

# Neural Simulation-Based Inference at the LHC

ML4Jets 2026, Vienna

Jay Sandesara



<https://indico.global/event/15240/>

**What is Simulation-Based Inference ?**

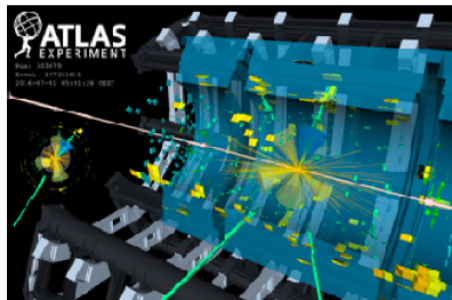
**What is Neural Simulation-Based Inference ?**

**How do we apply this to the LHC ?**

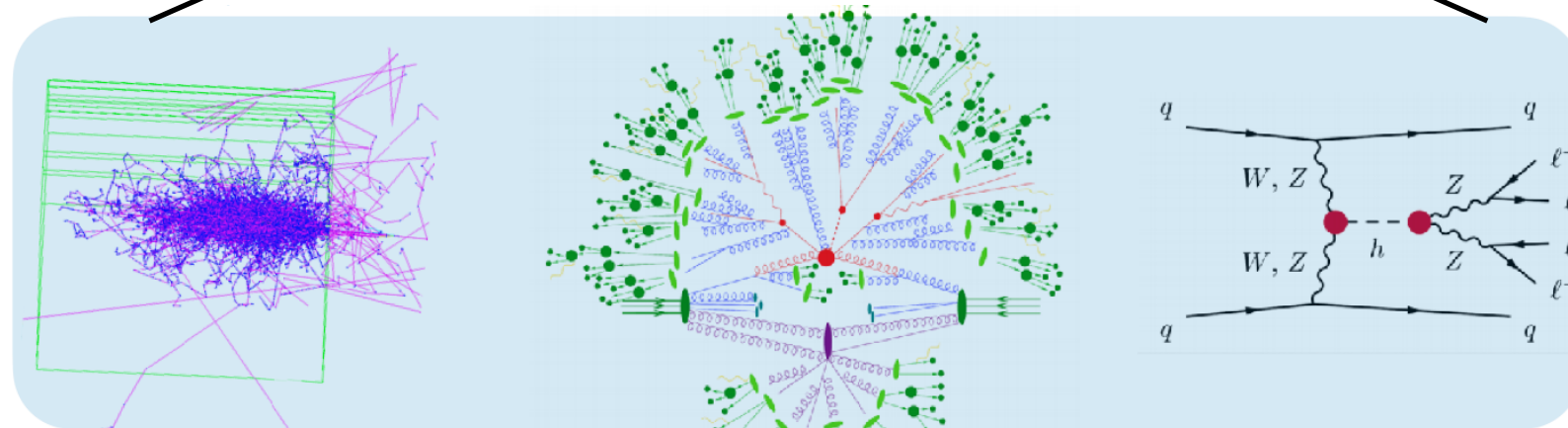
# Introduction

"Latent" or unobserved space

What we observe



What we measure



$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + i\bar{\psi}\not{D}\psi + h.c. + \bar{\psi}_i y_{ij} \psi_j \phi + h.c. + \frac{1}{2} \partial_\mu \phi^\dagger \partial^\mu \phi - V(\phi)$$

Data  $x$



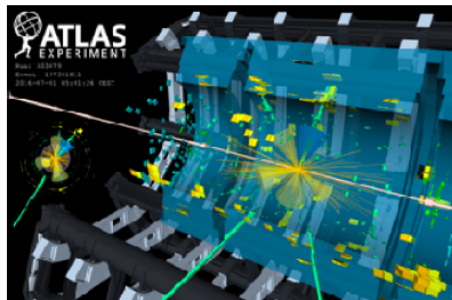
Theory model  
parameterized by  
some  $\mu$

$O(10^8)$ -dimensional  
 $x$ -space

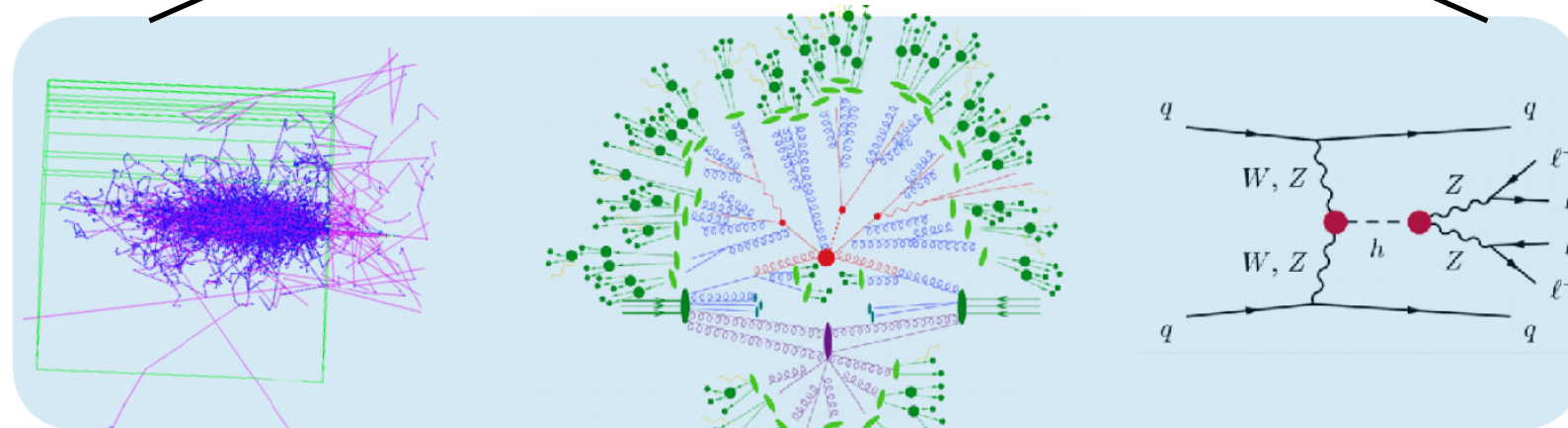
"High-dimensional"

# Introduction

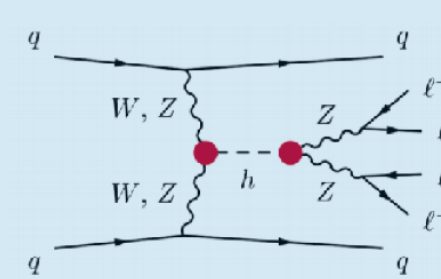
What we observe



"Latent" or unobserved space



What we measure



$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + i\bar{\psi}\not{D}\psi + h.c. + \bar{\psi}_i y_{ij} \psi_j \phi + h.c. + \frac{1}{2} D_\mu \phi^\dagger D^\mu \phi - V(\phi)$$

Data  $x$

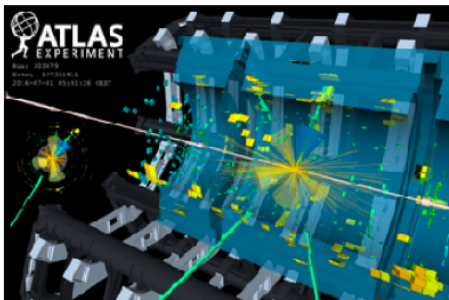


Theory model  
parameterized by  
some  $\mu$

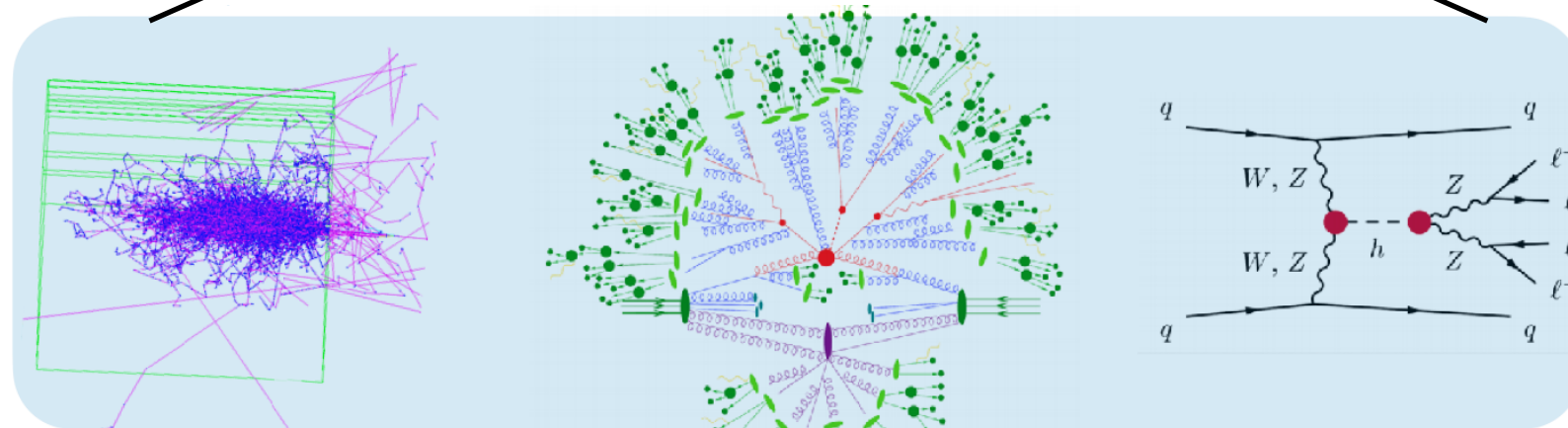
**Inference goal:** Likelihood  $L(\mu; x) = p(x | \mu)$

# Introduction

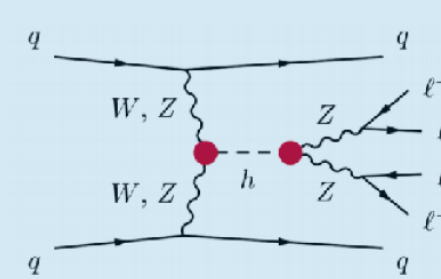
What we observe



"Latent" or unobserved space



What we measure



$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + i\bar{\psi}\not{D}\psi + h.c. + \bar{\psi}_i y_{ij} \psi_j \phi + h.c. + \frac{1}{2} \partial_\mu \phi^\dagger \partial^\mu \phi - V(\phi)$$

Data  $x$



Theory model  
parameterized by  
some  $\mu$

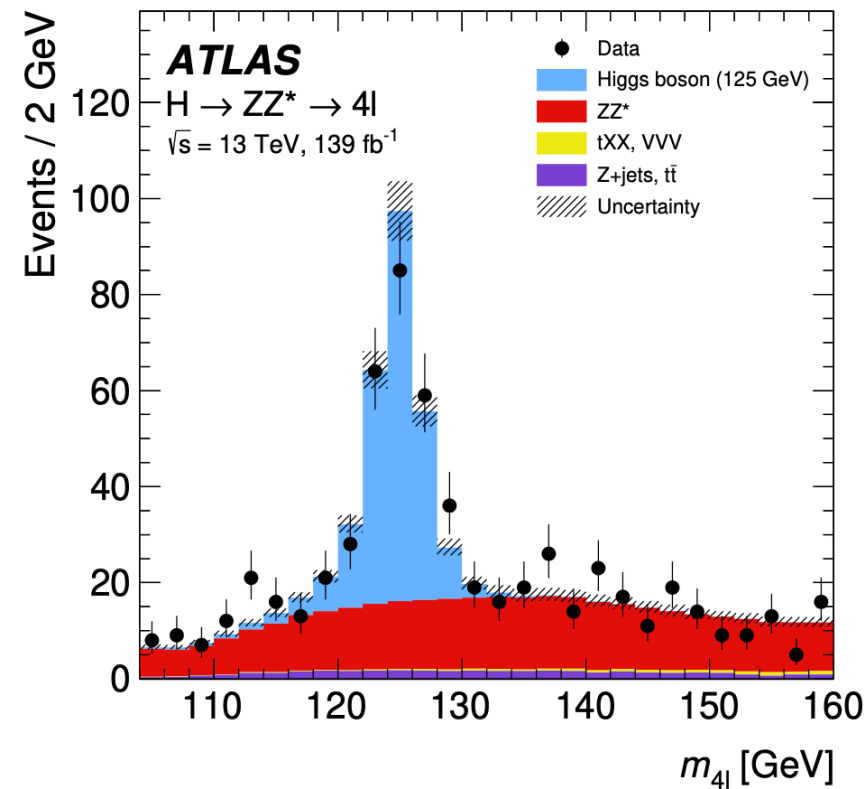
**Inference goal:** Likelihood  $L(\mu; x) = p(x | \mu)$

