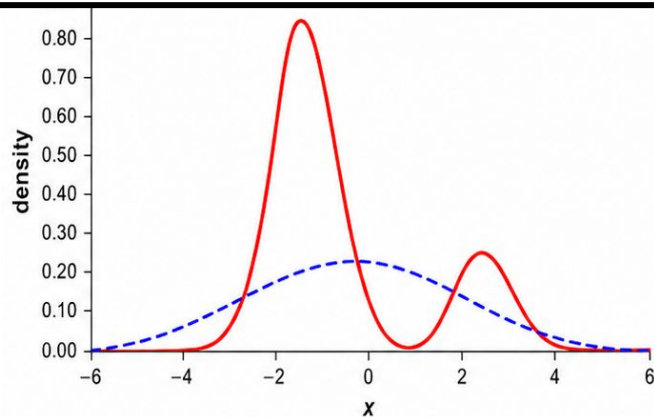
$$D_{\text{KL}}(q|P)[\theta] = \int q_{\theta}(x) \ln \left( \frac{q_{\theta}(x)}{P(x)} \right) dx$$

Forward KL:  $\sim$  approximation (blue) seeks out probability mass  
Approximation = Gaussian: 1st (mean) and 2nd (variance) are matched

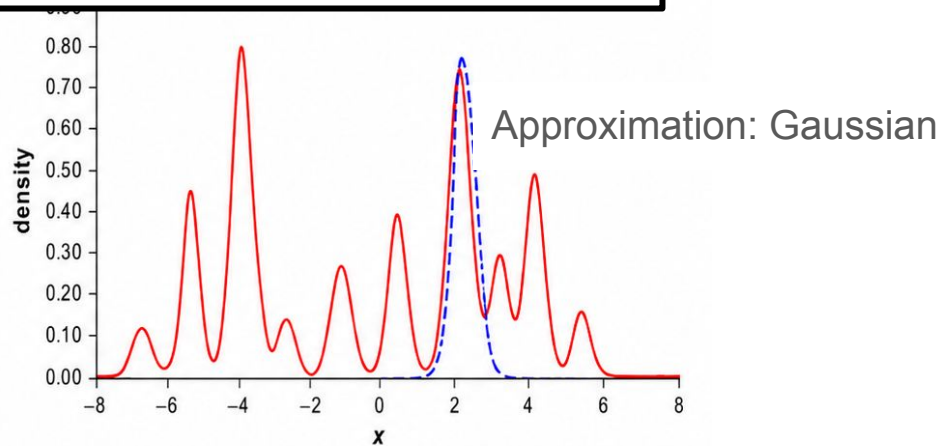
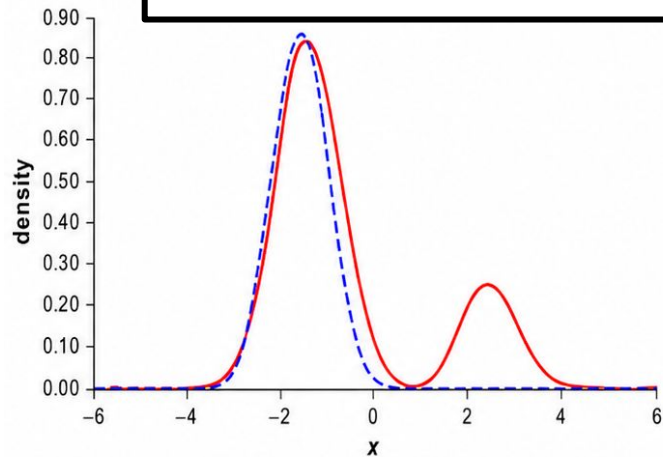


