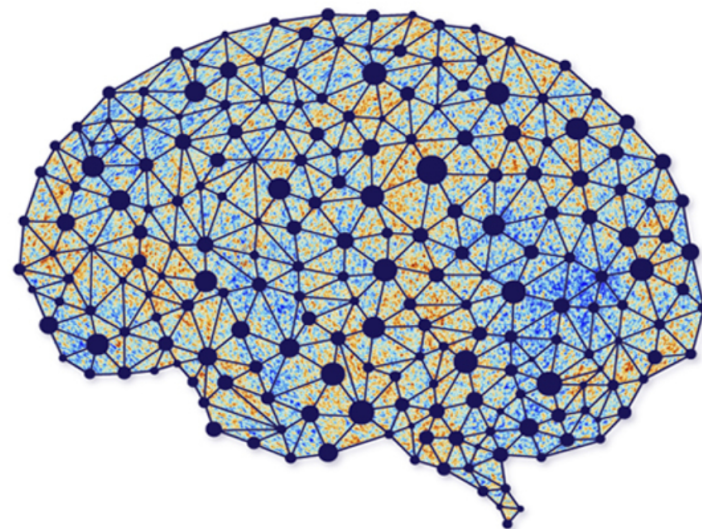


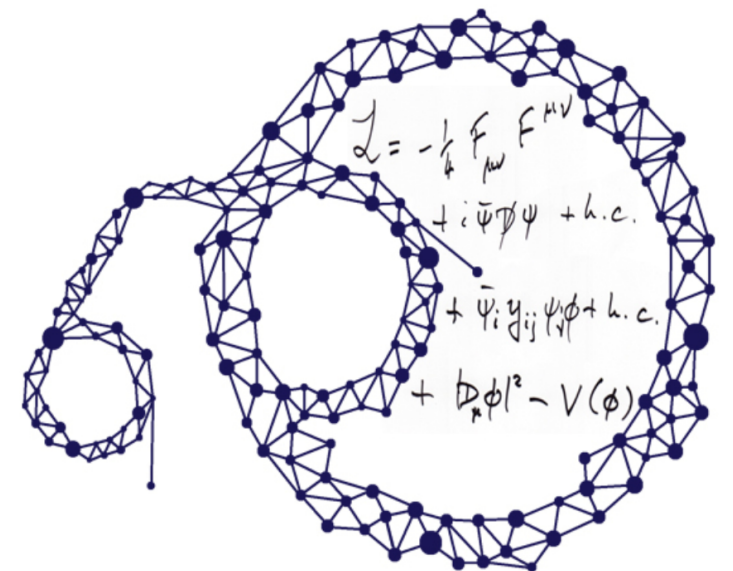
Physics 361 - Machine Learning in Physics

Lecture 13 – Simulation-Based Inference 2

March 4th 2025



AI
∩
Universe



Moritz Münchmeyer

Simulation-Based Inference

Examples of SBI

- We will use the “**sbi**” python library, a popular option. There are other codes, for example the more general probability framework “pyro”.

 **sbi** v0.23.3

Search

sbi-dev/sbi
v0.23.3 637 160

[sbi](#)
[Home](#)
[Tutorials and Examples](#)
[API Reference](#)
[FAQ](#)
[Contributing](#)
[Citation](#)
[Credits](#)

sbi : simulation-based inference toolkit

sbi is a Python package for simulation-based inference, designed to meet the needs of both researchers and practitioners. Whether you need fine-grained control or an easy-to-use interface, sbi has you covered.

With sbi, you can perform parameter inference using Bayesian inference: Given a simulator that models a real-world process, SBI estimates the full posterior distribution over the simulator's parameters based on observed data. This distribution indicates the most likely parameter values while additionally quantifying uncertainty and revealing potential interactions between parameters.

sbi provides access to simulation-based inference methods via a user-friendly interface:

```
import torch
from sbi.inference import NPE

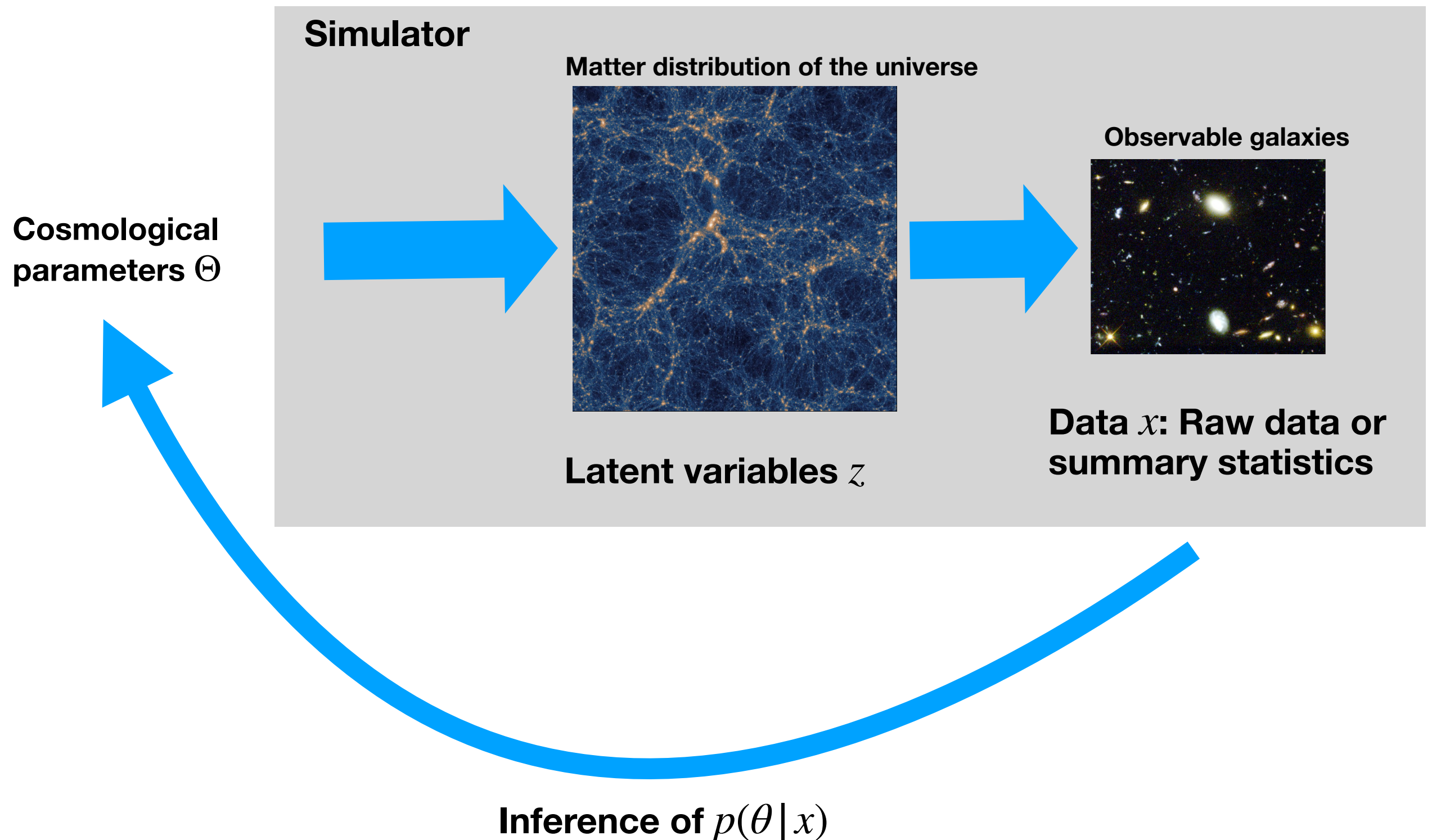
# define shifted Gaussian simulator.
def simulator(theta): return theta + torch.randn_like(theta)
# draw parameters from Gaussian prior.
theta = torch.randn(1000, 2)
# simulate data
x = simulator(theta)

# choose sbi method and train
inference = NPE()
inference.append_simulations(theta, x).train()

# do inference given observed data
x_o = torch.ones(2)
posterior = inference.build_posterior()
samples = posterior.sample((1000,), x=x_o)
```