**"Likelihood-Free" Regime at the LHC**

Simulate events  $x \sim p(x | \mu)$  ✓

Compute likelihood  $x, \mu \rightarrow p(x | \mu)$  ✗

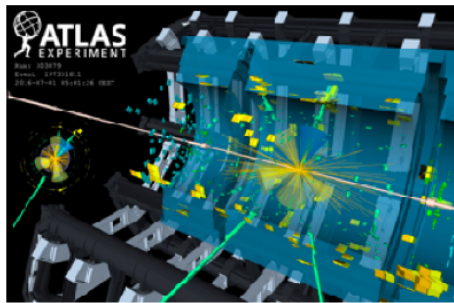
"Likelihood-Free" Regime → Simulation-Based Inference (SBI)



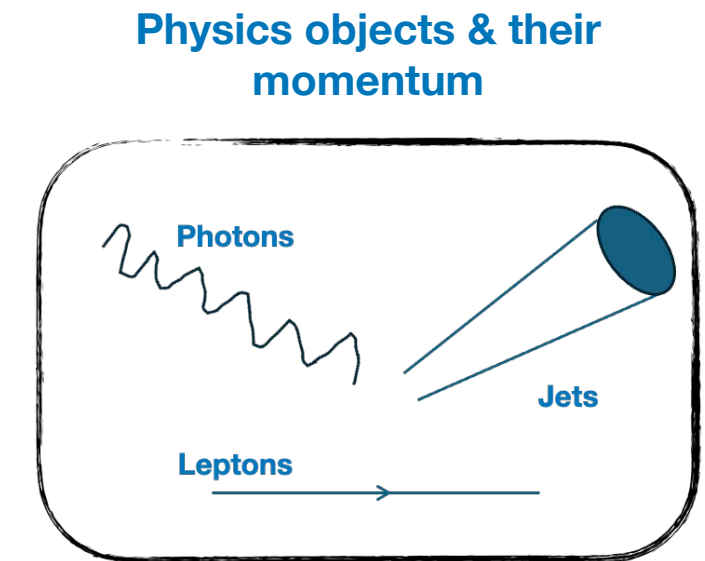
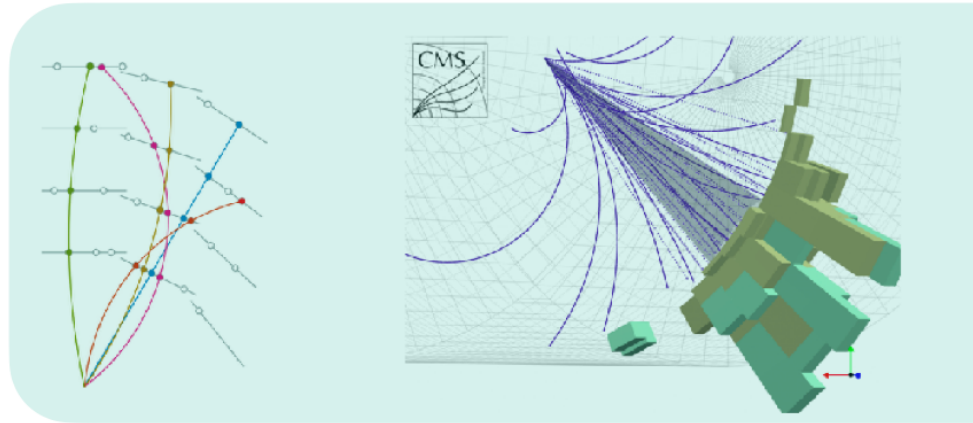
We have always been using SBI

# Traditional Inference

How do we fix the implicit model issue in traditional inference?



$O(10^8)$ -dimensional  
 $x$ -space

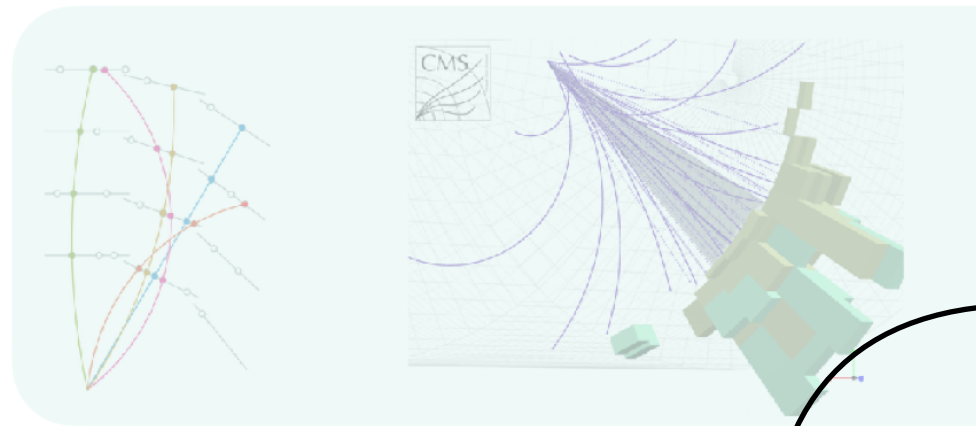
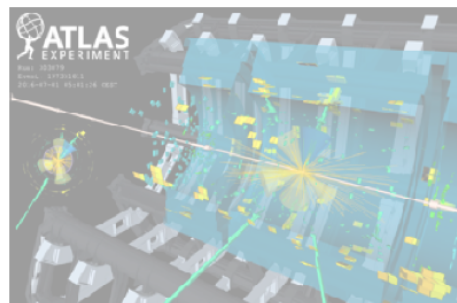


$O(10^2)$ -dimensional  
 $x$ -space

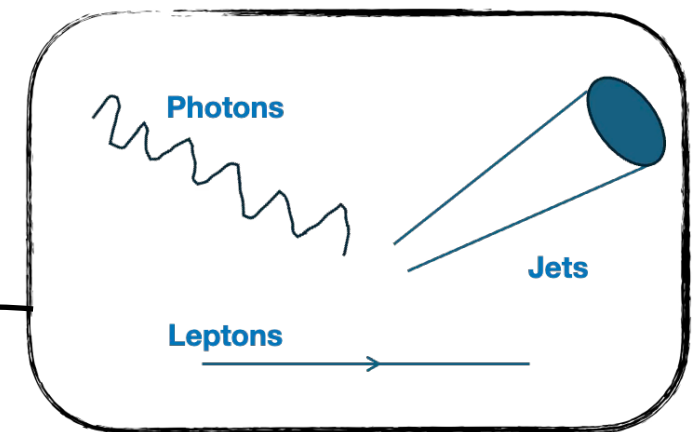
dimensionality reduction via reconstruction algorithms

# Traditional Inference

How do we fix the implicit model issue in traditional inference?



Physics objects & their momentum  $x \in \mathbb{R}^{O(10^2)}$



$$L(\mu, \nu) = \underbrace{\left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} \underbrace{p(x_e \mid \mu, \nu)}_{\text{Probability Density (intractable)}} \right]}_{\text{Extended Poisson term}} \cdot \underbrace{\prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)}_{\text{Constraint likelihood term}}$$

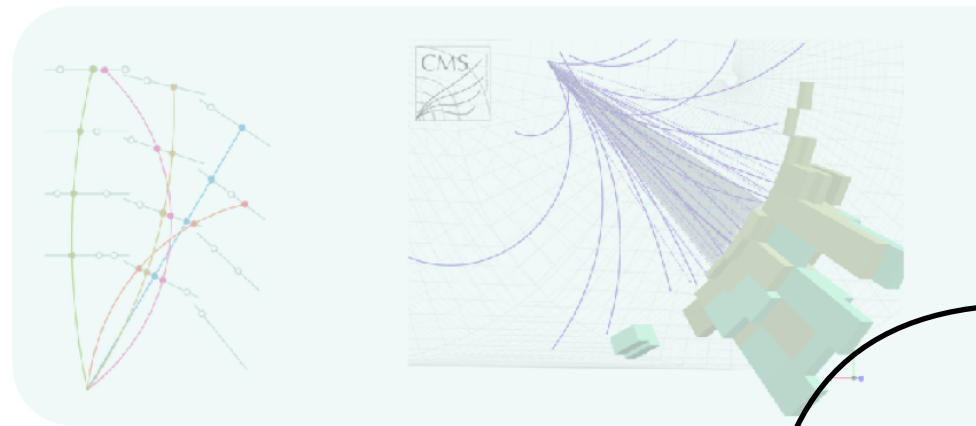
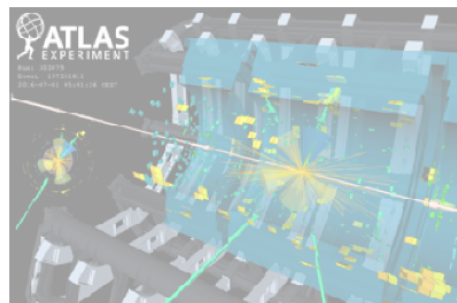
$n \rightarrow$  **observed # of events**  
 $\lambda \rightarrow$  **expected # of events**

$\mu \rightarrow$  **parameters of interest**  
 $\nu \rightarrow$  **nuisance parameters**

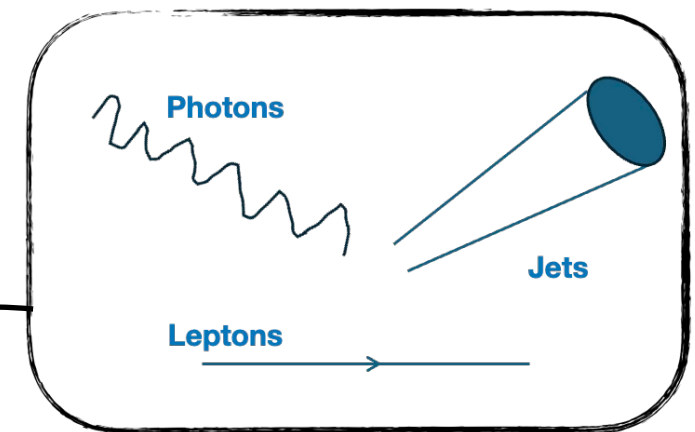
$a_m \rightarrow$  **global observable associated with  $m^{\text{th}}$  nuisance parameter**

# Traditional Inference

How do we fix the implicit model issue in traditional inference?



Physics objects & their momentum  $x \in \mathbb{R}^{O(10^2)}$



$$L(\mu, \nu) = \underbrace{\left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} \underbrace{p(x_e \mid \mu, \nu)}_{\text{Probability Density (intractable)}} \right]}_{\substack{\text{Extended Poisson term} \\ \checkmark}} \cdot \underbrace{\prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)}_{\substack{\text{Constraint likelihood term} \\ \checkmark}}$$

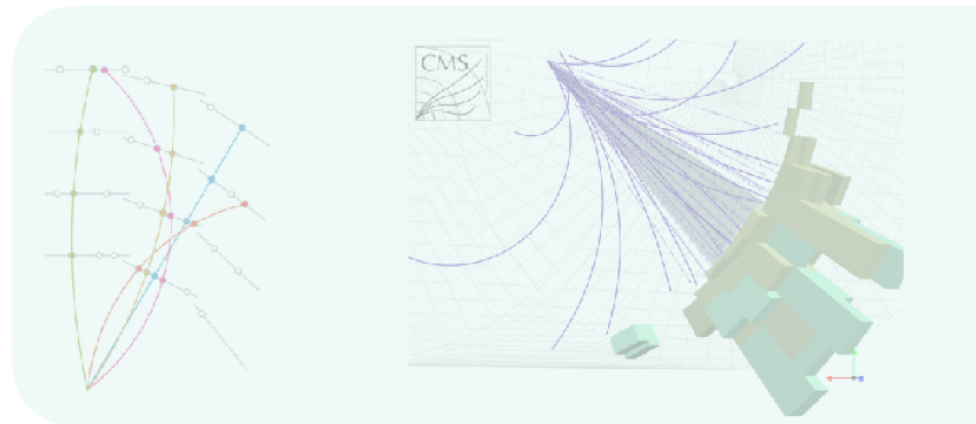
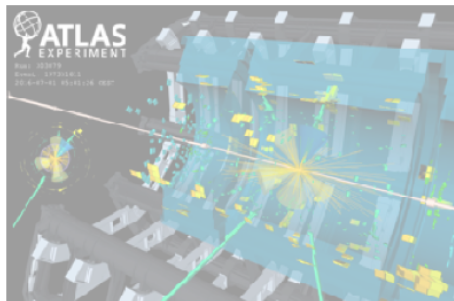
$n \rightarrow$  **observed # of events**  
 $\lambda \rightarrow$  **expected # of events**

$\mu \rightarrow$  **parameters of interest**  
 $\nu \rightarrow$  **nuisance parameters**

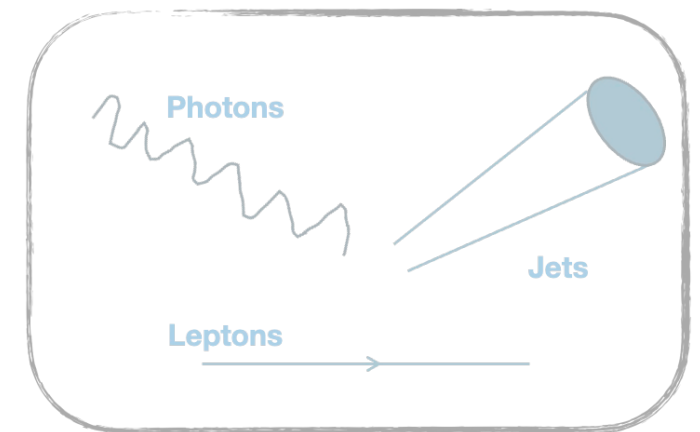
$a_m \rightarrow$  **global observable associated with  $m^{\text{th}}$  nuisance parameter**

# Traditional Inference

How do we fix the implicit model issue in traditional inference?