5

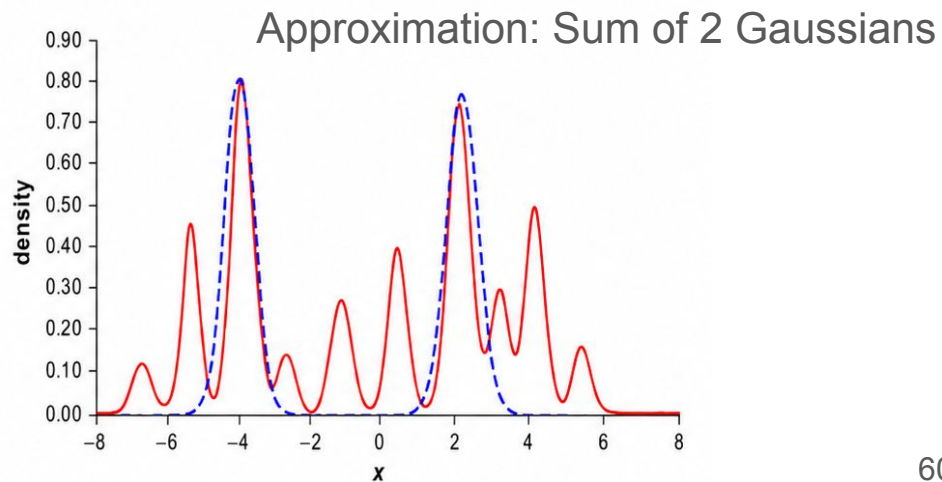
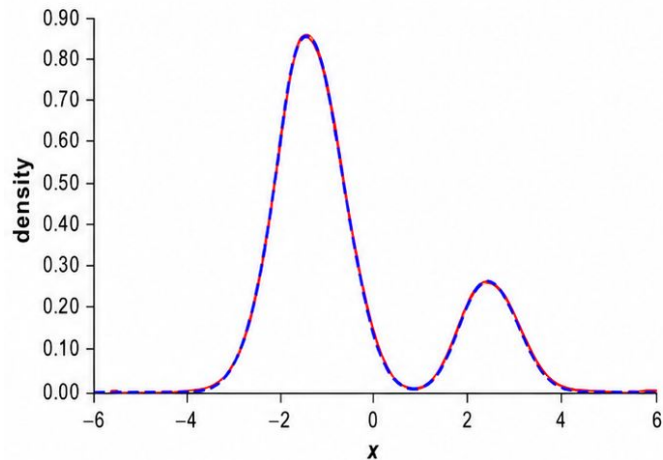
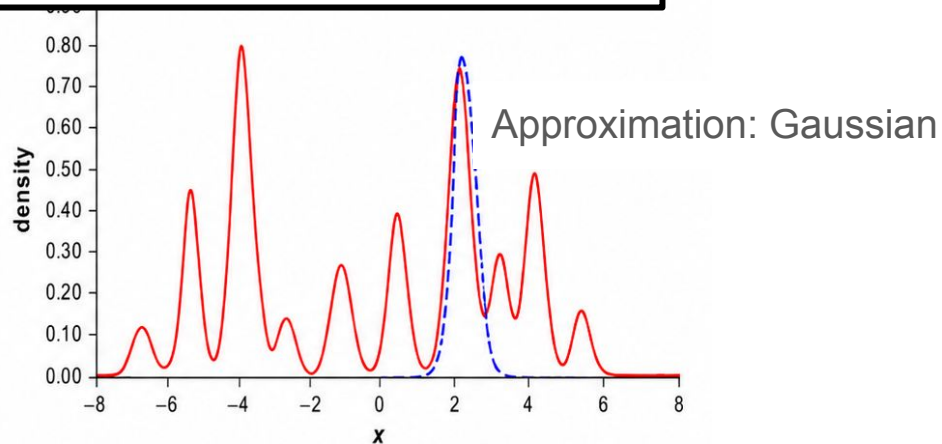
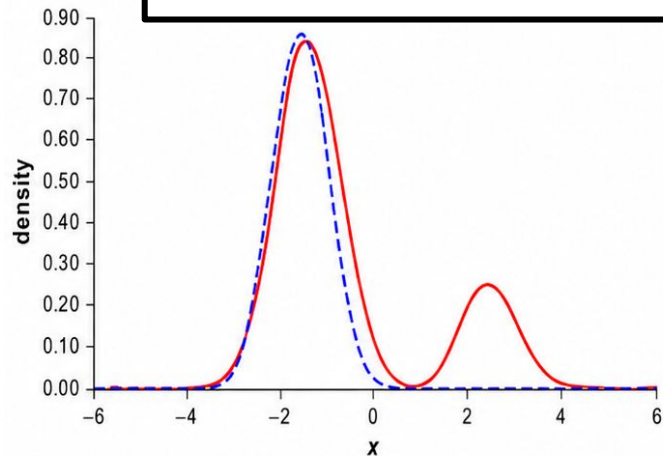
Forward KL:  $\sim$  approximation (blue) seeks out probability mass  
Approximation = Gaussian: 1st (mean) and 2nd (variance) are matched



## Reverse KL: $\sim$ approximation (blue) seeks out modes



## Reverse KL: $\sim$ approximation (blue) seeks out modes



# PDF-approximation matching via “forward” KL divergence

In practice .. **We do not know  $P(x)$  -> we have samples from  $P(x)$**

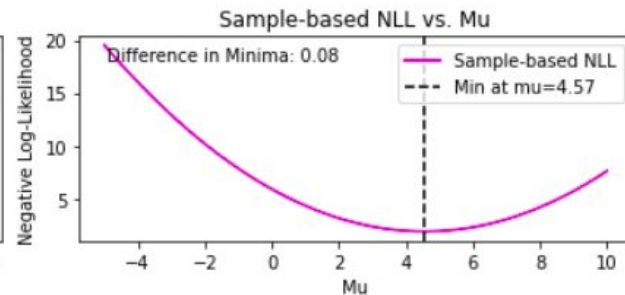
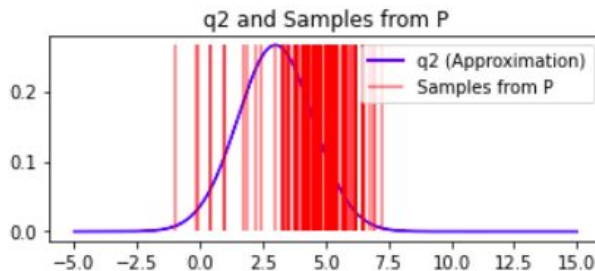
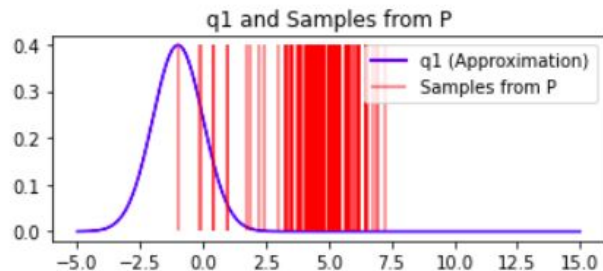
True distribution:  $P(x)$

Approximation:  $q(x)$

$$D_{\text{KL}}(P|q)[\theta] = \int P(x) \ln \left( \frac{P(x)}{q_{\theta}(x)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i)}{q_{\theta}(x_i)} \right)$$

$$\longrightarrow \underset{\theta}{\operatorname{argmin}} \left( \frac{1}{N} \sum_i \ln \left( \frac{P(x_i)}{q_{\theta}(x_i)} \right) \right) \quad \underset{\theta}{\operatorname{argmin}} \left( \frac{1}{N} \sum_i -\ln (q_{\theta}(x_i)) \right)$$