[Table of contents](#)
[Overview](#)
[Motivation and approach](#)
[Implemented algorithms](#)
 Posterior estimation ((S)NPE)
 Likelihood-estimation ((S)NLE)
 Likelihood-ratio-estimation ((S)NRE)
 Diagnostics

Recall example from cosmology



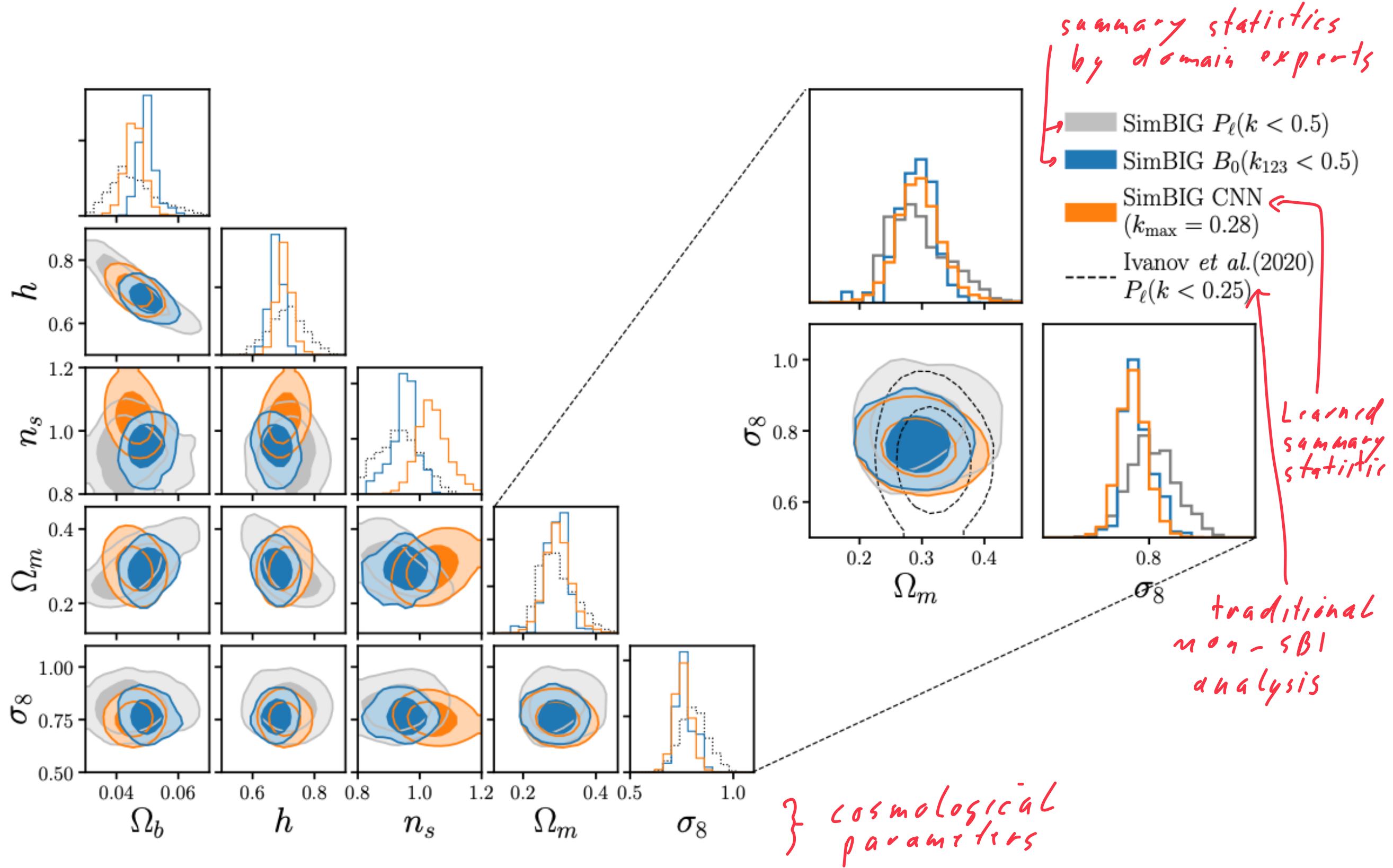


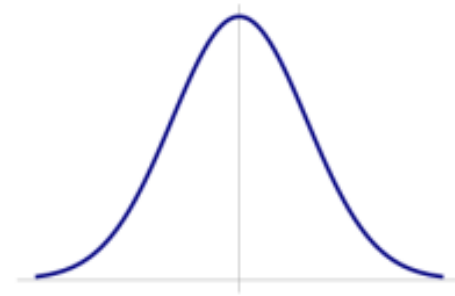
Figure 1. *Left:* Posteriors of cosmological parameters inferred from B_0 (blue) and CNN (orange) using SIMBIG. All posteriors include a ω_b prior from BBN studies. The contours mark the 68 and 95 percentiles. For comparison, we include the posterior from the SIMBIG P_ℓ analysis (gray). *Right:* We focus on the posteriors of Ω_m and σ_8 , the parameters that can be most significantly constrained by galaxy clustering. We include the posterior from the Ivanov *et al.* (2020) PT-based $P_\ell(k < 0.25 h/\text{Mpc})$ analysis for reference (black dashed). The SIMBIG B_0 and CNN constraints are significantly tighter yet consistent with P_ℓ constraints.