Physics objects & their momentum  $x \in \mathbb{R}^{O(10^2)}$



$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} p(x_e \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

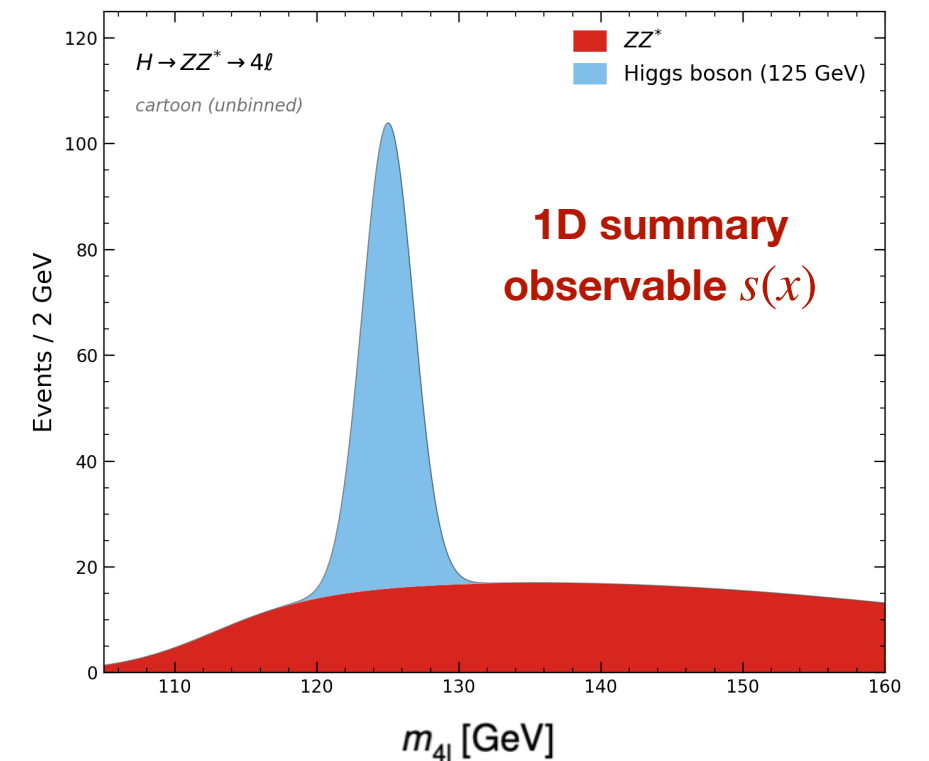
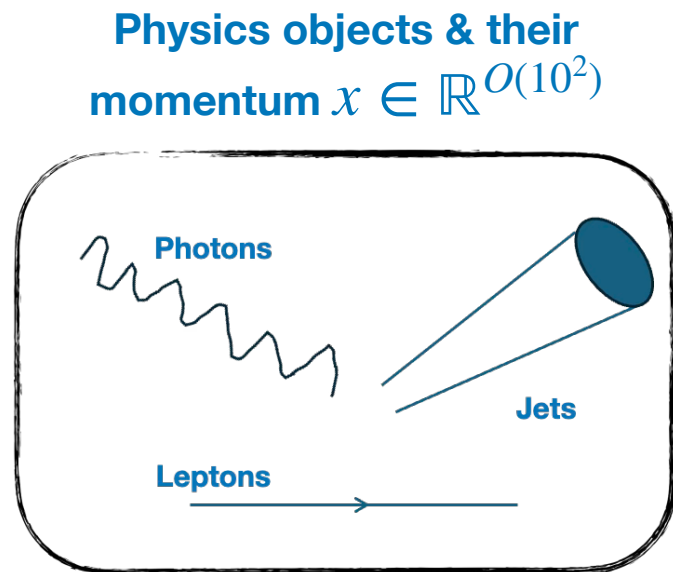
$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} p(s(x_e) \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Rewrite with a dimensionality reduction mapping function  $x \rightarrow s(x)$  e.g.  $m_{4\ell}(x)$

# Traditional Inference

$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{Events}} p(x_e \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{Events}} p(s(x_e) \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$



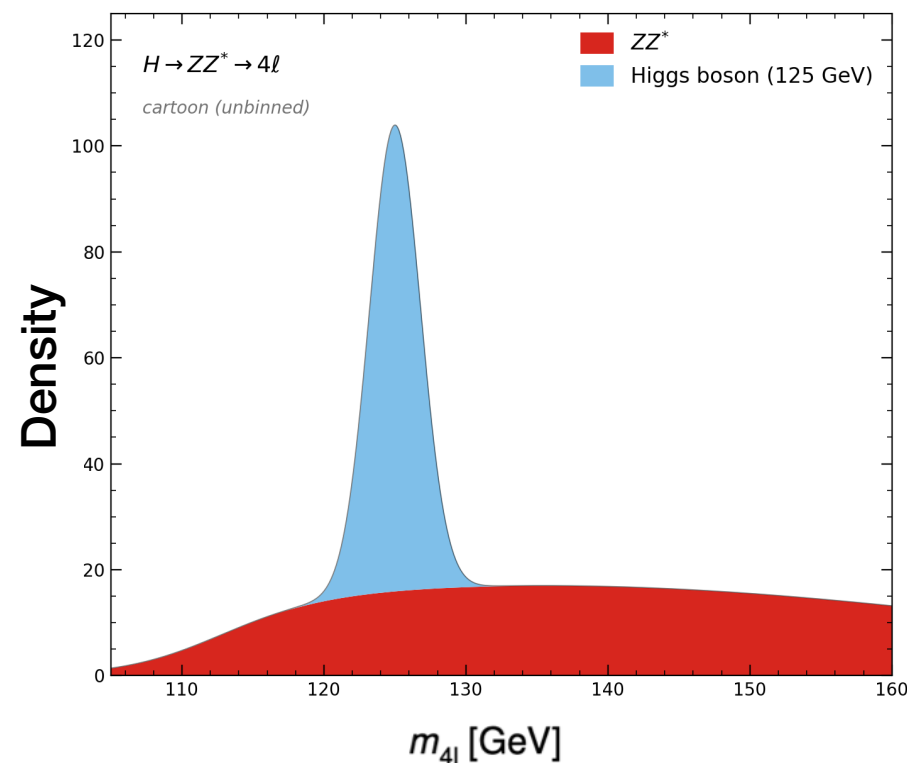
$$p(x_e \mid \mu, \nu) = \frac{1}{\sigma(\mu, \nu)} \frac{d\sigma}{dx}(x; \mu, \nu)$$

$$p(s(x_e) \mid \mu, \nu) = \frac{1}{\sigma(\mu, \nu)} \frac{d\sigma}{ds}(s; \mu, \nu)$$

# Traditional Inference

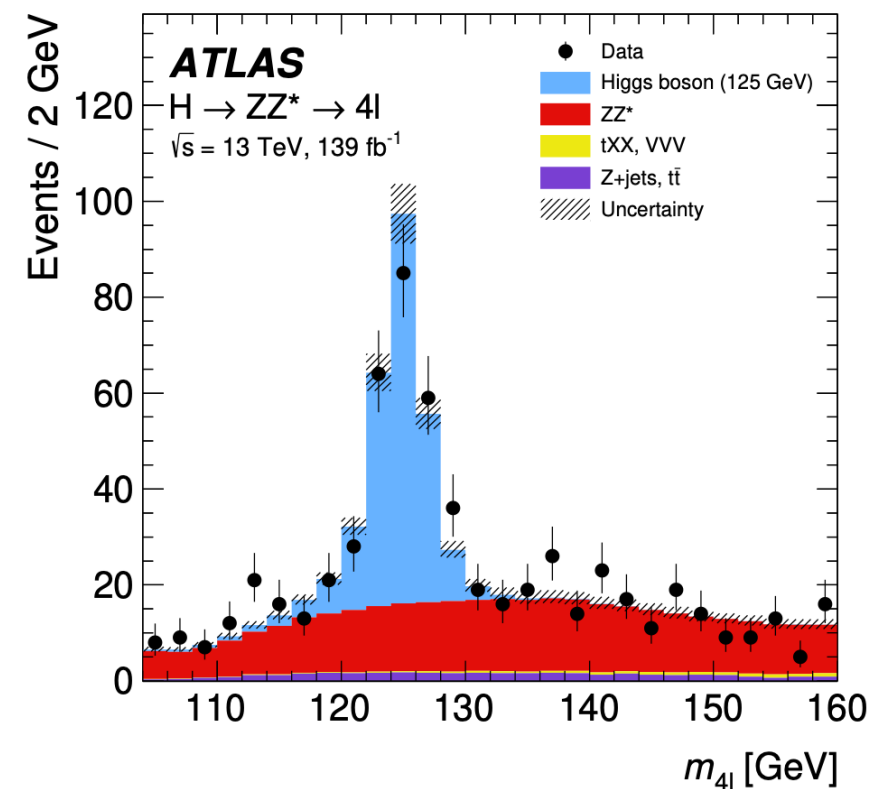
Continuous truth

$$p(s(x_e) | \mu, \nu) = \frac{1}{\sigma(\mu, \nu)} \frac{d\sigma}{ds}(s; \mu, \nu)$$



Binned model for density estimation

$$q_{\text{hist}}(s(x_e) | \mu, \nu) = \frac{1}{\lambda(\mu, \nu)} \frac{\lambda_b(\mu, \nu)}{\Delta_b}$$



The bin counts are modeled using simulation histograms

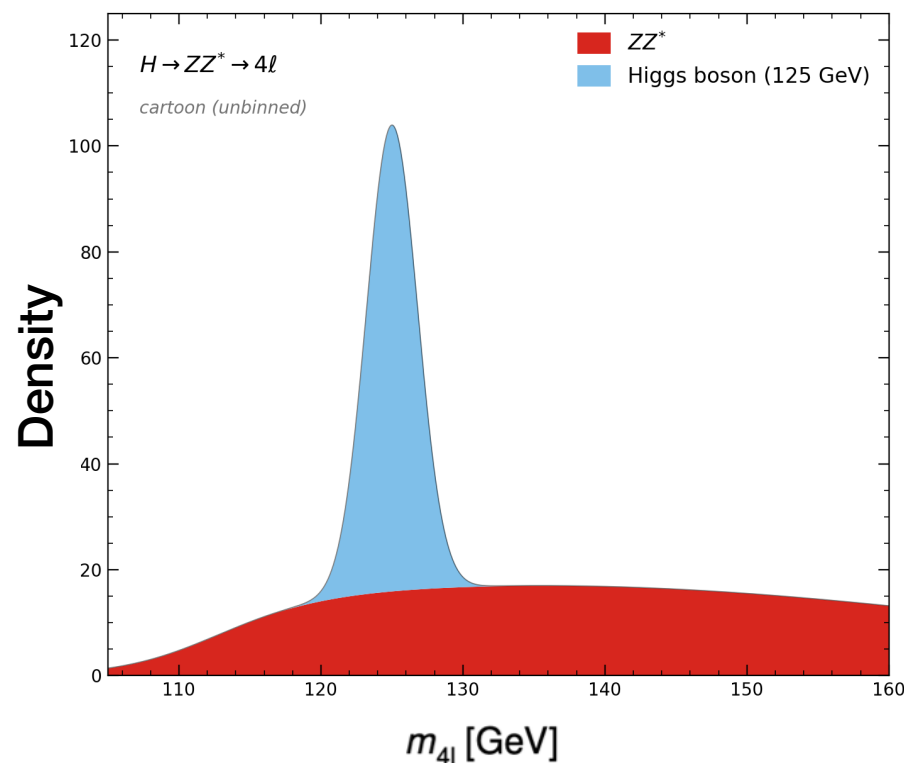
$\lambda_b \rightarrow$  expected yield in bin b