# PDF-approximation matching via “forward” KL divergence

In practice .. **We do not know  $P(x)$  -> we have samples from  $P(x)$**

True distribution:  $P(x)$

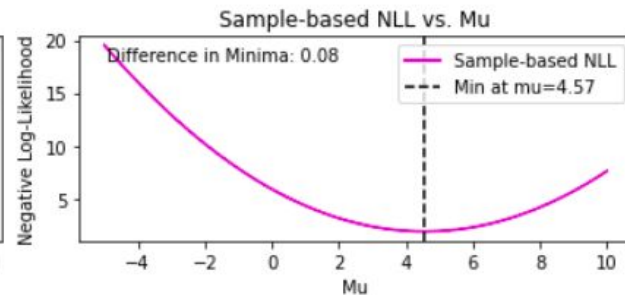
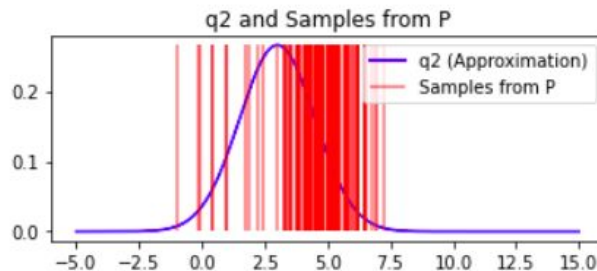
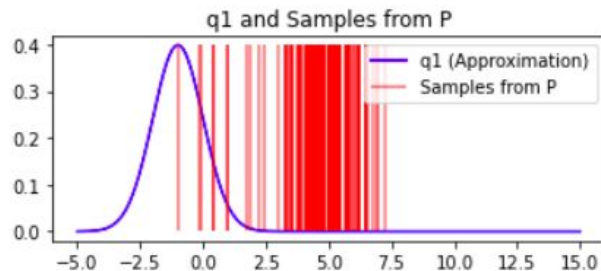
Approximation:  $q(x)$

$$D_{\text{KL}}(P|q)[\theta] = \int P(x|\theta_{\text{true}}) \ln \left( \frac{P(x|\theta_{\text{true}})}{q(x|\theta)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right)$$

(Slight rewrite! Same thing..)

$$\longrightarrow \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right) = \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i -\ln(q(x_i|\theta))$$



# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x)$

Approximation:  $q(x)$

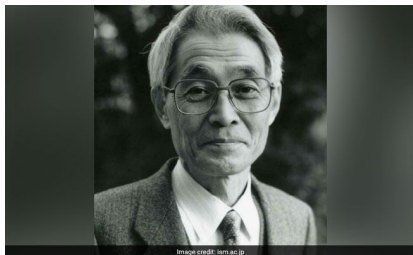
$$D_{\text{KL}}(P|q)[\theta] = \int P(x|\theta_{\text{true}}) \ln \left( \frac{P(x|\theta_{\text{true}})}{q(x|\theta)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right)$$

$$\longrightarrow \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right) = \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i -\ln(q(x_i|\theta))$$

**Max. Likelihood for  $L(\theta)=q(x|\theta)$**

1973 (Akaike): **Maximum likelihood derivable from KL divergence!**



$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x)$

Approximation:  $q(x)$

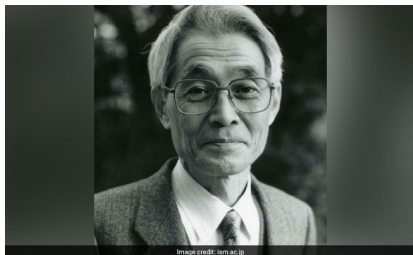
$$D_{\text{KL}}(P|q)[\theta] = \int P(x|\theta_{\text{true}}) \ln \left( \frac{P(x|\theta_{\text{true}})}{q(x|\theta)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right)$$

$$\longrightarrow \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right) = \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i -\ln(q(x_i|\theta))$$