Simulation-Based Inference

SBI with Normalizing Flows

Illustration of the MAF flow

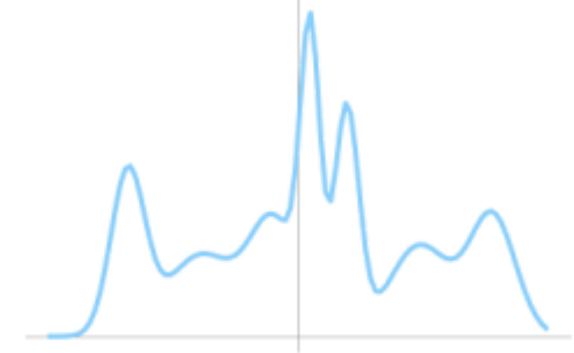
Source: 2310.03741



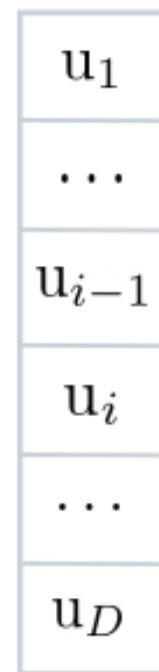
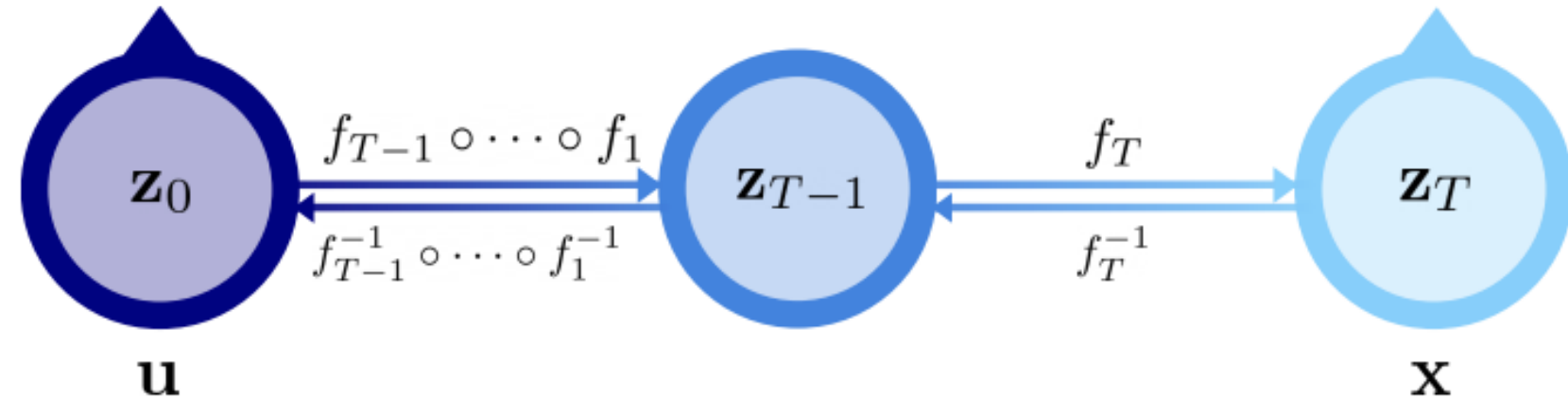
$$\mathbf{z}_0 = \mathbf{u} \sim \pi(\mathbf{u})$$

$$\mathbf{x} = f_T \circ \dots \circ f_1(\mathbf{u}) = f(\mathbf{u})$$

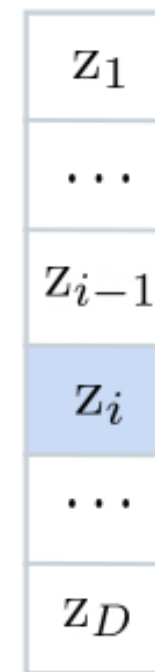
$$p(\mathbf{x}) = \pi(f^{-1}(\mathbf{x})) \left| \det \left(\frac{\partial f^{-1}}{\partial \mathbf{x}} \right) \right|$$



$$\mathbf{z}_T = \mathbf{x} \sim p(\mathbf{x})$$



...

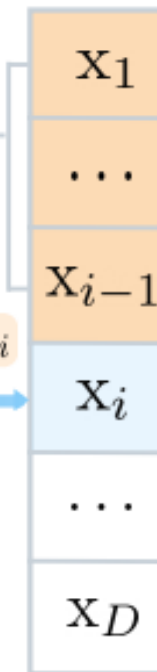


$$c_i(\mathbf{x}_{<i})$$

$$\mathbf{h}_i = \{\alpha_i, \mu_i\}$$

$$\mathbf{x}_i = \mathbf{z}_i \exp(\alpha_i) + \mu_i$$

$$\mathbf{z}_i = (\mathbf{x}_i - \mu_i) \exp(-\alpha_i)$$

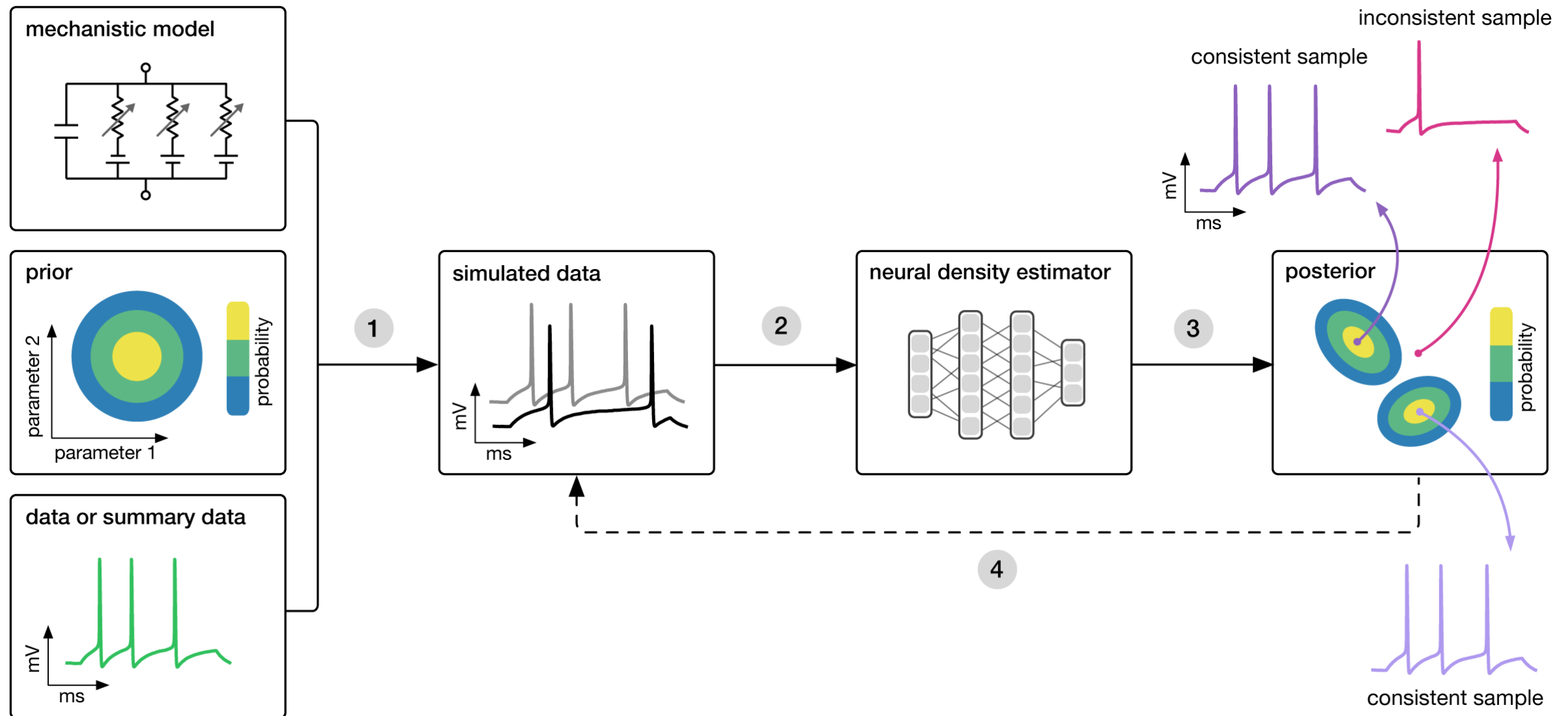


$\{\alpha_i, \mu_i\}$ depend
on $\mathbf{x}_{<i}$ AND
the conditioner,
through a neural
network