"Simulation-based" Inference

# Traditional Inference

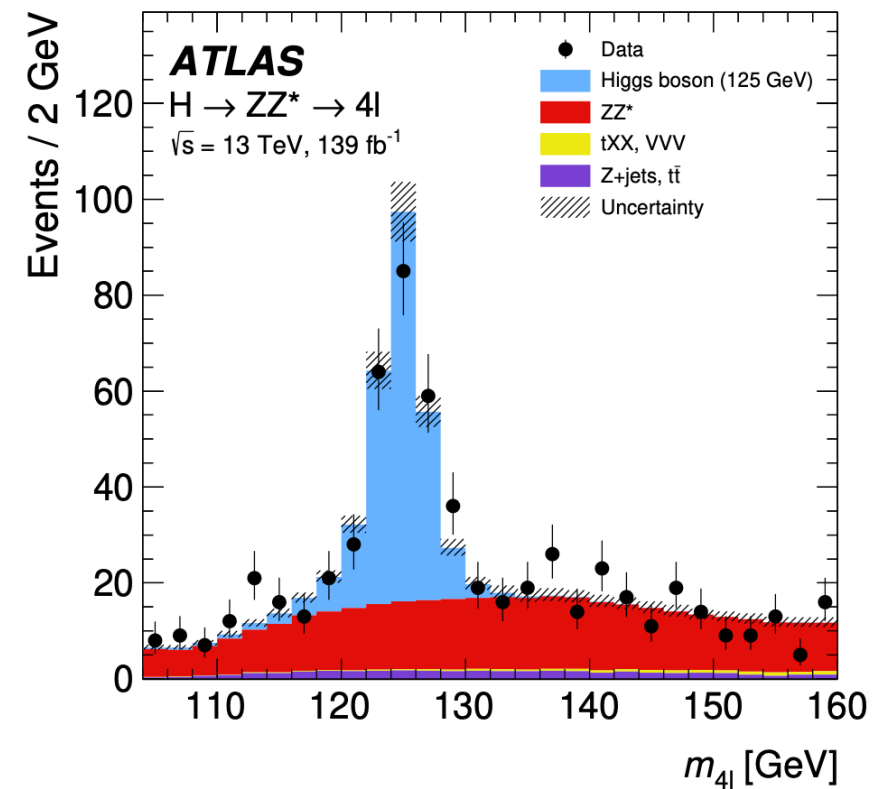
Continuous truth

$$p(s(x_e) | \mu, \nu) = \frac{1}{\sigma(\mu, \nu)} \frac{d\sigma}{ds}(s; \mu, \nu)$$



Binned model for density estimation

$$q_{\text{hist}}(s(x_e) | \mu, \nu) = \frac{1}{\lambda(\mu, \nu)} \frac{\lambda_b(\mu, \nu)}{\Delta_b}$$



The "reduced" model then uses histograms as density estimators:

$$L(\mu, \nu) = \left[ \text{Pois}(n | \lambda(\mu, \nu)) \prod_{e \in \text{events}} q_{\text{hist}}(s(x_e) | \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m | \nu_m)$$

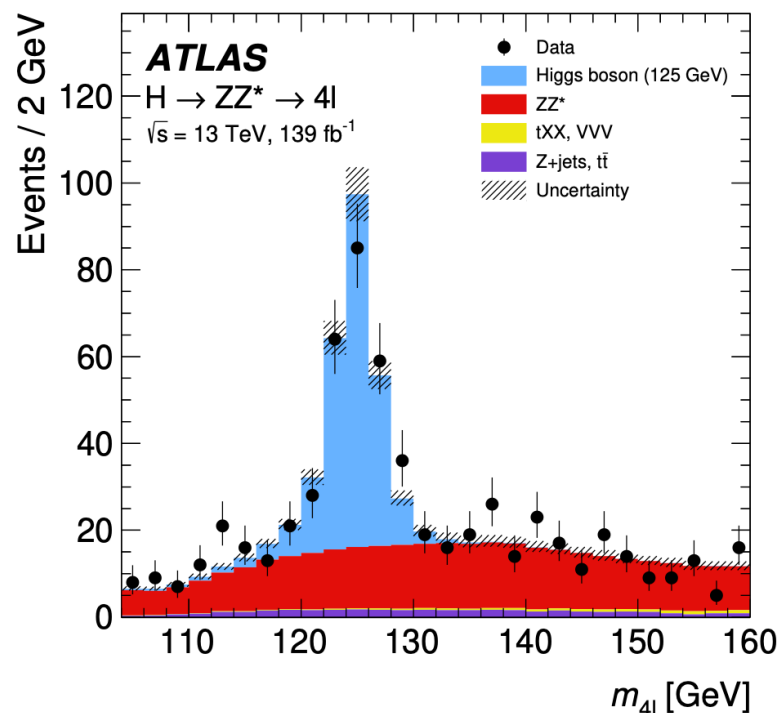
# Traditional Inference

The "reduced" model then uses histograms as density estimators:

$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} q_{\text{hist}}(s(x_e) \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Equivalent to the product of binned Poisson likelihoods

$$L(\mu, \nu) = \text{const} \cdot \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$



Per-bin  
observed  
event count

Per-bin  
expected  
event count

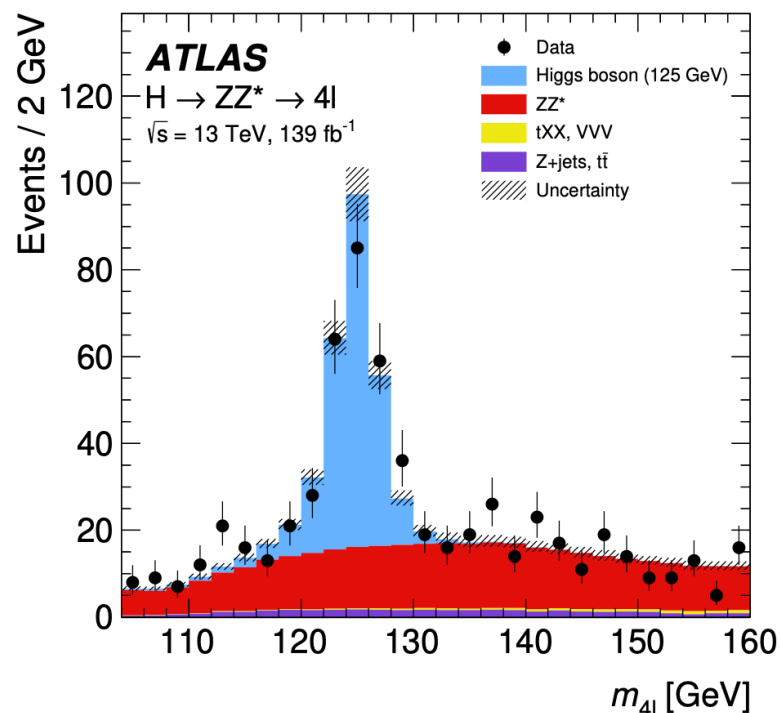
# Traditional Inference

The "reduced" model then uses histograms as density estimators:

$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} q_{\text{hist}}(s(x_e) \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Equivalent to the product of binned Poisson likelihoods

$$L(\mu, \nu) = \text{const} \cdot \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$



**Frequentist test statistic**

$$t_\mu = -2 \ln \frac{L(\mu, \hat{\nu})}{L(\hat{\mu}, \hat{\nu})}$$

**Allows us to perform parameter inference**

# Traditional Inference

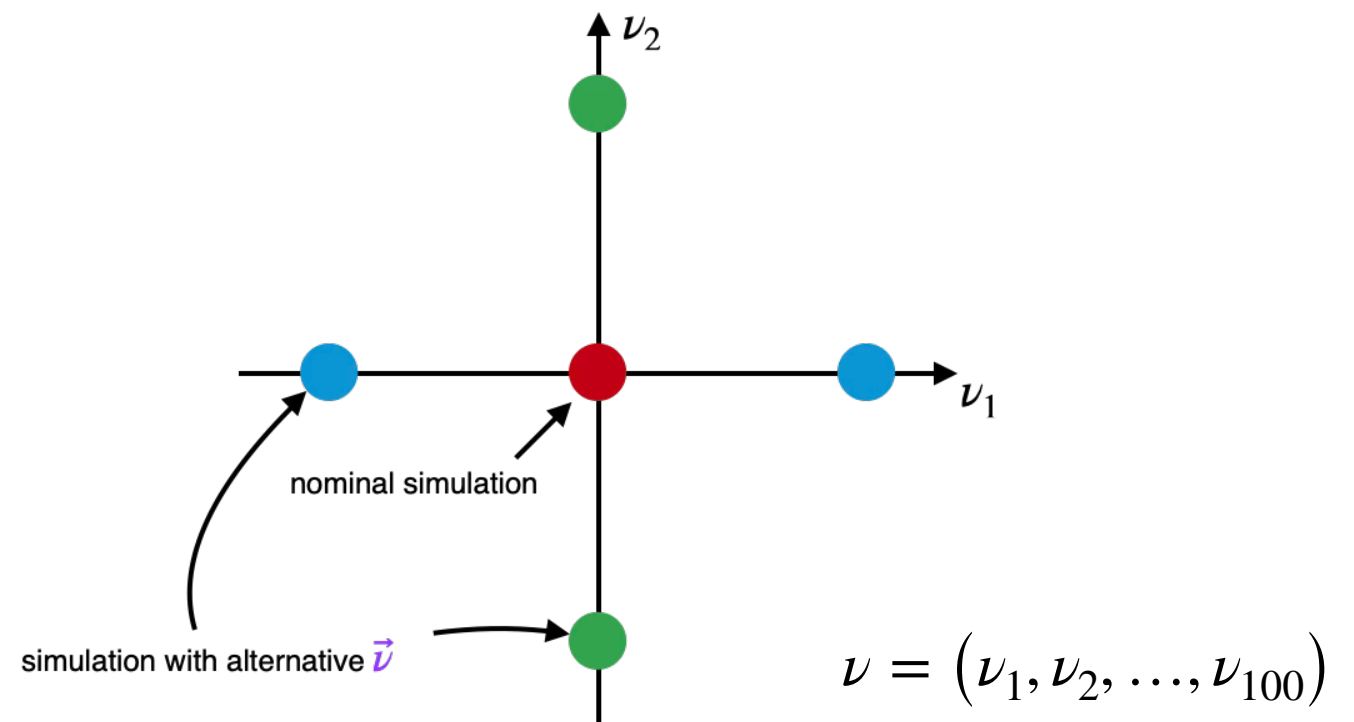
The expected rate needs to be modeled continuously as a function of parameters

$$L(\mu, \nu) = \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Need to build a continuous model as a function of  
**high-dimensional parameter space**

**What we have:**

expected event counts  $\lambda_b$  at  
discrete values of high-  
dimensional parameter space



# Traditional Inference

The expected rate needs to be modeled continuously as a function of parameters

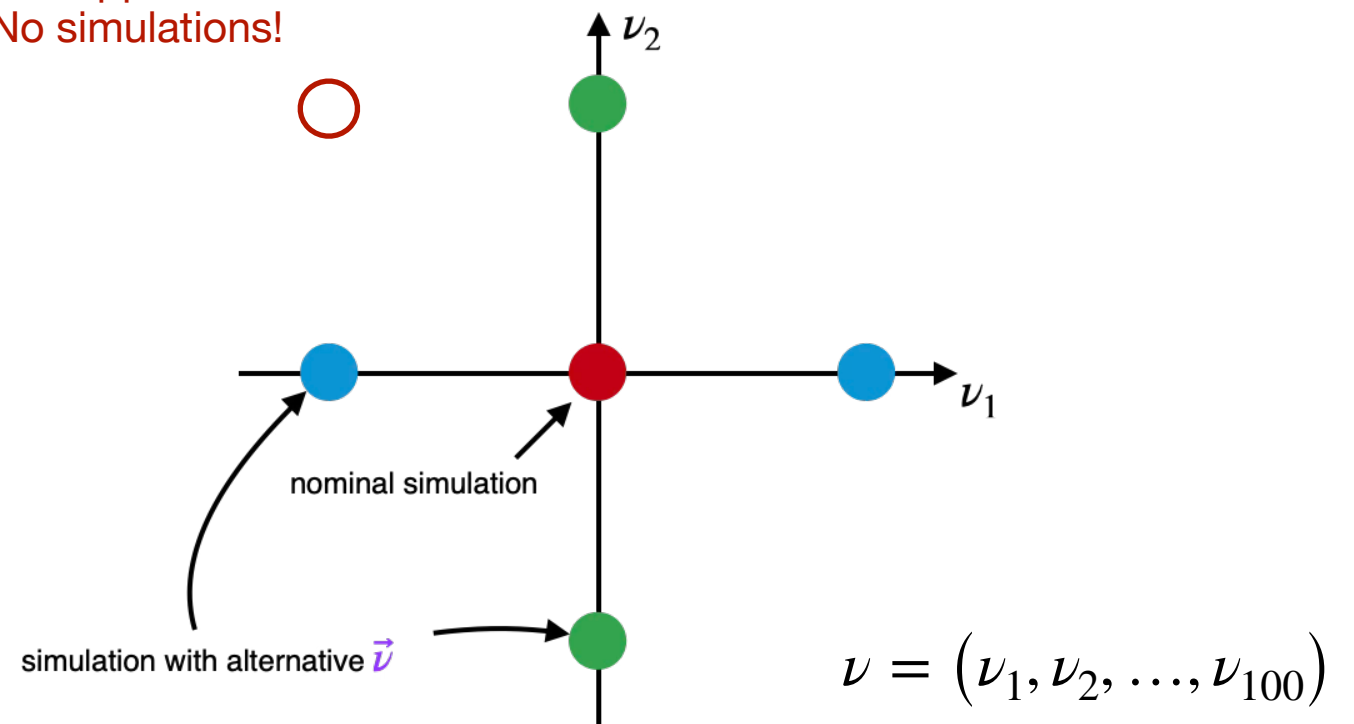
$$L(\mu, \nu) = \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Need to build a continuous model as a function of  
**high-dimensional parameter space**

**What we have:**

expected event counts  $\lambda_b$  at  
discrete values of high-  
dimensional parameter space

What happens here?  
No simulations!