**Max. Likelihood for  $L(\theta)=q(x|\theta)$**

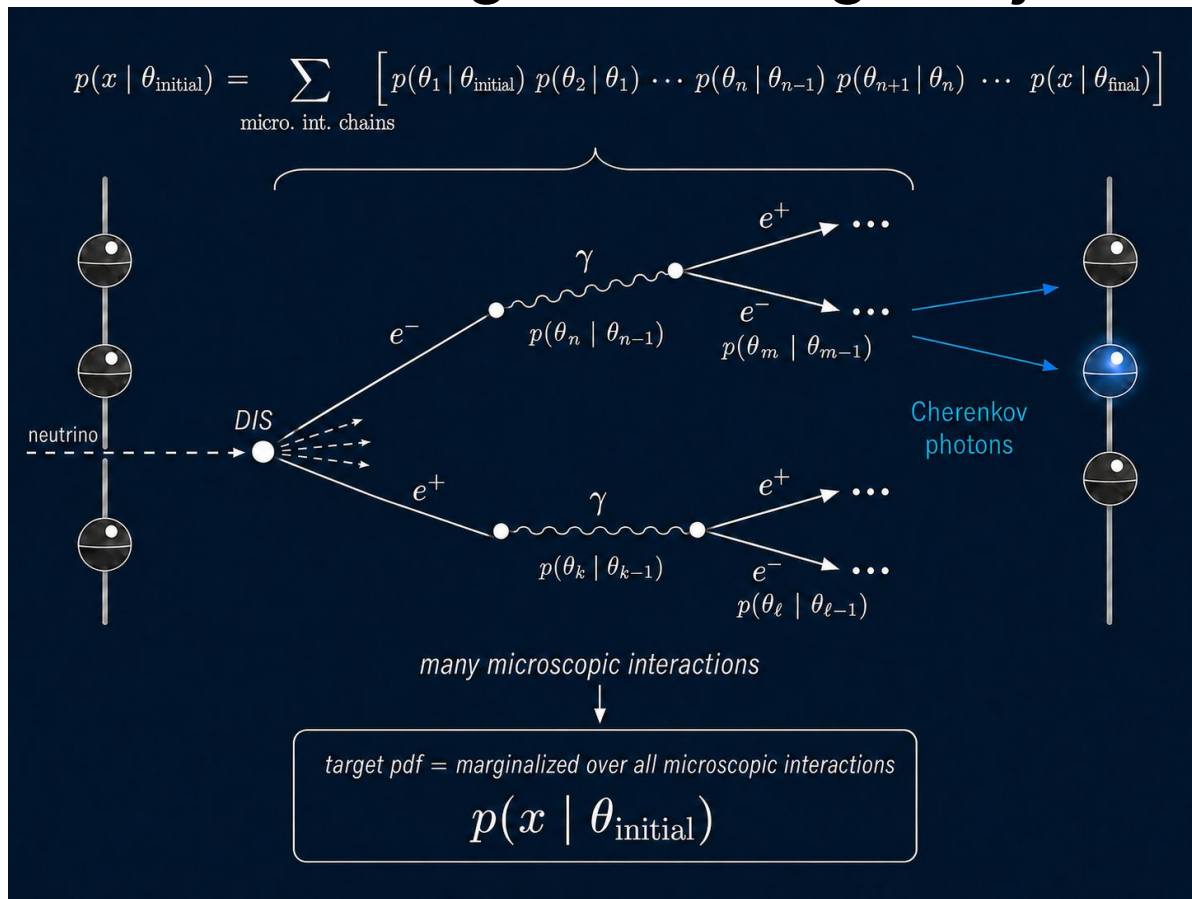
1973 (Akaike): **Maximum likelihood derivable from KL divergence!**



**But:** We still have to assume a parametrization of the Likelihood! **What is it?**

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

# The KL-divergence using the joint distribution in $(x, \theta)$



Prior

$p(\theta_{\text{initial}})$

(energy distribution,  
zenith distribution,...)

Monte Carlo  
implements  
sampling  
from the joint:

$p(\theta_{\text{initial}}) p(x \mid \theta_{\text{initial}})$

$= p(\theta_{\text{initial}}, x)$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right] + \mathbb{E}_{P(\theta)} \left[ \log \frac{P(\theta)}{q(\theta)} \right]$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right] + \mathbb{E}_{P(\theta)} \left[ \log \frac{P(\theta)}{q(\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(\theta)} \left[ D_{\text{KL}} \left( P(x|\theta) \parallel q_{\omega}(x|\theta) \right) \right] + \text{const}$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right] + \mathbb{E}_{P(\theta)} \left[ \log \frac{P(\theta)}{q(\theta)} \right]$$

**data-distribution KL-divergence**

$$= \arg \min_{\omega} \mathbb{E}_{P(\theta)} \left[ D_{\text{KL}} \left( P(x|\theta) \parallel q_{\omega}(x|\theta) \right) \right] + \text{const}$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \quad \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right] + \mathbb{E}_{P(\theta)} \left[ \log \frac{P(\theta)}{q(\theta)} \right] \quad \begin{array}{l} \text{data-distribution} \\ \text{KL-divergence} \end{array}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(\theta)} \left[ D_{\text{KL}} \left( P(x|\theta) \parallel q_{\omega}(x|\theta) \right) \right] + \text{const}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} [-\log q_{\omega}(x|\theta)] \quad (\text{drop terms constant in } \omega)$$

$$\omega^* = \arg \min_{\omega} D_{\text{KL}} \left( P(x, \theta) \parallel q(\theta) q_{\omega}(x|\theta) \right) \longleftarrow \text{“Joint” KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x, \theta)}{q(\theta) q_{\omega}(x|\theta)} \right]$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} \left[ \log \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right] + \mathbb{E}_{P(\theta)} \left[ \log \frac{P(\theta)}{q(\theta)} \right] \text{data-distribution KL-divergence}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(\theta)} \left[ D_{\text{KL}} \left( P(x|\theta) \parallel q_{\omega}(x|\theta) \right) \right] + \text{const}$$

$$= \arg \min_{\omega} \mathbb{E}_{P(x, \theta)} [-\log q_{\omega}(x|\theta)] \quad (\text{drop terms constant in } \omega)$$

$$\approx \arg \min_{\omega} \frac{1}{N} \sum_i^N -\log q_{\omega}(x_i|\theta_i) \equiv \arg \min_{\omega} \mathcal{L}_{\text{ML}}(\omega).$$

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x)$

Approximation:  $q(x)$

$$D_{\text{KL}}(P|q)[\theta] = \int P(x|\theta_{\text{true}}) \ln \left( \frac{P(x|\theta_{\text{true}})}{q(x|\theta)} \right) dx$$

$$D_{\text{KL}}(P|q)[\theta] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right)$$

$$\longrightarrow \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_{\text{true}})}{q(x_i|\theta)} \right) = \operatorname{argmin}(\theta) \quad \frac{1}{N} \sum_i -\ln(q(x_i|\theta))$$

**Max. Likelihood for  
 $L(\theta)=q(x|\theta)$**

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

$$D_{\text{KL}}(P|q)[\omega] = \int P(x|\theta) \cdot P(\theta) \ln \left( \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right) dx + \text{Const}$$