Figure 2. Diagram of how normalizing flows work, with the specific example of Masked Autoregressive Flows. The samples from the vector $\mathbf{z}_0 = \mathbf{u}$, sampled from the simple distribution $\pi(\mathbf{u})$, are deformed through the sequence of transformations $f = f_T \circ \dots \circ f_1$ into those of $\mathbf{z}_T = \mathbf{x}$, which follow a more complex distribution $p(\mathbf{x})$. In the lower panel, we illustrate the conditioner that “masks out” the connections between $\mathbf{z}_i$ and $\mathbf{h}_{<i}$, as well as the affine functions applied to the vector components.

Simulation-based inference with NDEs

inputs: A candidate mechanistic model, prior knowledge or constraints on model parameters, and observational data (or summary statistics thereof).



Goal: Algorithmically identify mechanistic models (simulators) which are consistent with data.

Recall: Neural Likelihood Estimation (NLE)

- In neural likelihood estimation we learn the likelihood from simulated pairs of model parameters and data:

$$\{(\theta_n, x_n)\}_{n=1}^N \quad \longrightarrow \quad \mathcal{L}(x \mid \theta)$$

- On the next slide we will learn the required training objective for the normalizing flow.
- After training the NDE on our simulated data, we can then evaluate the likelihood of observed data from our measurement.

$$\mathcal{L}(x^{obs} \mid \theta)$$

- Now we can proceed with normal Bayesian data analysis. That usually means that we sample from the posterior with MCMC:

$$p(\theta \mid x^{obs}) \propto \mathcal{L}(x^{obs} \mid \theta)p(\theta)$$

Neural Likelihood estimation with NFs

- From a simulated data set $\{(\theta_n, x_n)\}_{n=1}^N$ drawn from a proposal density $\tilde{p}(\theta)$ we want to approximate the target likelihood $p(x|\theta)$ by a conditional normalizing flow $q_\phi(x|\theta)$, where $\tilde{p}(\theta, x) \stackrel{(2)}{=} p(x|\theta)\tilde{p}(\theta)$. $\leftarrow$ joint PDF of our samples $\{\theta, x\}$
- We thus minimize the expected KL-divergence with respect to the flow parameters ϕ .

Recall from lecture 3: $D_{\text{KL}}(P||Q) \stackrel{(1)}{=} \mathbb{E}_{X \sim P} \left[\log \frac{P(X)}{Q(X)} \right]$

$$\begin{aligned} \min_{\phi} \mathbb{E}_{\tilde{p}(\theta)} \left[D_{\text{KL}} \left[p(\mathbf{x}|\theta) \parallel q_{\phi}(\mathbf{x}|\theta) \right] \right] &\stackrel{(1)}{=} \min_{\phi} \int d\theta \tilde{p}(\theta) \int d\mathbf{x} p(\mathbf{x}|\theta) \log \left(\frac{p(\mathbf{x}|\theta)}{q_{\phi}(\mathbf{x}|\theta)} \right) \\ &\stackrel{(2)}{=} \min_{\phi} \int d\theta d\mathbf{x} \tilde{p}(\theta, \mathbf{x}) \log \left(\frac{p(\mathbf{x}|\theta)}{q_{\phi}(\mathbf{x}|\theta)} \right) = \log p - \log q \\ &= \min_{\phi} - \mathbb{E}_{\tilde{p}(\theta, \mathbf{x})} \left[\log q_{\phi}(\mathbf{x}|\theta) \right] + \text{const.} \\ &\approx \min_{\phi} - \frac{1}{N_{\text{sim}}} \sum_{n=1}^{N_{\text{sim}}} \log q_{\phi}(\mathbf{x}_n|\theta_n) + \text{const.}, \end{aligned}$$

eval expectation value via sampling

- Minimizing the expected KL is thus the same as maximum likelihood training (see lecture 3).

Recall: Neural Posterior Estimation (NPE)

- You might wonder why why learn the likelihood and not the posterior which is our ultimate goal. Learning the posterior is indeed a possibility.
- From a simulated data set

$$\{(\theta_n, x_n)\}_{n=1}^N$$

drawn from a proposal density $\tilde{p}(\theta)$ it is possible to directly learn the posterior

$$p(\theta | x)$$

- On the next slide we will learn the training objective that makes this possible.
- An advantage of learning the posterior directly is that we do not need to run an MCMC anymore. The model directly outputs the desired posterior, i.e. our parameter measurement. A disadvantage is that it is difficult to explore different prior distributions of the parameters.

Neural posterior estimation with NFs

- From a simulated data set $\{(\theta_n, x_n)\}_{n=1}^N$ drawn from a prior $p(\theta)$ we want to approximate the target posterior $p(\theta|x)$ by a conditional normalizing flow $q_\phi(\theta|x)$.

$p(\theta, x)$: joint PDF of our samples

- We minimize the KL-divergence with the joint PDF as follows.

$p(\theta|x)p(x)$

$$\min_{\phi} D_{KL} \left(p(\theta, x) \parallel \underbrace{q_\phi(\theta|x)p(x)}_{q_\phi(\theta, x)} \right) = \min_{\phi} \int p(\theta, x) \log \frac{p(\theta, x)}{q_\phi(\theta|x)p(x)} d\theta dx$$

$q_\phi(\theta, x)$

$$= \min_{\phi} \int p(\theta, x) \log \frac{p(\theta|x)}{q_\phi(\theta|x)} d\theta dx$$

*eval expectation
value via sampling*

$$\approx \min_{\phi} \sum_i \log p(\theta_i|x_i) - \log q_\phi(\theta_i|x_i)$$

$$\approx \min_{\phi} \sum_i -\log q_\phi(\theta_i|x_i)$$

- As in the previous case we can thus learn the posterior by sampling from the simulator.

Some examples

Real-time gravitational-wave science with neural posterior estimation

Maximilian Dax,^{1,*} Stephen R. Green,^{2,†} Jonathan Gair,^{2,‡}
Jakob H. Macke,^{1,3} Alessandra Buonanno,^{2,4} and Bernhard Schölkopf¹

¹*Max Planck Institute for Intelligent Systems, Max-Planck-Ring 4, 72076 Tübingen, Germany*

²*Max Planck Institute for Gravitational Physics (Albert Einstein Institute), Am Mühlenberg 1, 14476 Potsdam, Germany*

³*Machine Learning in Science, University of Tübingen, 72076 Tübingen, Germany*

⁴*Department of Physics, University of Maryland, College Park, MD 20742, USA*

We demonstrate unprecedented accuracy for rapid gravitational-wave parameter estimation with deep learning. Using neural networks as surrogates for Bayesian posterior distributions, we analyze eight gravitational-wave events from the first LIGO-Virgo Gravitational-Wave Transient Catalog and find very close quantitative agreement with standard inference codes, but with inference times reduced from $O(\text{day})$ to 20 seconds per event. Our networks are trained using simulated data, including an estimate of the detector-noise characteristics near the event. This encodes the signal and noise models within millions of neural-network parameters, and enables inference for any observed data consistent with the training distribution, accounting for noise nonstationarity from event to event. Our algorithm—called “DINGO”—sets a new standard in fast-and-accurate inference of physical parameters of detected gravitational-wave events, which should enable real-time data analysis without sacrificing accuracy.