# Traditional Inference

The expected rate needs to be modeled continuously as a function of parameters

$$L(\mu, \nu) = \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

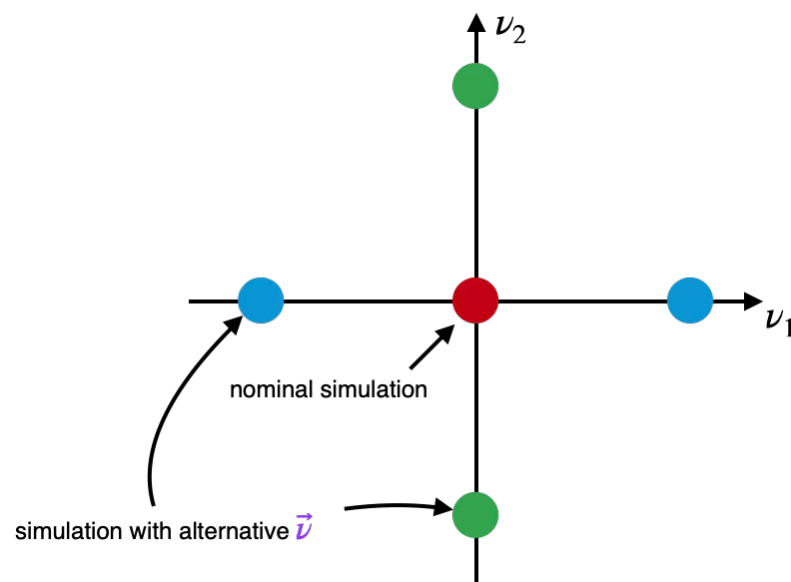
$$\sum_{s \in \text{samples}} \kappa_s(\mu) \cdot \prod_m \lambda_{s,b}(\nu_m)$$

"HistFactory" parametric model for continuous bin morphing:

factorised response and parametric interpolation

$$\nu = (\nu_1, \nu_2, \dots, \nu_{100})$$

$\kappa_s(\mu) \rightarrow$  known factorisable, analytic dependence



Assume impacts from each  $\nu_m$  are independent

$$\lambda_{s,b}(\nu) = \prod_m \lambda_{s,b}(\nu_m)$$

# Traditional Inference

The expected rate needs to be modeled continuously as a function of parameters

$$L(\mu, \nu) = \prod_{b \in \text{bins}} \text{Pois}(n_b \mid \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

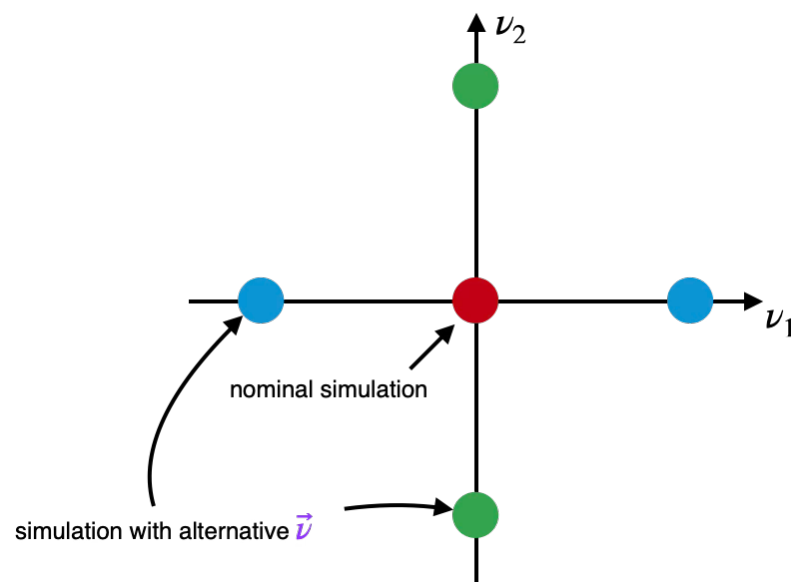
$$\sum_{s \in \text{samples}} \kappa_s(\mu) \cdot \prod_m \lambda_{s,b}(\nu_m)$$

"HistFactory" parametric model for continuous bin morphing:

factorised response and parametric interpolation

$$\nu = (\nu_1, \nu_2, \dots, \nu_{100})$$

$\kappa_s(\mu) \rightarrow$  known factorisable, analytic dependence



Assume impacts from each  $\nu_m$  are independent

$$\lambda_{s,b}(\nu) = \prod_m \lambda_{s,b}(\nu_m)$$

Remaining goal: estimate  $\lambda_{s,b}(\nu_m)$  from sparse simulations

$$\nu_m = -1, 0, 1$$

# Traditional Inference

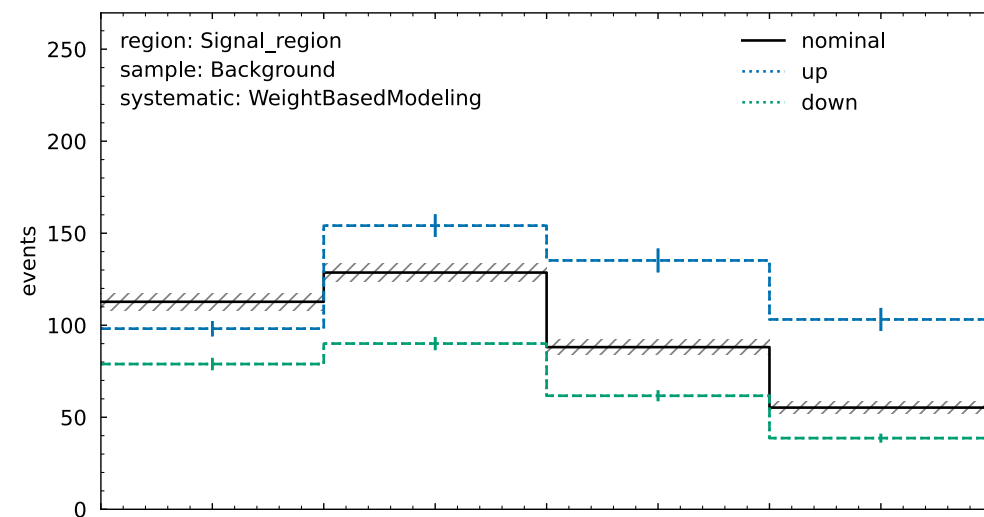
The expected rate needs to be modeled continuously as a function of parameters

$$L(\mu, \nu) = \prod_{b \in \text{bins}} \text{Pois}(n_b | \lambda_b(\mu, \nu)) \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m | \nu_m)$$

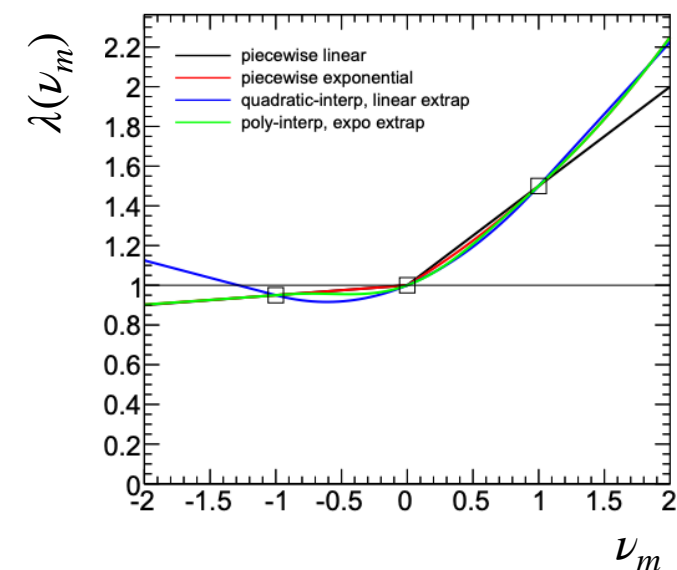
$$\sum_{s \in \text{samples}} \kappa_s(\mu) \cdot \prod_m \lambda_{s,b}(\nu_m)$$

"HistFactory" parametric model for continuous bin morphing:

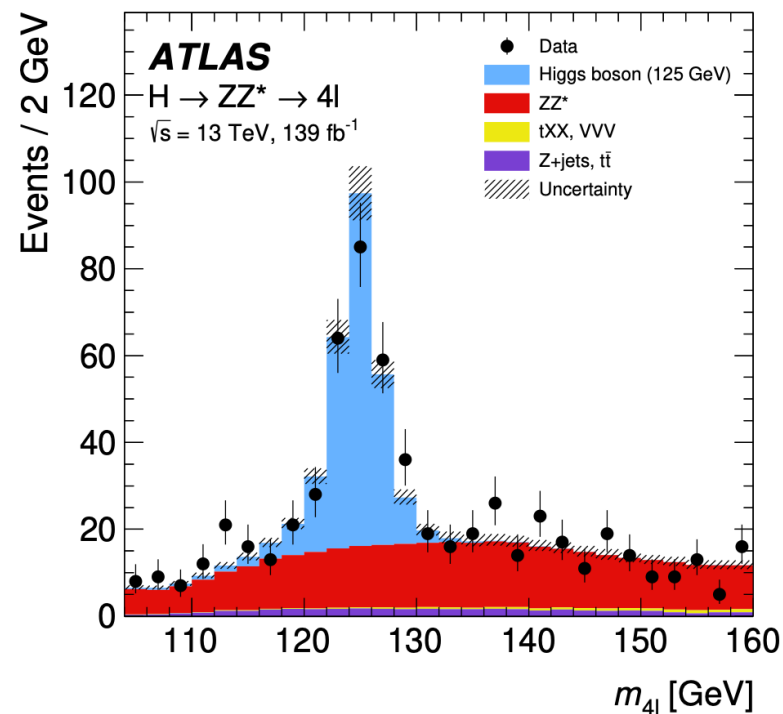
factorised response and parametric interpolation



toy example: templates for each  $\nu_m = -1, 0, +1$



parametric fit between templates per-bin and per  $\nu_m$



## Summary-based Histogram SBI

**Pros:** easy, cheap and asymptotically unbiased density estimation

**Cons:** lose sensitive information from high-dimensional data

## Why use neural SBI?

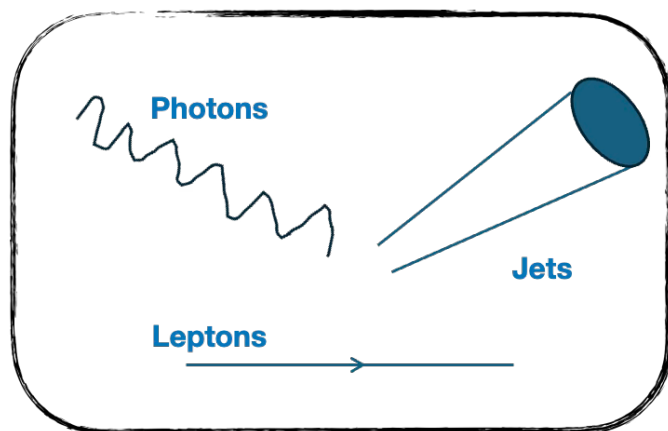
What if we can exploit the fully differential high-dimensional data to perform statistical inference?