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) + \text{Const}$$

$$\longrightarrow \operatorname{argmin}(\omega) \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) = \operatorname{argmin}(\omega) \frac{1}{N} \sum_i -\ln(q_{\omega}(x_i|\theta_i))$$

**(higher order) Max. Likelihood**  
**“Likelihood fit” of the likelihood function**

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

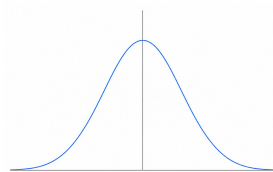
$$D_{\text{KL}}(P|q)[\omega] = \int P(x|\theta) \cdot P(\theta) \ln \left( \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right) dx + \text{Const}$$

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) + \text{Const}$$

$$\longrightarrow \operatorname{argmin}(\omega) \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) = \operatorname{argmin}(\omega) \frac{1}{N} \sum_i -\ln(q_{\omega}(x_i|\theta_i))$$

Event Params:

$\theta$



$$\begin{pmatrix} \mu \\ \sigma \end{pmatrix} = \vec{f}_{\omega}(\theta)$$

$$q_{\omega}(x | \theta) = \frac{1}{\sqrt{2\pi} \sigma_{\omega}[\theta]} \exp \left( -\frac{(x - \mu_{\omega}[\theta])^2}{2\sigma_{\omega}[\theta]^2} \right)$$

**(higher order) Max. Likelihood**

**“Likelihood fit” of the likelihood function**

**“Learn all possible data generating distributions”**

# PDF-approximation matching via “forward” KL divergence

Event Generator (2107.12110)

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

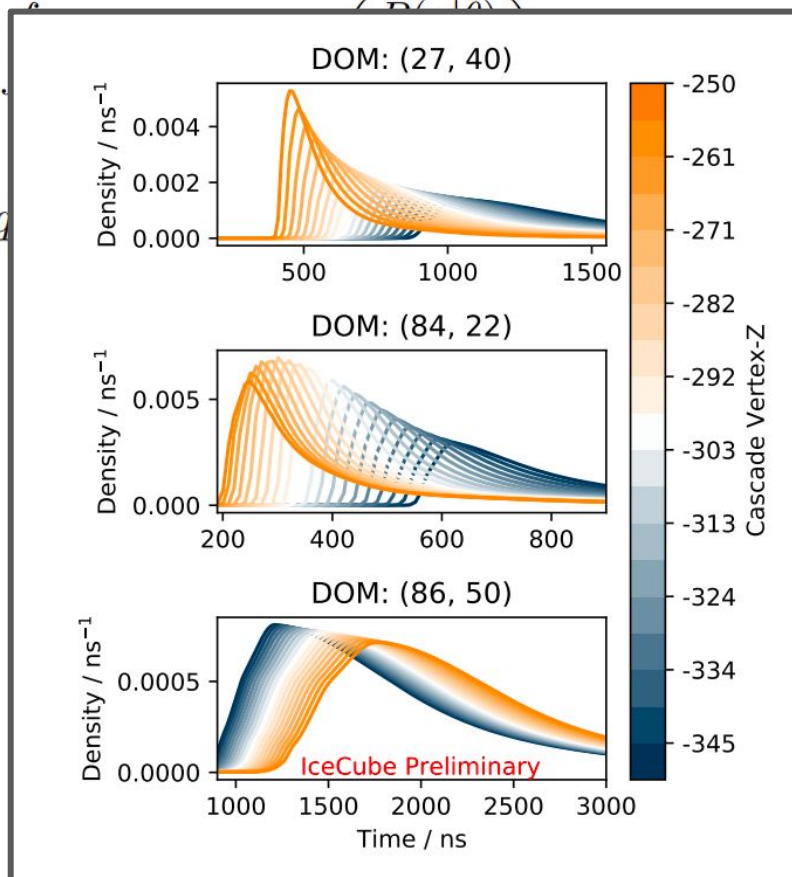
$$D_{\text{KL}}(P|q)[\omega] =$$

$$D_{\text{KL}}(P|q)$$

$$\longrightarrow \text{argmin}(\omega) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_\omega(x_i|\theta_i)} \right)$$

EventGenerator  
 $\omega$ =Neural Network  
params

“Learn all possible data generating distributions”



same  
minimum  
ver  
->  
infinity)

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

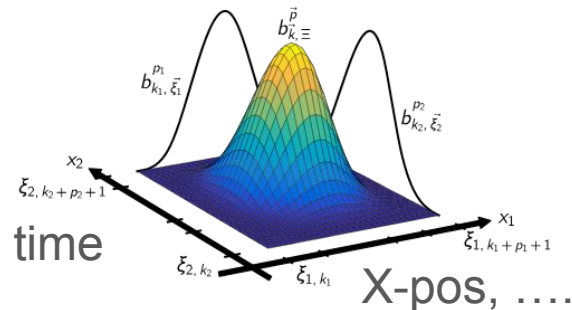
$$D_{\text{KL}}(P|q)[\omega] = \int P(x|\theta) \cdot P(\theta) \ln \left( \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right) dx + \text{Const}$$

SplineReco (table fit)  
 $\omega$ =B-Spline coefficients

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) + \text{Const}$$

$$\longrightarrow \underset{\omega}{\text{argmin}} \left( \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) \right) = \underset{\omega}{\text{argmin}} \left( \frac{1}{N} \sum_i -\ln (q_{\omega}(x_i|\theta_i)) \right)$$

6-d Tensor-Product Spline (arXiv:1301.2184)



**(higher order) Max. Likelihood**

**“Likelihood fit” of the likelihood function**

**“Learn all possible data generating distributions”**

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

$$D_{\text{KL}}(P|q)[\omega] = \int P(x|\theta) \cdot P(\theta) \ln \left( \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right) dx + \text{Const}$$

SplineReco (table fit)  
 $\omega$ =B-Spline coefficients

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) + \text{Const}$$

$$\longrightarrow \text{argmin}(\omega) \quad \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) = \text{argmin}(\omega) \quad \frac{1}{N} \sum_i -\ln(q_{\omega}(x_i|\theta_i))$$

EventGenerator  
 $\omega$ =Neural Network  
params

**(higher order) Max. Likelihood**

**“Likelihood fit” of the likelihood function**

**“Learn all possible data generating distributions”**

Hint: standard parametrization

## Interactive Question 4

Can we switch the order of parametrization?