<https://arxiv.org/pdf/2106.12594.pdf>

 > astro-ph > arXiv:2301.06575

Search..
Help | Ad

Astrophysics > Earth and Planetary Astrophysics

[Submitted on 16 Jan 2023 (v1), last revised 10 Feb 2023 (this version, v2)]

Neural posterior estimation for exoplanetary atmospheric retrieval

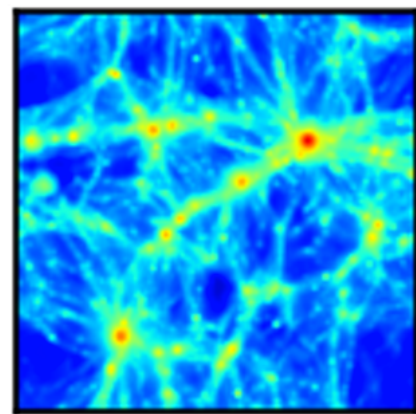
Malavika Vasist, François Rozet, Olivier Absil, Paul Mollière, Evert Nasedkin, Gilles Louppe

Retrieving the physical parameters from spectroscopic observations of exoplanets is key to understanding their atmospheric properties. Exoplanetary atmospheric retrievals are usually based on approximate Bayesian inference and rely on sampling-based approaches to compute parameter posterior distributions. Accurate or repeated retrievals, however, can result in very long computation times due to the sequential nature of sampling-based algorithms. We aim to amortize exoplanetary atmospheric retrieval using neural posterior estimation (NPE), a simulation-based inference algorithm based on variational inference and normalizing flows. In this way, we aim (i) to strongly reduce inference time, (ii) to scale inference to complex simulation models with many nuisance parameters or intractable likelihood functions, and (iii) to enable the statistical validation of the inference results. We evaluate NPE on a radiative transfer model for exoplanet spectra petitRADTRANS, including the effects of scattering and clouds. We train a neural autoregressive flow to quickly estimate posteriors and compare against retrievals computed with MultiNest. NPE produces accurate posterior approximations while reducing inference time down to a few seconds. We demonstrate the computational faithfulness of our posterior approximations using inference diagnostics including posterior predictive checks and coverage, taking advantage of the quasi-instantaneous inference time of NPE. Our analysis confirms the reliability of the approximate posteriors produced by NPE. The accuracy and reliability of the inference results produced by NPE establishes it as a promising approach for atmospheric retrievals. Amortization of the posterior inference makes repeated inference on several observations computationally inexpensive since it does not require on-the-fly simulations, making the retrieval efficient, scalable, and testable.

<https://arxiv.org/abs/2301.06575>

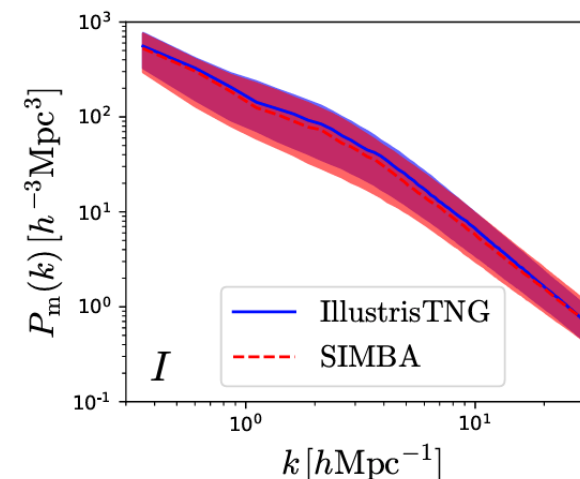
Our problem set: SBI on CAMELS simulations

- We will re-use the cosmological data we explored in the last problem set to perform SBI.
- The input of our analysis will be summary statistics of the CAMELS data: The matter power spectrum.
- The output will be cosmological parameters.
- We will apply NPE to estimate cosmological parameters from simulations.



simulations
(fct. of cosmo.
parameters λ)