# High-dimensional Inference

The full likelihood model in an ATLAS or CMS analysis looks like this:

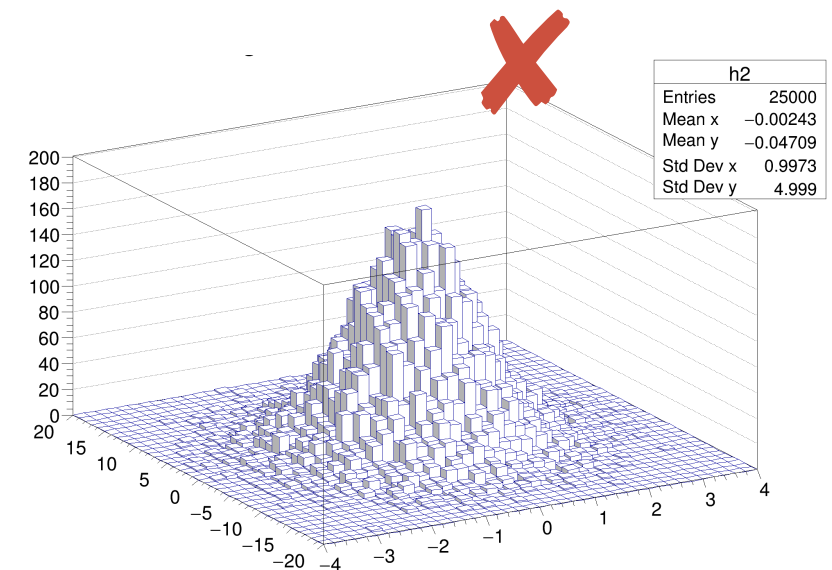
$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} p(x_e \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Estimation goal:  
per i.i.d. event density function



**Physics objects & their momentum**  
~10-100 dimensional vector

$x \sim p(x \mid \mu, \nu)$   
Simulated high-dimensional data



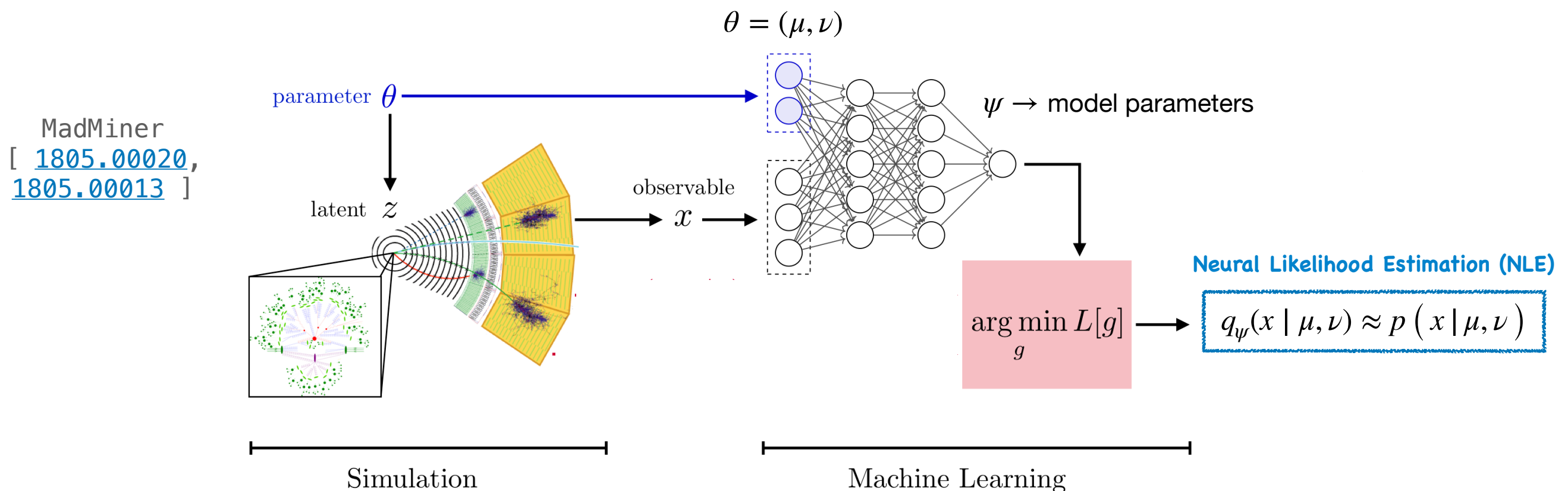
$$q_{\text{hist}}(x \mid \mu, \nu) \approx p(x \mid \mu, \nu)$$

Curse of dimensionality - need more simulations to model  $x$  than feasible

# Neural Simulation-Based Inference

The full likelihood model in an ATLAS or CMS analysis looks like this:

$$L(\mu, \nu) = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} p(x_e \mid \mu, \nu) \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

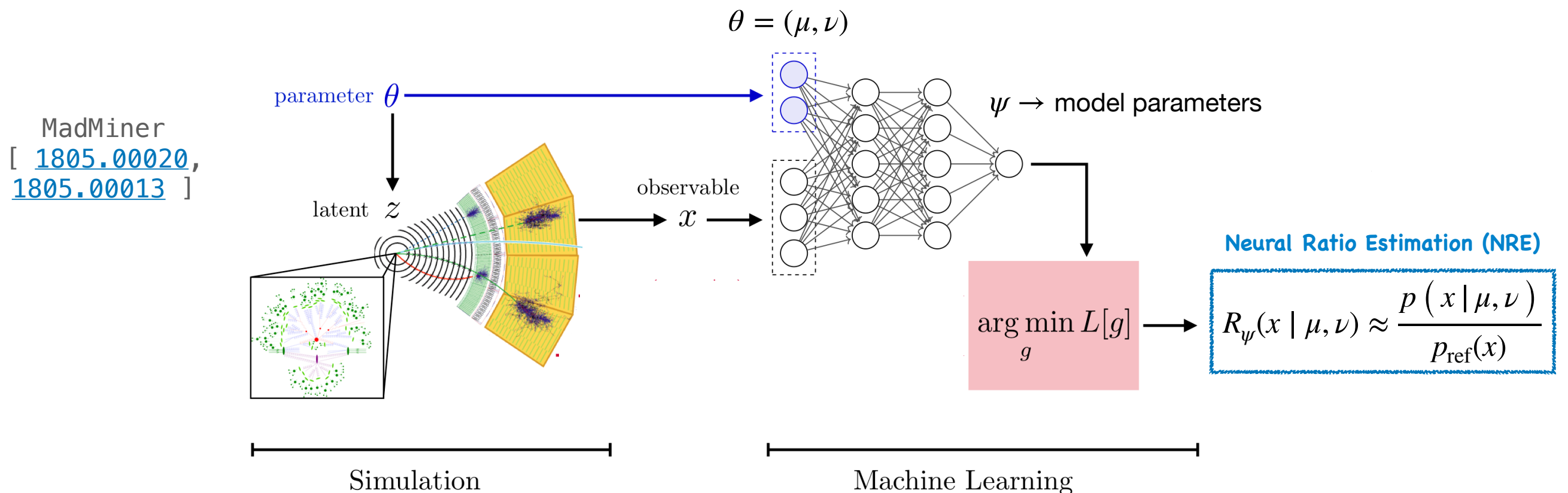


**neural SBI in a nutshell: use deep learning surrogates for the modeling of probability density**

# Neural Simulation-Based Inference

The full likelihood model in an ATLAS or CMS analysis looks like this:

$$\frac{L(\mu, \nu)}{\prod_e p_{\text{ref}}(x)} = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} \frac{p(x_e \mid \mu, \nu)}{p_{\text{ref}}(x)} \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$



**neural SBI in a nutshell: use deep learning surrogates for the modeling of probability density (ratios)**

# Neural Simulation-Based Inference

For a frequentist test, the parameter-independent  $p_{\text{ref}}$  can be arbitrarily chosen

$$\frac{L(\mu, \nu)}{\prod_e p_{\text{ref}}(x_e)} = \left[ \text{Pois}(n \mid \lambda(\mu, \nu)) \prod_{e \in \text{events}} \frac{p(x_e \mid \mu, \nu)}{p_{\text{ref}}(x_e)} \right] \cdot \prod_{m \in \text{systs}} p_{\text{constr.}}(a_m \mid \nu_m)$$

Density ratios can be estimated using  
well-trained classification models.  
Easier than density modelling

Test statistic remains unaffected

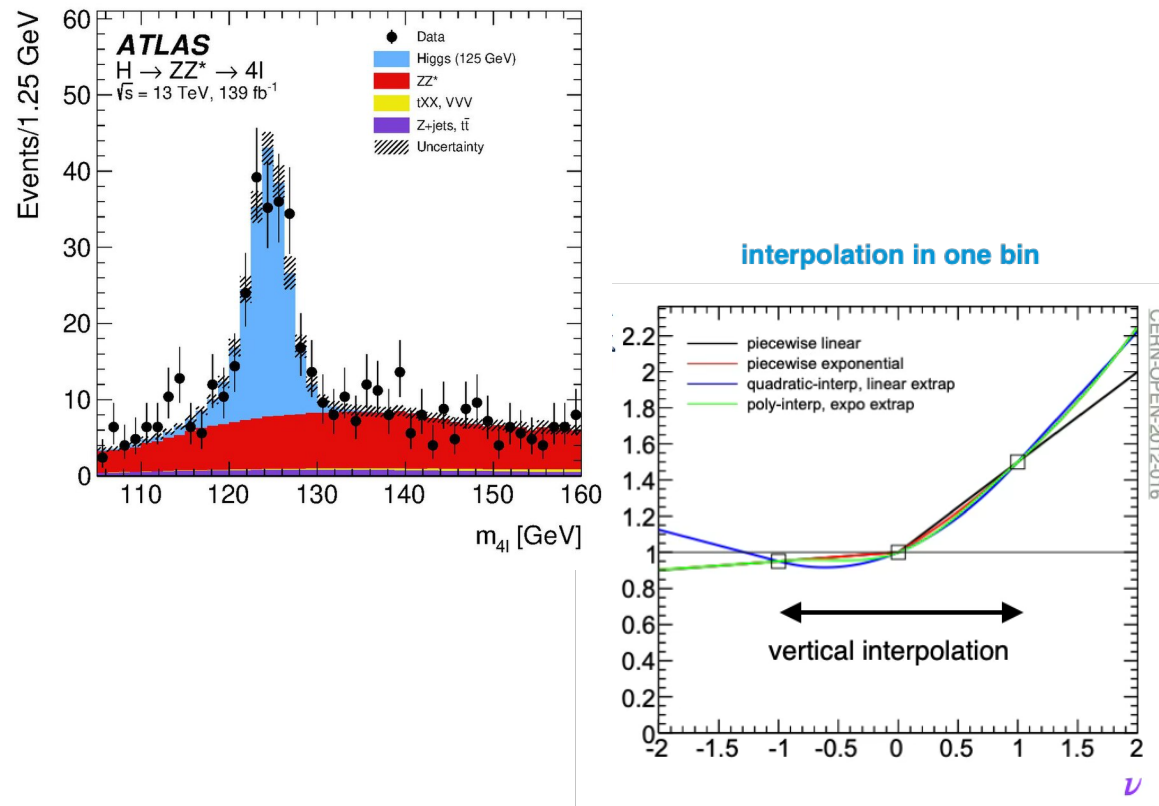
$$t_\mu = -\ln \frac{L(\mu, \hat{\nu}) / \cancel{p_{\text{ref}}(x_e)}}{L(\hat{\mu}, \hat{\nu}) / \cancel{p_{\text{ref}}(x_e)}}$$

$p_{\text{ref}}$  can be chosen arbitrarily  
eventually cancels out

The test retains all the desirable asymptotic properties that come with profile negative likelihood ratio test [[1007.1727](#)]