$q(x|\theta) \rightarrow q(\theta|x)$

What might be difficulties?

Would that help us?

$$D_{\text{KL}}(P \parallel q) = \int P(x, \theta) \log \frac{P(x, \theta)}{q_{\omega}(x|\theta)q(\theta)} dx$$


$$\frac{1}{N} \sum_i -\ln (q_{\omega}(x_i|\theta_i))$$

# Learning objectives lecture 1:

- Bayes' Theorem and the “joint” distribution
- Frequentist / Bayesian parameter estimation
- Information theory: KL divergence and its interpretation
- **likelihood-based vs likelihood-free methods**

# PDF-approximation matching via “forward” KL divergence

True distribution:  $P(x|\theta)$

Approximation:  $q(x|\theta)[\omega]$

$$D_{\text{KL}}(P|q)[\omega] = \int P(x|\theta) \cdot P(\theta) \ln \left( \frac{P(x|\theta)}{q_{\omega}(x|\theta)} \right) dx + \text{Const}$$

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) + \text{Const}$$

$$\longrightarrow \operatorname{argmin}(\omega) \frac{1}{N} \sum_i \ln \left( \frac{P(x_i|\theta_i)}{q_{\omega}(x_i|\theta_i)} \right) = \operatorname{argmin}(\omega) \frac{1}{N} \sum_i -\ln(q_{\omega}(x_i|\theta_i))$$

**(higher order) Max. Likelihood**

**“Likelihood fit” of the likelihood function**

**“Learn all possible data generating distributions”**

# (Posterior) PDF-approximation matching via “forward” KL divergence

True distribution:  $P(\theta|x)$

Approximation:  $q(\theta|x)[\omega]$

$$D_{\text{KL}}(P|q)[\omega] = \int P(\theta|x) \cdot P(x) \ln \left( \frac{P(\theta|x)}{q_{\omega}(\theta|x)} \right) dx + \text{Const}$$

$$D_{\text{KL}}(P|q)[\omega] \approx \frac{1}{N} \sum_i \ln \left( \frac{P(\theta_i|x_i)}{q_{\omega}(\theta_i|x_i)} \right) + \text{Const}$$

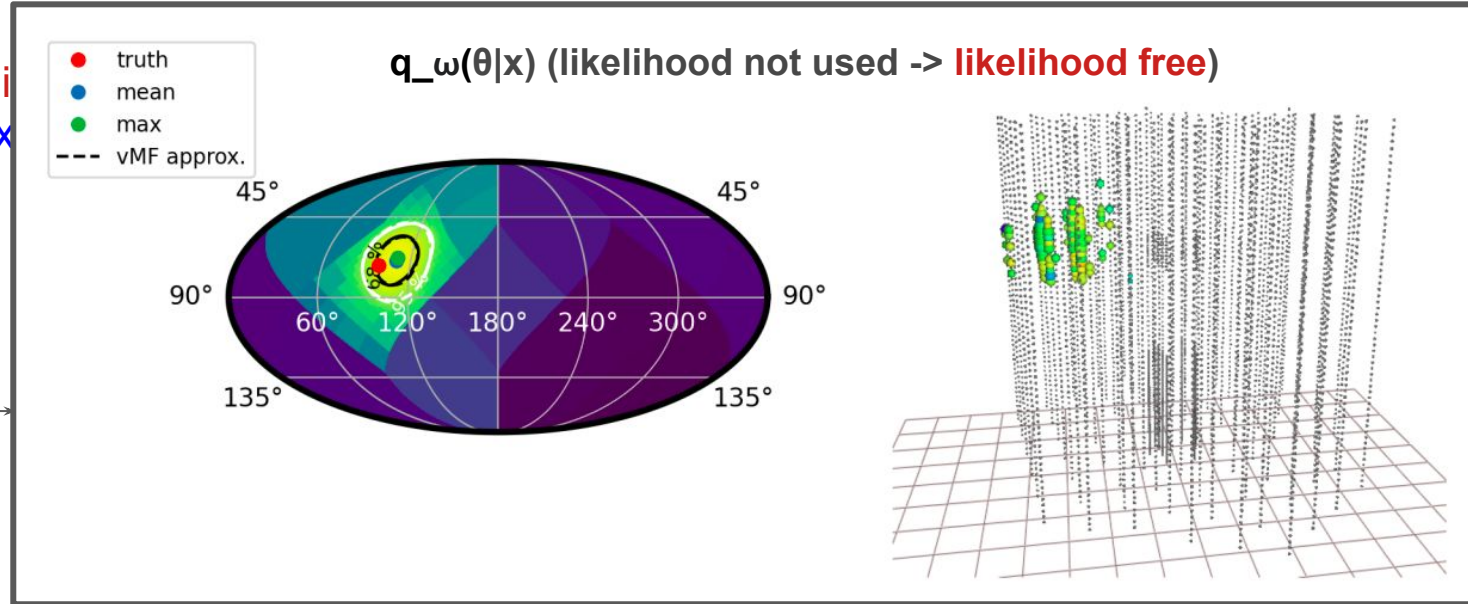
$$\longrightarrow \operatorname{argmin}(\omega) \frac{1}{N} \sum_i \ln \left( \frac{P(\theta_i|x_i)}{q_{\omega}(\theta_i|x_i)} \right) = \operatorname{argmin}(\omega) \frac{1}{N} \sum_i -\ln (q_{\omega}(\theta_i|x_i))$$