summary
statistic
→
already
done



Power spectrum

set of
numbers
 $\{P(k_i)\}$

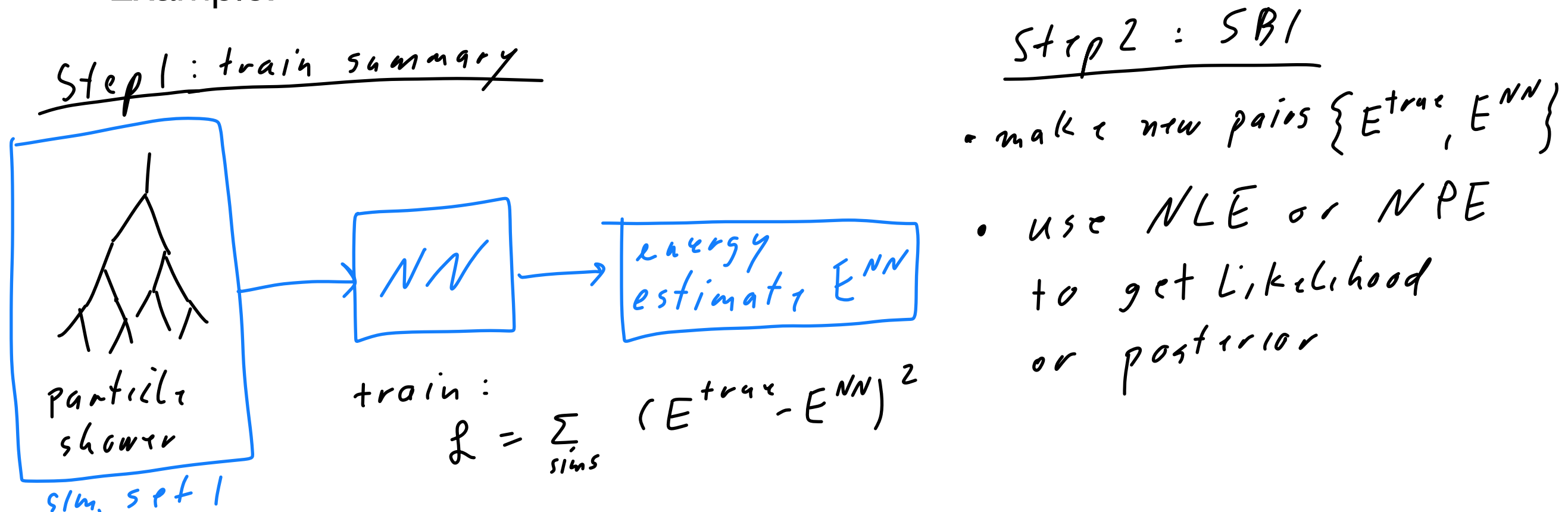
Goal: Learn $P(\lambda | \{P(k_i)\})$

Simulation-Based Inference

More on Uncertainty with
Neural Networks

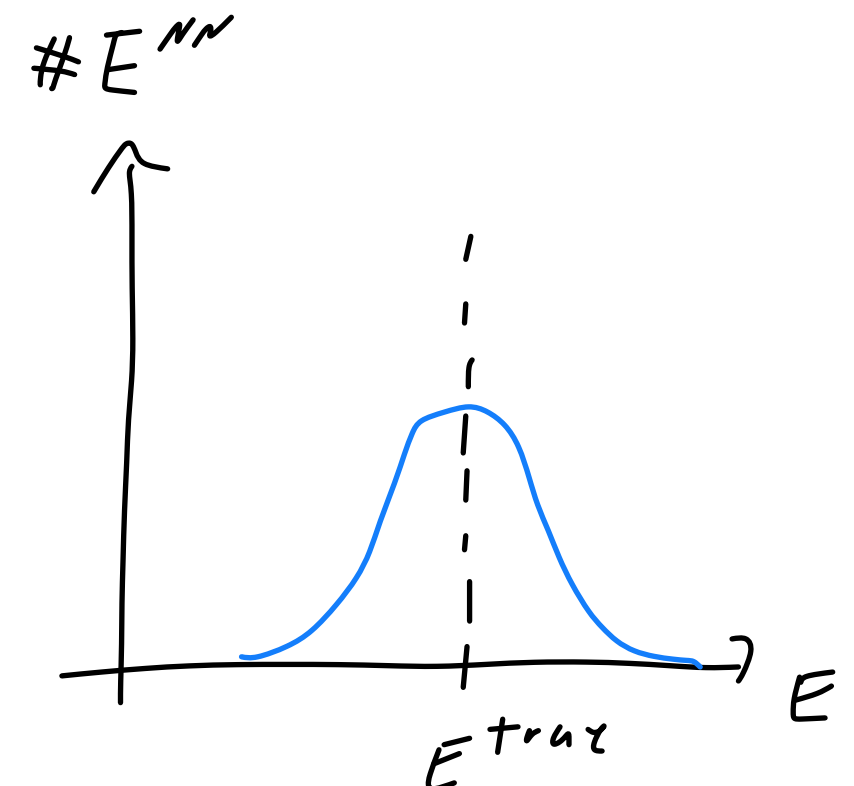
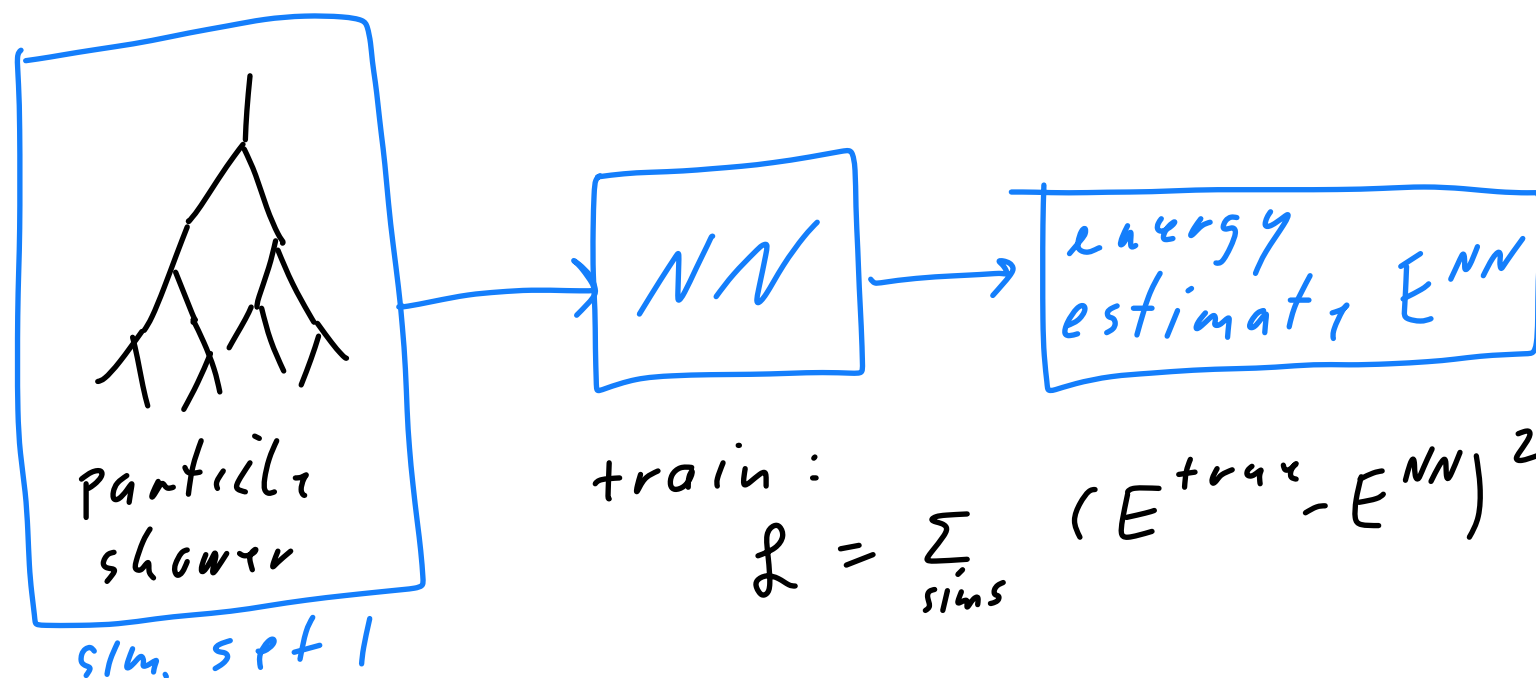
Neural networks as summary statistics

- Above we have developed SBI. Usually the data that goes into SBI is not the raw data but some summary statistic thereof.
- If we simply “train a neural network” to learn a parameter, the neural network can be considered as such a summary statistic.
- We first train the neural network to extract the parameters, e.g. using MSE loss. Then we learn the posterior or the likelihood of the neural network observable with SBI (e.g. using a flow). In this way we **get statistical “error bars” for the neural network measurement**.
- Example:



Summary statistics are often Gaussian

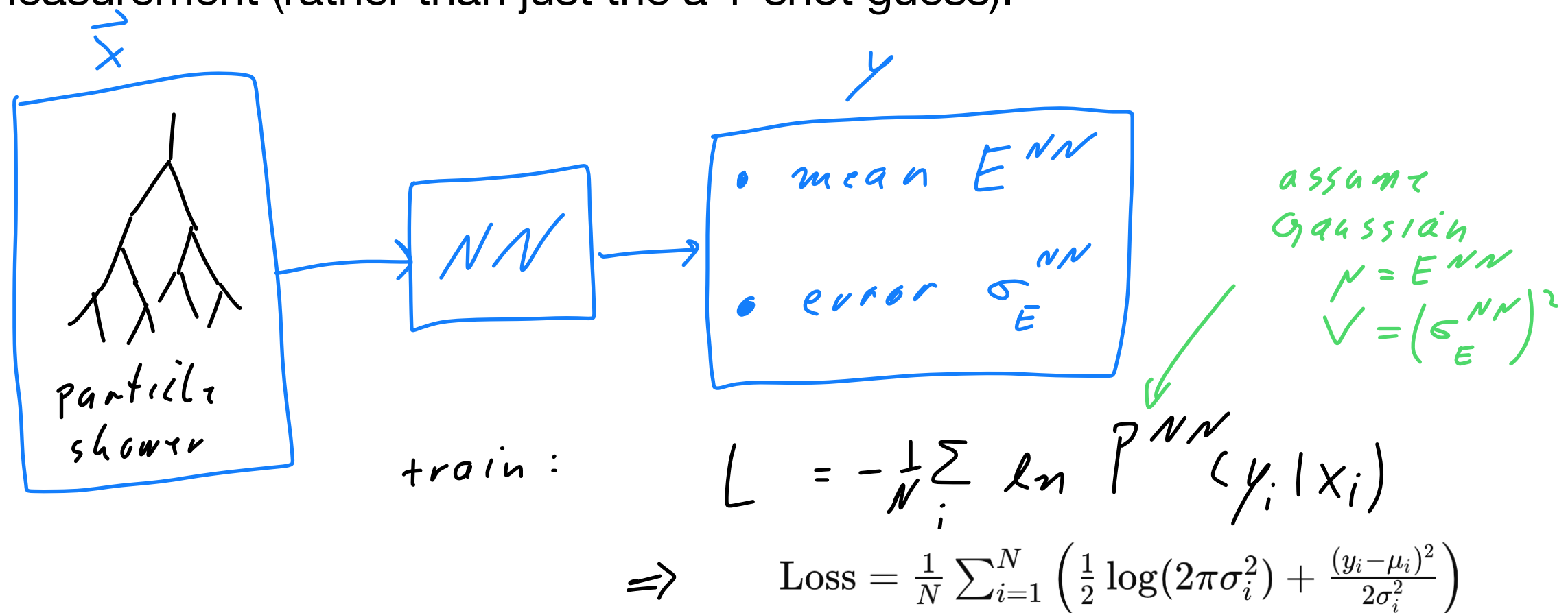
- **Summary statistics are often Gaussian** due to the Central Limit Theorem. This is also true for the output of neural networks (i.e. neural network summary statistics).
- If we can assume Gaussianity (which we can test using simulations), then we can **estimate the mean and covariance of the trained neural network** output by running many simulations through the trained network. **In this case full SBI with a normalizing flow is not necessary.**



The trained NN gives estimates that are Gaussian around the truth.

Neural networks can be trained to output error bars

- The method we advocated before to assign error bars is to first train the neural network, then keep the weights fixed, and get error bars from SBI.
- However **it is also possible to train the neural network to output error bars**. For example, for some samples the network may be more confident in the measurement than for other samples. This confidence can be learned from the data.
- A typical strategy is to make the neural network output the mean and the variance of the measurement (rather than just the a 1-shot guess).



- Note however that **there is no formal guarantee that the learned error bars will be correct**. SBI is a somewhat more systematic approach.

Types of uncertainty

- As you know, there are several types of uncertainty in an analysis:
 - **Statistical uncertainty** (also called **aleatoric uncertainty** in the context of machine learning): this is due to the inherent limitation of the data (e.g. detector noise) and the analysis method (e.g. summary statistics information loss). It can be reduced by getting more data.
 - **Systematic uncertainty**: this is due to biases and flaws in the measurement. It cannot be reduced by taking more data but only by making better measurements with less bias.
- In machine learning we often use a third concept, **epistemic uncertainty**.
 - Epistemic uncertainty is similar to systematic uncertainty but it is not exactly the same thing.
 - **Epistemic uncertainty deals with the uncertainty due to the model's inadequacy** while systematic uncertainty usually means flaws in the measurement method.
 - Example of epistemic uncertainty: **Several neural networks trained on the same data give slightly different measurements** and error bars. We do not know which (if any) is correct.

Epistemic uncertainty in neural networks

- The **basic idea is to have an ensemble of neural network predictions**. If all agree, epistemic uncertainty is low; if they disagree, it is high. There are different versions of that idea:
- **Bayesian Neural Networks**
 - Bayesian Neural Networks (BNNs) are an approach to model epistemic uncertainty by **placing a probability distribution over the weights** of the network. By doing this, BNNs not only provide predictions but also give a measure of uncertainty in those predictions based on the posterior distribution of the weights.
 - **Advantages:** Provides a principled way of quantifying uncertainty.
 - **Challenges:** Computationally expensive and often complex to implement due to the need for approximations like variational inference or Monte Carlo methods to compute the posterior distributions.
 - Typical weight distribution: Gaussian, uncorrelated. But there is no one-size-fits-all answer, and the choice of prior can significantly affect the model's performance and its inferred posterior distributions.

Epistemic uncertainty in neural networks

- **Monte Carlo Dropout**

- Monte Carlo (MC) Dropout is a practical and widely used method for estimating uncertainty in neural networks. It involves applying dropout not only during training but also **at test time**, allowing the model to generate different outputs by randomly dropping units during multiple forward passes.
- **Advantages:** Easy to implement and less computationally intensive than full Bayesian methods.
- **Challenges:** The estimates of uncertainty may depend significantly on the dropout rate and the number of stochastic forward passes.

- **Deep Ensembles**

- Deep Ensembles involve **training multiple versions of the same model on the same data, but with different initializations or even slightly different model architectures**. The variance in the predictions from these models can be used as a measure of epistemic uncertainty.
- **Advantages:** Often provides improved prediction accuracy and uncertainty estimation over the other methods.
- **Challenges:** Requires more computational resources since multiple models need to be trained and stored.

Epistemic uncertainty in neural networks

- Some comments on these methods:
 - <https://arxiv.org/abs/2004.10710> **Deeply Uncertain:** Comparing Methods of Uncertainty Quantification in Deep Learning Algorithms
 - Often one reports the statistical and epistemic uncertainty as two different errors.

$$\sigma_{pr} = \sqrt{\sigma_{al}^2 + \sigma_{ep}^2}$$

- <https://www.sciencedirect.com/science/article/pii/S1566253521001081> A review of uncertainty quantification in deep learning: Techniques, applications and challenges
- **Epistemic uncertainty in SBI:** By treating a NN as a summary statistic it is not important that the NN is exact. A bias would be corrected for in the likelihood or posterior (though the NN could be sub-optimal if not trained out or too little capacity). However the ML model used in the SBI, i.e. the normalizing flow, can give incorrect probabilities. One could train several of them to test their agreement (as part of the tests of SBI results).
- If the **data is out of the training data distribution** (i.e. the simulations are “wrong”) all bets are off (**“unknown unknowns”**). So we put a lot of effort into making them reliable in the range where we use them, and also sample over uncertain parameters (**“known unknowns”**).