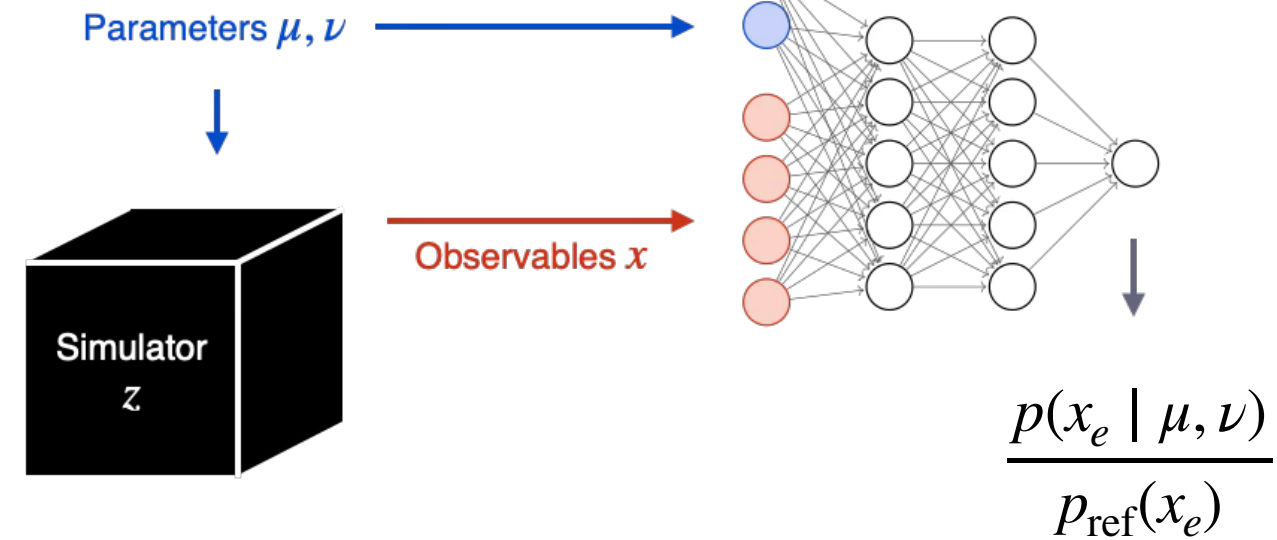
# SBI Landscape

## Traditional SBI



parametric = simple analytical dependency

## NSBI



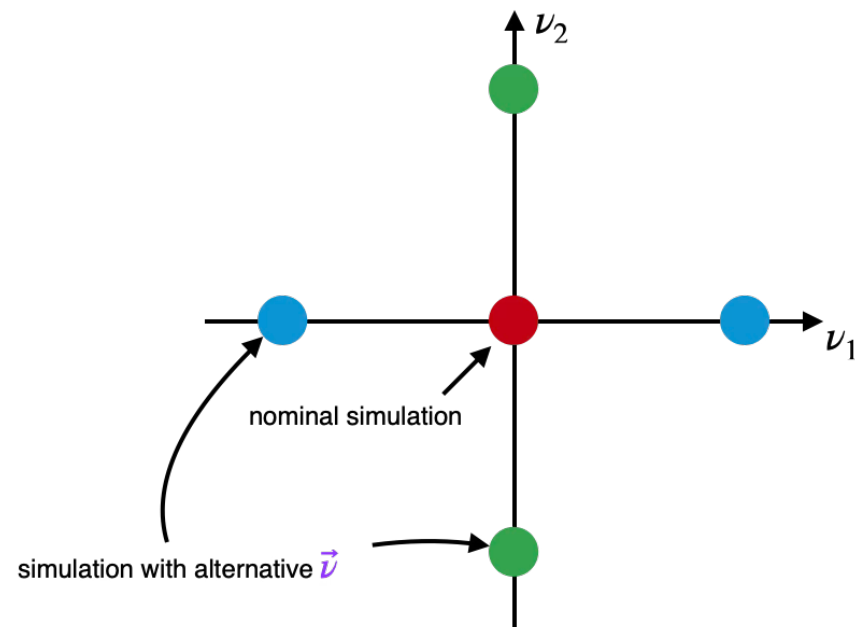
non-parametric = NNs, Gaussian Processes, etc.

|                               | $x$            | $\mu$          | $\nu$          |
|-------------------------------|----------------|----------------|----------------|
| Traditional Statistical Model | parametric     | parametric     | parametric     |
| Full Neural SBI               | non-parametric | non-parametric | non-parametric |

**Can we apply NSBI to LHC analysis?**

# Sparse simulations

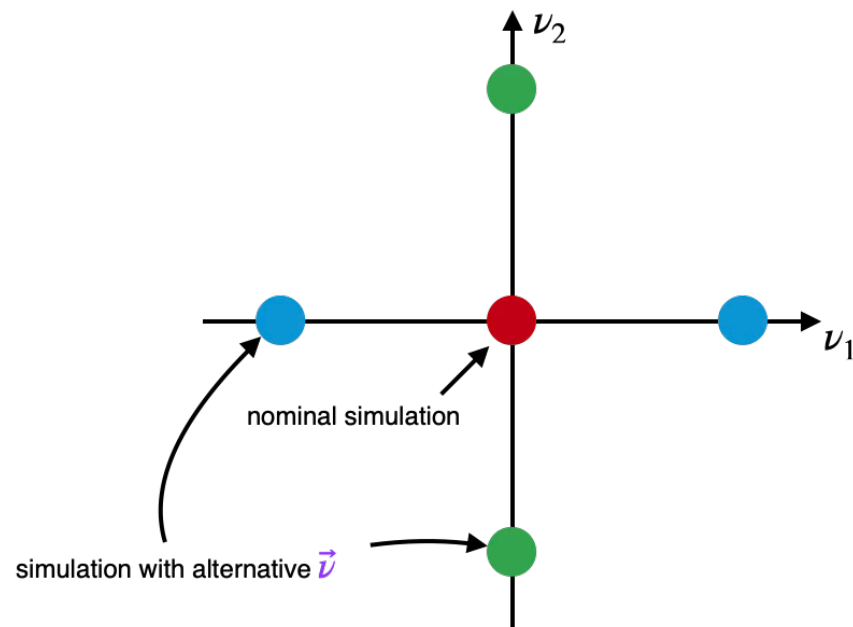
$$\nu = (\nu_1, \nu_2, \dots, \nu_{100})$$



**Challenge:** Only sparsely populated parameter space in 100–1000 dimensions for event samples in ATLAS and CMS.

# Sparse simulations

$$\nu = (\nu_1, \nu_2, \dots, \nu_{100})$$



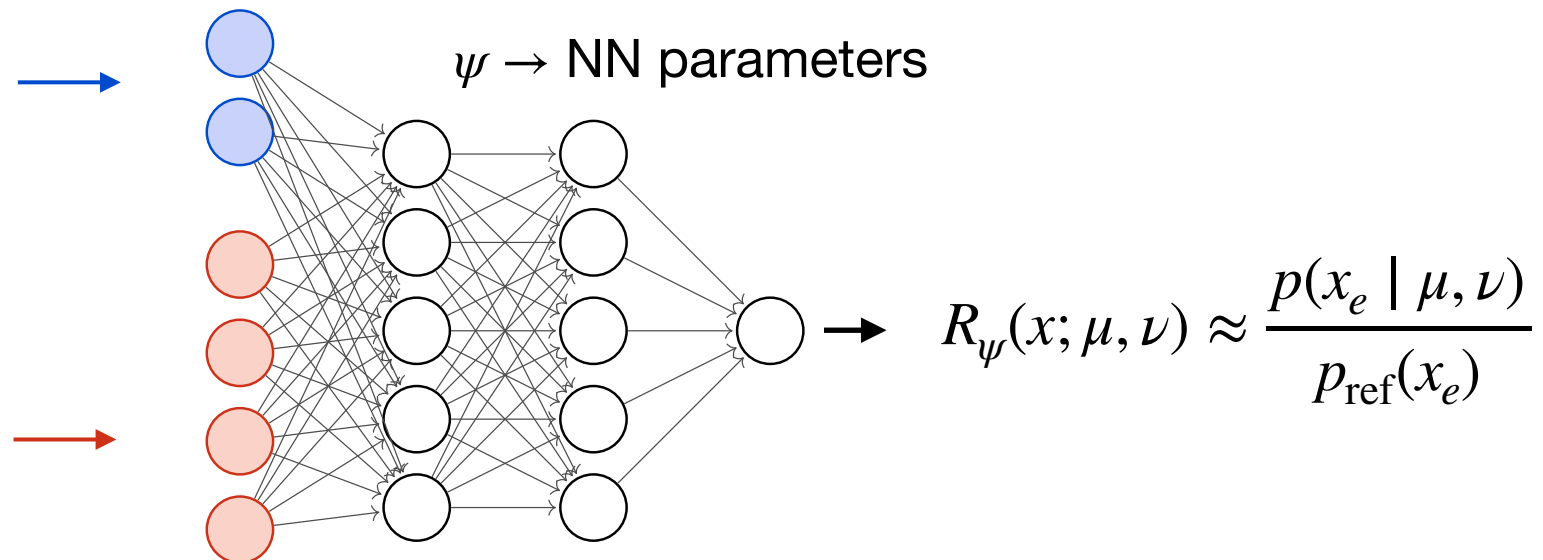
**Challenge:** Only sparsely populated parameter space in 100-1000 dimensions for event samples in ATLAS and CMS.

Neural Networks would not know how to model parameter space where there are no simulations

$$\mu, \nu \sim \pi(\mu, \nu)$$

$$x \sim p(x | \mu, \nu)$$

$$x \sim p_{\text{ref}}(x)$$



**Need to sample continuously and densely from some  $\pi(\mu, \nu)$  to model the ratios**

# Factorised Approach

The HistFactory model for per-event density ratio:

**Differential cross-section term  
(high-dimensional inference)**

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \frac{p_s(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)}$$

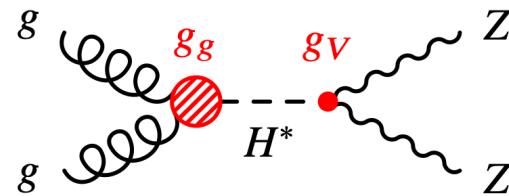
**Total yields  
(computed via HistFactory)**



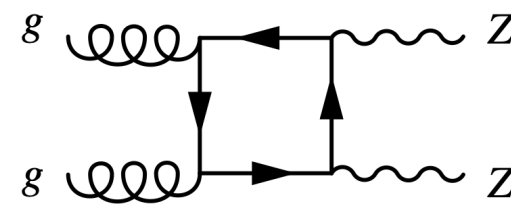
# Factorised Approach

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \left[ \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \frac{p_s(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} \right]$$



Signal to measure  $p_s = p_{\text{Sig}}$



Background with same final state  $p_s = p_{\text{Bkg}}$

E.g.: "signal strength" parameter

$$p(x | \mu) = \frac{1}{\mu \cdot \lambda_{\text{Sig}} + \lambda_{\text{Bkg}}} \left[ \mu \cdot \lambda_{\text{Sig}} \cdot p_{\text{Sig}}(x) + \lambda_{\text{Bkg}} \cdot p_{\text{Bkg}}(x) \right]$$

# Factorised Approach

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \frac{p_s(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)}$$



$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \frac{p_s(x_e | \nu)}{p_s(x_e)} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

**Known  
dependencies  
( $\mu$ -dependence)**

**Systematic  
uncertainty  
reweight factor  
( $x, \nu$ -dependence)**

**Nominal  
Density ratio  
( $x$ -dependence)**

**Goal: "factorised" targets to tackle independently**

# Factorised Approach

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \frac{p_s(x_e | \nu)}{p_s(x_e)} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

Systematic uncertainty  
reweight factor



Known  
dependencies



Nominal density ratio  
(parameter-independent)

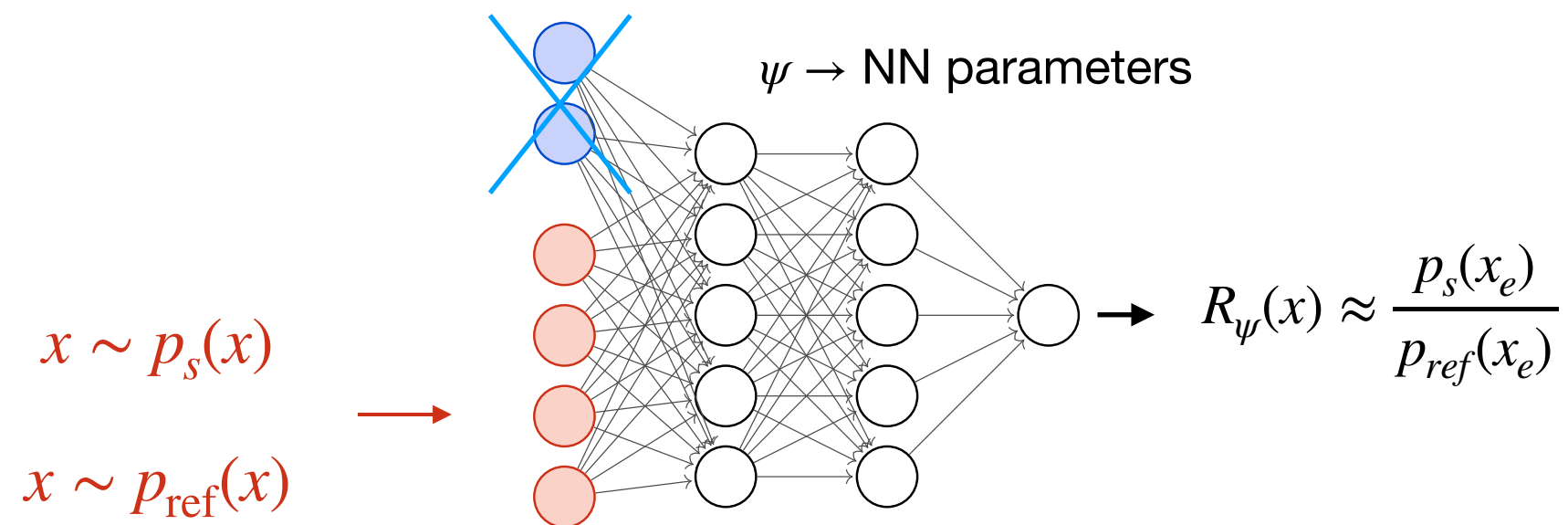
# Factorised Approach

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \frac{p_s(x_e | \nu)}{p_s(x_e)} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

↑  
Nominal density ratio  
(parameter-independent)

Parameter-independent model  
no dense parameter sampling needed!