**Likelihood-free inference**

“Learn all possible **posterior distributions**”

# (Posterior) PDF-approximation matching via “forward” KL divergence

True dist  
Approx



Same  
Minimum  
Over  
 $\omega$   
( $N \rightarrow$   
Infinity)

Inference

Neural posterior estimation  
of IceCube neutrino direction  
2604.19846

(“amortized neural posterior estimation”)  
Papamakarios (2016)

“Learn all possible **posterior**  
distributions”

“Likelihood-free” inference

Neural Posterior Estimation (**NPE**)

$$q_{\omega}(\theta|x)$$

Neural Ratio Estimation (**NRE**)

$$r_{\phi}(\theta, x) \approx \frac{p(x|\theta)}{p(x)} = \frac{p(\theta|x)}{p(\theta)}$$

Bayes' Theorem

$$p(\theta|x) = \frac{p(x|\theta) \cdot p(\theta)}{p(x)}$$

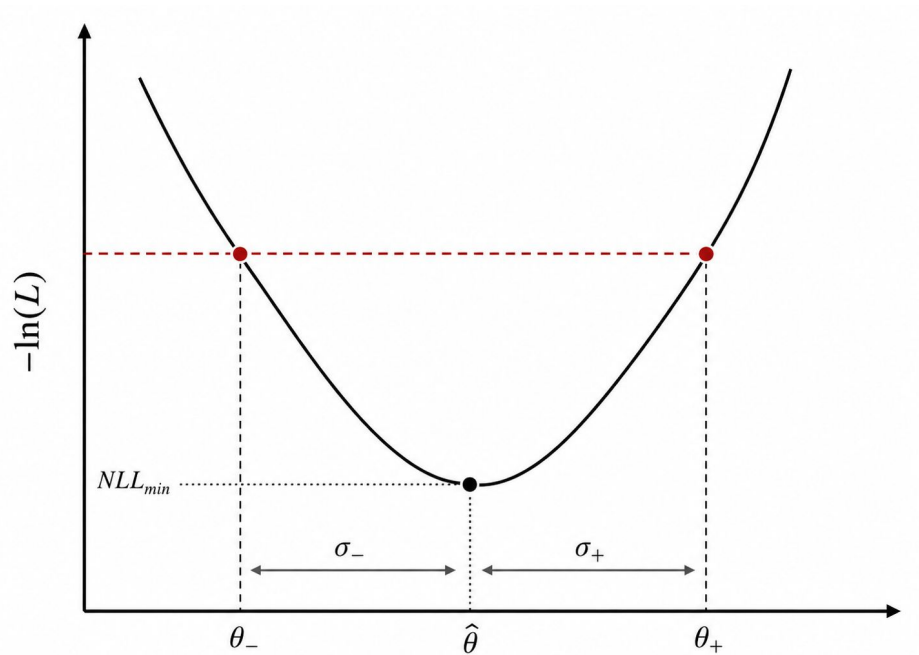
Trained with classifier

Proportional to likelihood, sampling still required

# Summary:

## Likelihood-based

Max. Likelihood (Frequentist)

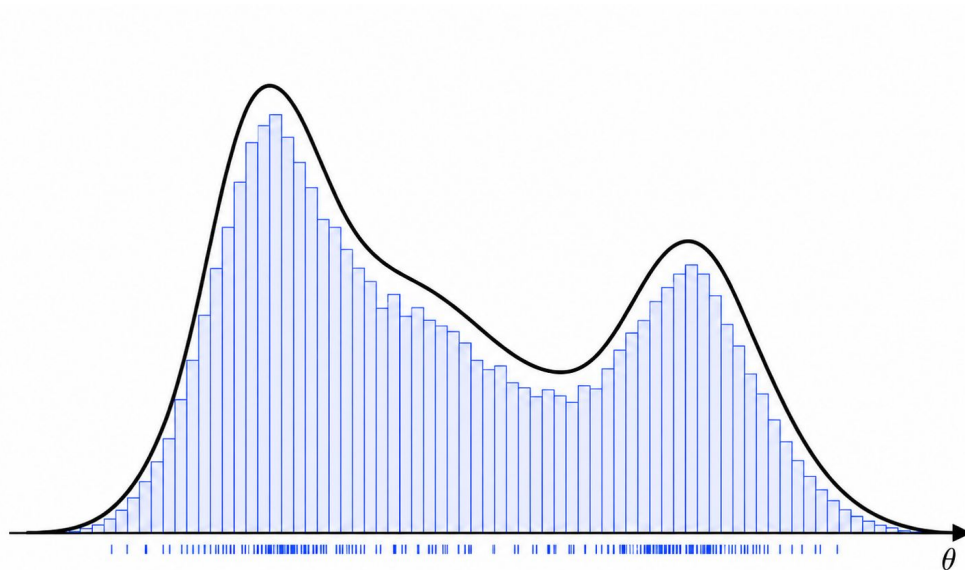


# Summary:

## Likelihood-based

Max. Likelihood (Frequentist)