Simulation-Based Inference

A quick look into explicit inference with
Differentiable Simulations

Recall: Implicit vs Explicit inference

- In most situations a simulation does not provide a probability density (likelihood) $\mathcal{L}(x | \theta)$ of observations given parameters. Such simulations are sometimes called **implicit models**.
- Implicit means that their **likelihood cannot be computed explicitly**, i.e. it is not computationally tractable. We only get samples of the simulation.
- On the other hand, models or simulations that do provide a likelihood are called **explicit models**. Recall for example Gaussian likelihoods.
- A key problem in **explicit inference is to marginalize over the latent variables**, such as the random initial conditions of a simulation.

$$p(x|\theta) = \int dz \, p(x, z|\theta)$$

- For comparison to our study of SBI I now want to discuss a specific example of explicit inference. This approach is also sometimes called “**forward modelling**” or **Bayesian hierarchical modelling**. This is a general approach to many problems that is important to know.

Explicit inference with probabilistic forward modelling

- Assume that we have a simulator, called the “**forward model**” f , depending on **parameters** θ that describes the evolution of a system starting from a **signal** $\mathbf{s}$ (e.g. the initial conditions of the forward process):

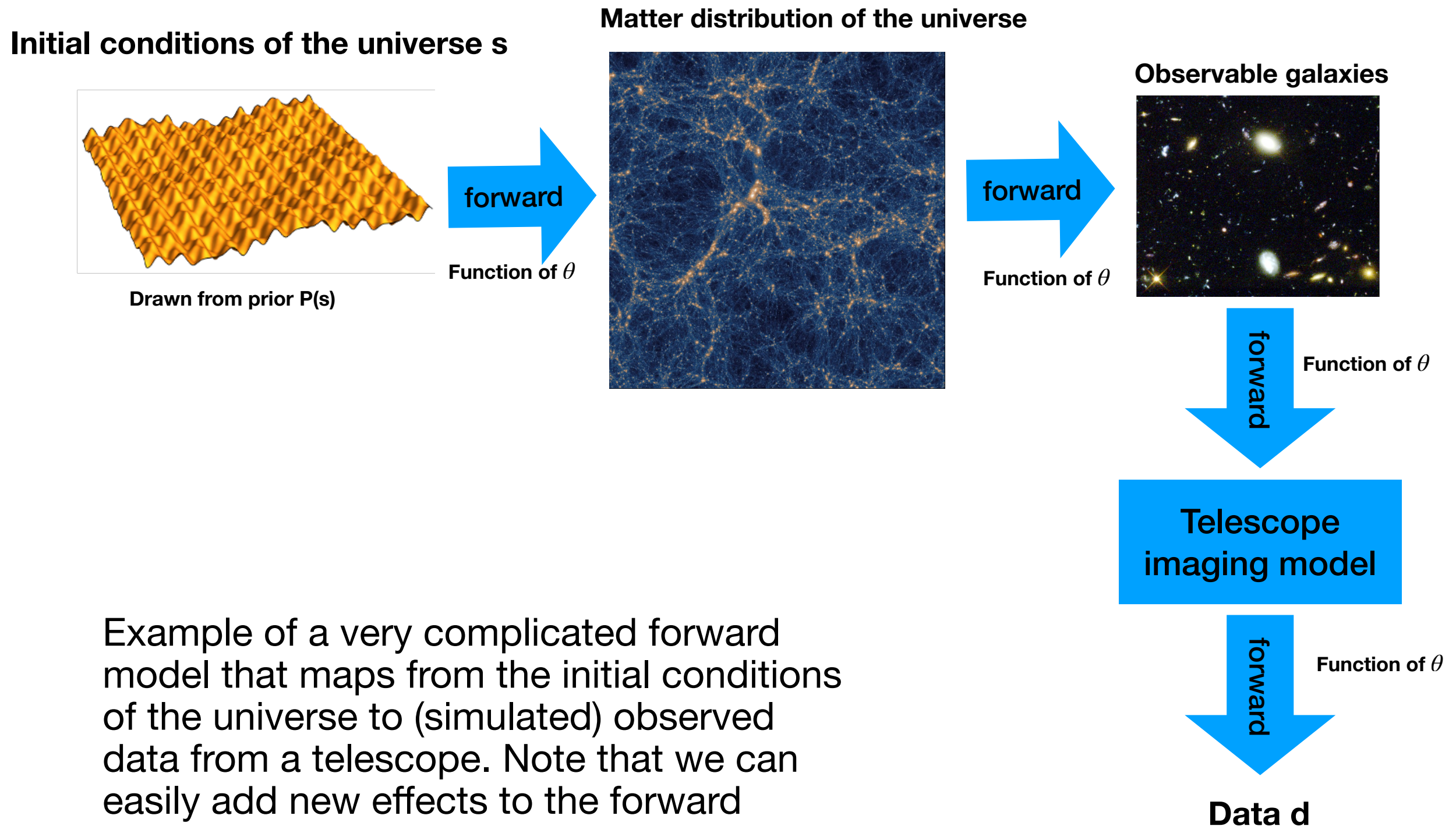
$$f(\mathbf{s}, \Theta)$$

- We want to infer the **parameters** θ and the **signal** $\mathbf{s}$ from **data** $\mathbf{d}$, which we take to be a **the forward evolved signal plus some observational noise**.

$$\mathbf{d} = f(\mathbf{s}) + n$$

- This setup is also called an “**inverse problem**”, i.e. we want to reconstruct the signal from the data, assuming that we know the forward model as a function of some parameters. Inverse problems are generally ill-posed (need to regularize).

Example: Cosmology forward model



Example of a very complicated forward model that maps from the initial conditions of the universe to (simulated) observed data from a telescope. Note that we can easily add new effects to the forward model. Clearly the computational challenge is enormous.

Example: Cosmology forward model

- Assuming the observational noise is Gaussian we can write an explicit likelihood of the form

$$\log \mathcal{L}(\mathbf{d}|\mathbf{s}, \Theta) = -\frac{1}{2}(\mathbf{f}(\mathbf{s}, \Theta) - \mathbf{d}^{obs})^T \mathbf{N}^{-1}(\mathbf{f}(\mathbf{s}, \Theta) - \mathbf{d}^{obs}) + \text{const.}$$

- To complete the posterior we need to add a prior for the parameters θ and $\mathbf{s}$ which we want to infer.
- Then we need to sample the posterior. Since $\mathbf{s}$ is usually very high dimensional, normal MCMC will not work.
- However there are sampling methods that work in very high dimensions, in particular **Hamiltonian Monte Carlo (HMC)** and versions of Variational Inference (VI). These **require the forward model to be differentiable**.
- Optimization in very high dimensions requires derivatives to find the minimum. For this reason **differentiable simulations** have become an important topic all over physics.
- More details about this approach in cosmology and references can be found in my cosmology lecture notes.

Examples: differentiable cosmology simulations

- <https://arxiv.org/abs/2010.11847> FlowPM: Distributed TensorFlow Implementation of the FastPM Cosmological N-body Solver
- <https://arxiv.org/abs/2211.09815> Differentiable Cosmological Simulation with Adjoint Method
- https://www.youtube.com/watch?v=Epsgh6vr0qs&ab_channel=ParticleMeshWithDerivatives
- <https://arxiv.org/abs/2002.00965> Bayesian de-lensing delight: sampling-based inference of the primordial CMB and gravitational lensing

Course logistics

- **Reading for this lecture:**
 - This lecture was based in part on <https://arxiv.org/abs/1911.01429>