MCMC (Bayesian)



$$\alpha = \min \left( 1, \frac{\pi(\theta')}{\pi(\theta_t)} \right)$$

# Summary:

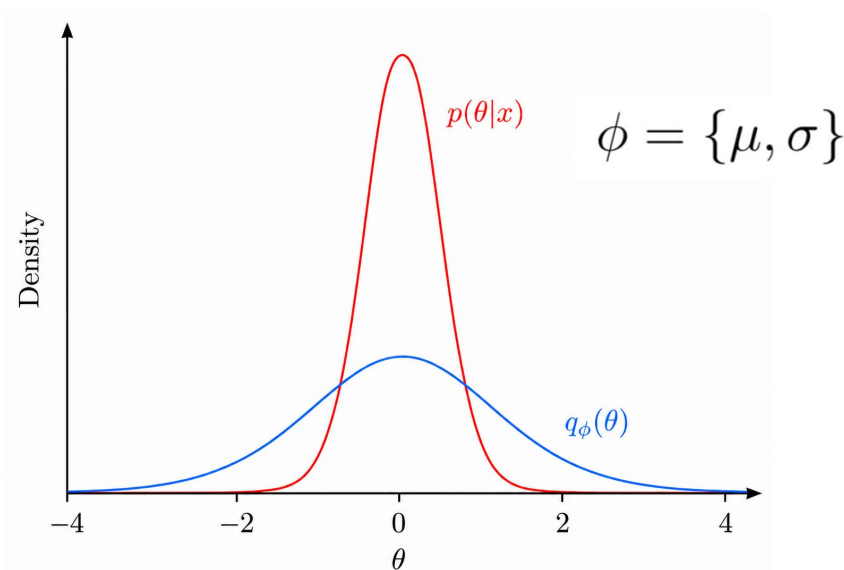
## Likelihood-based

Max. Likelihood (Frequentist)

MCMC (Bayesian)

ELBO-based

Variational Inference (Bayesian)



$$p(x) \geq \underbrace{\int_{\theta} q_{\phi}(\theta) \cdot \left[ \ln p(x; \theta) - \ln \frac{q_{\phi}(\theta)}{p(\theta)} \right] d\theta}_{\text{evidence lower bound (ELBO)}}$$

# Summary:

## Information Theory!

-> KL divergence of **joint** distribution

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$

### Likelihood-based

Max. Likelihood (Frequentist)

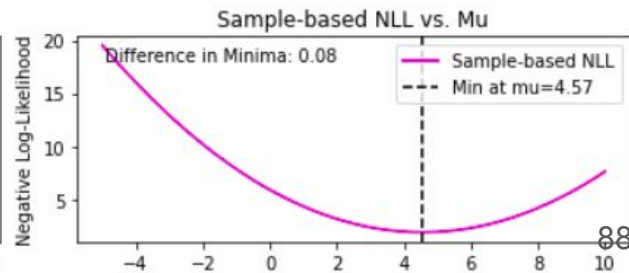
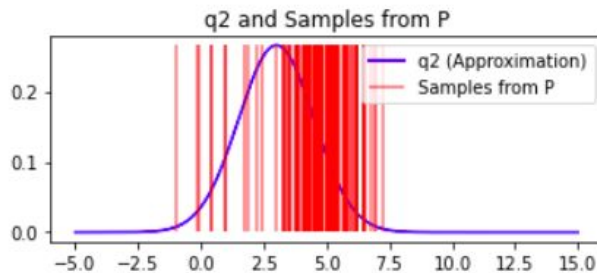
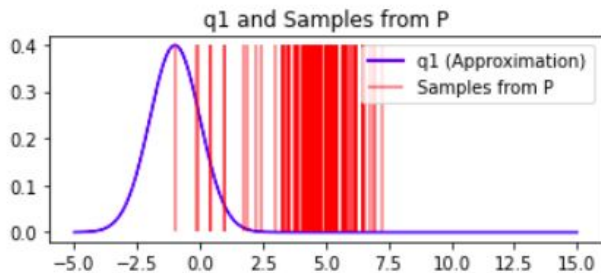
MCMC (Bayesian)

ELBO-based

Variational Inference (Bayesian)

$$q_{\omega}(x|\theta)$$

Learn “likelihood”



# Summary:

## Information Theory!

-> KL divergence of **joint** distribution

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$

### Likelihood-based

Max. Likelihood (Frequentist)

MCMC (Bayesian)

ELBO-based

Variational Inference (Bayesian)

$$q_{\omega}(x|\theta)$$

Learn “likelihood”

Per event

**Extra numerical step:**

Use likelihood to get

Posterior / Confidence Interval

Lecture 2/3

# Summary:

## Information Theory!

### Likelihood-based

- Max. Likelihood (Frequentist)
- MCMC (Bayesian)
- ELBO-based Variational Inference (Bayesian)

Per event

**Extra numerical step:**  
Use likelihood to get  
Posterior / Confidence Interval

Lecture 2/3

-> KL divergence of **joint** distribution

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$

$$q_{\omega}(x|\theta)$$

Learn "likelihood"

Learn all posteriors

$$q_{\omega}(\theta|x)$$

### Likelihood-free

Amortized neural posterior estimation

# Summary:

## Information Theory!

### Likelihood-based

Max. Likelihood (Frequentist)

MCMC (Bayesian)

ELBO-based  
Variational Inference (Bayesian)

Per event

**Extra numerical step:**

Use likelihood to get  
Posterior / Confidence Interval

Lecture 2/3

-> KL divergence of **joint** distribution

$$p(\theta, x) = p(\theta|x) \cdot p(x) = p(x|\theta) \cdot p(\theta)$$

$$q_{\omega}(x|\theta)$$

Learn "likelihood"

Learn all posteriors

$$q_{\omega}(\theta|x)$$

**Likelihood-free**

**No extra step:**

-> direct Posterior  
-> Standard NN  
prediction special  
case

Amortized neural  
posterior estimation