# Semi-parametric SBI

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \frac{p_s(x_e | \nu)}{p_s(x_e)} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

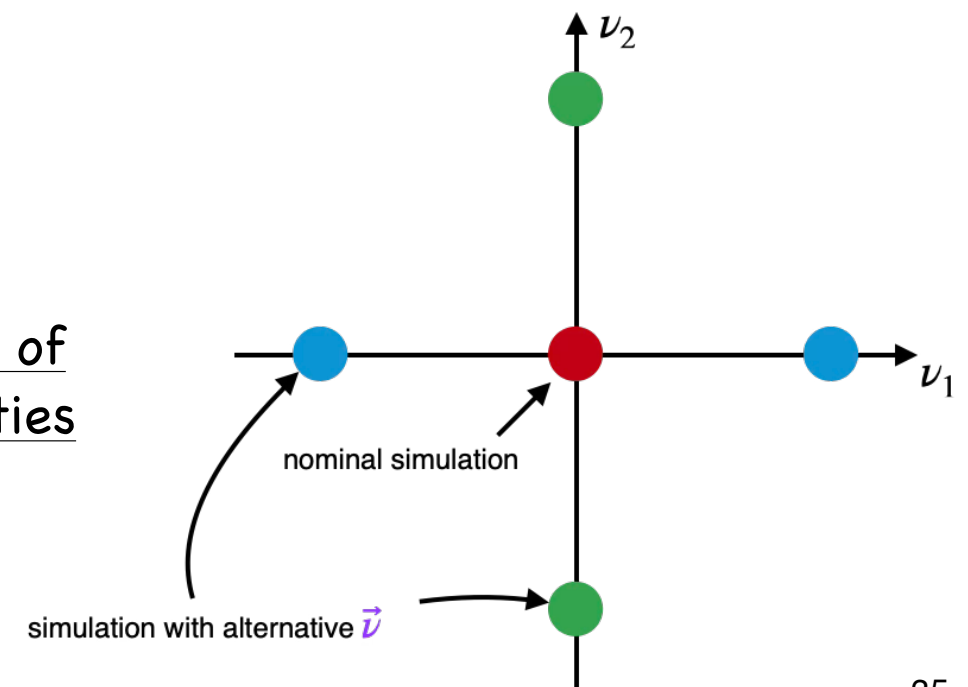
$$\approx \prod_m^{N_{\text{syst}}} \left[ \frac{p_s(x_e | \nu_m)}{p_s(x_e)} \right]$$

$\nu = (\nu_1, \nu_2, \dots, \nu_{N_{\text{syst}}})$

**HistFactory-style model**

**Assumption 1:**

impact from  
independent sources of  
systematic uncertainties  
factorise



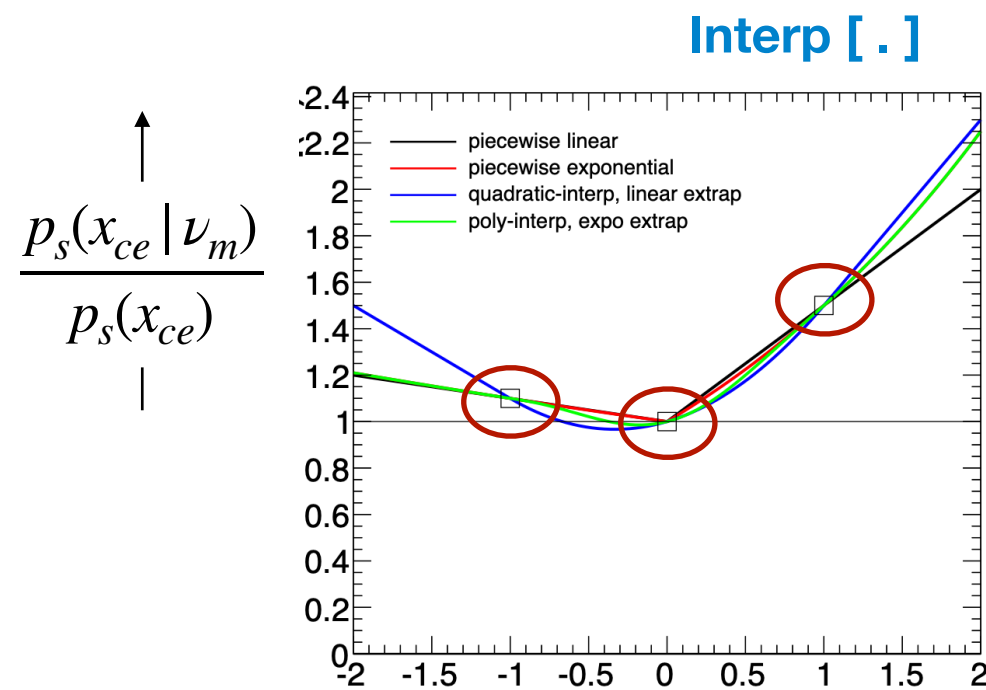
# Semi-parametric SBI

The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \frac{p_s(x_e | \nu)}{p_s(x_e)} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

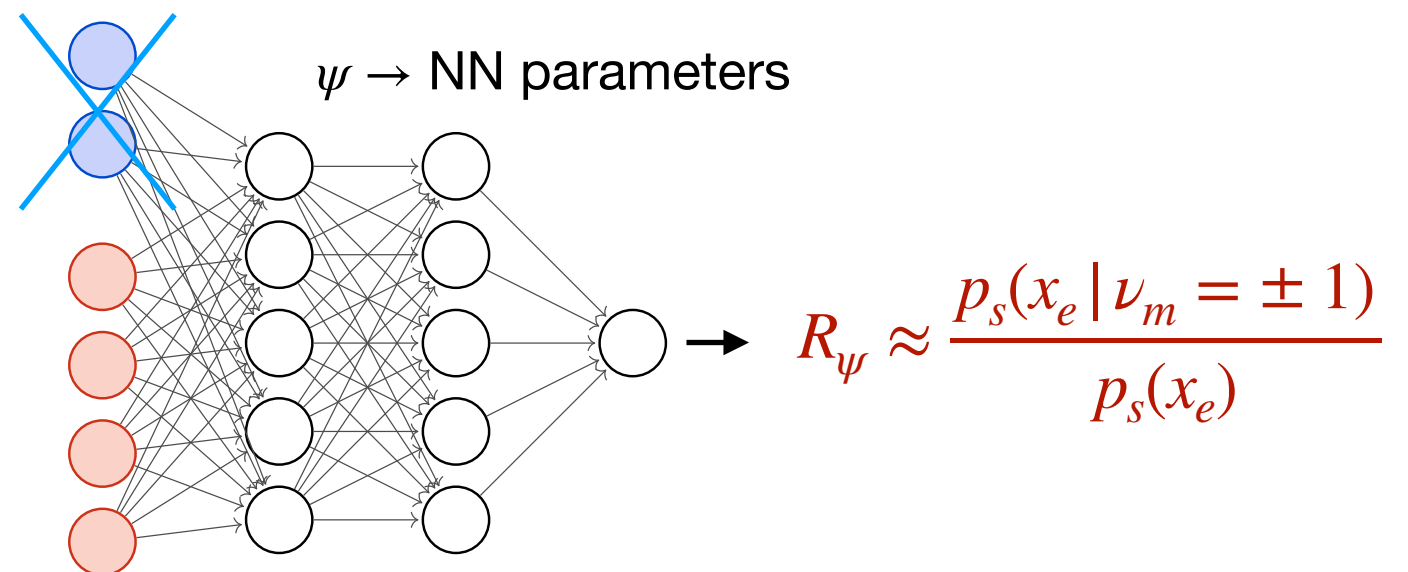
$$\approx \prod_m^{N_{\text{syst}}} \text{Interp} \left[ \frac{p_s(x_e | \nu_m = +1)}{p_s(x_e)}, 1, \frac{p_s(x_e | \nu_m = -1)}{p_s(x_e)} \right]$$

HistFactory-style model



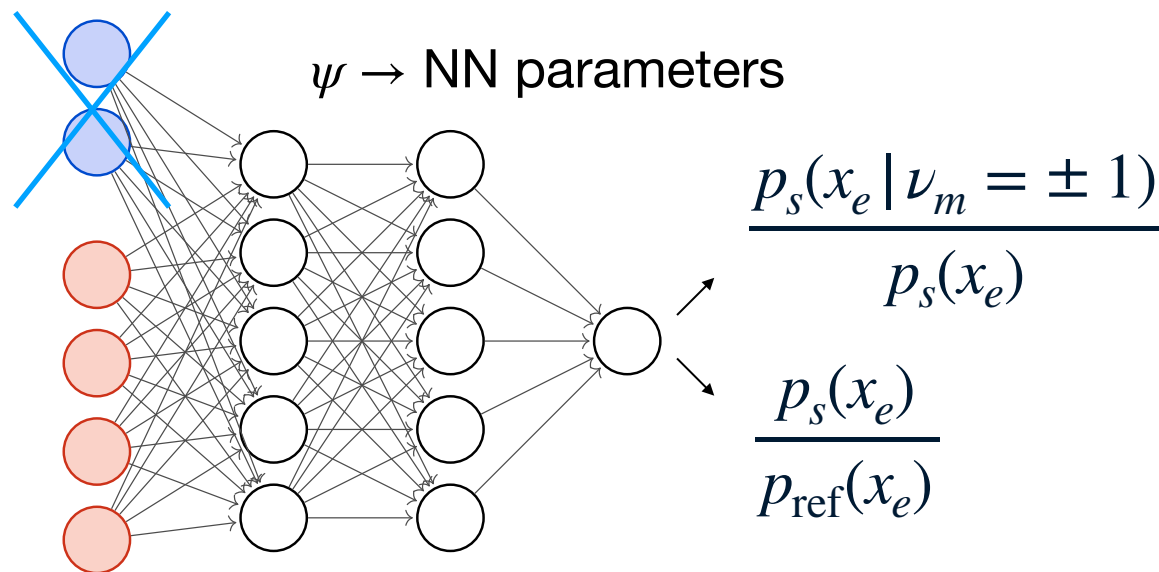
Assumption 2:

semi-parametric model  
for parameterised  
density ratios



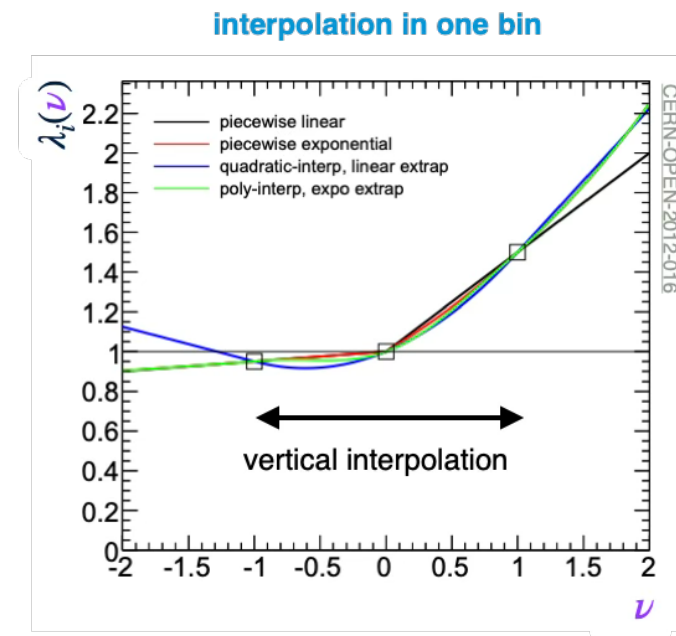
Train two NNs per nuisance-parameter

# SBI Landscape



non-parametric models for modeling  $x$

+

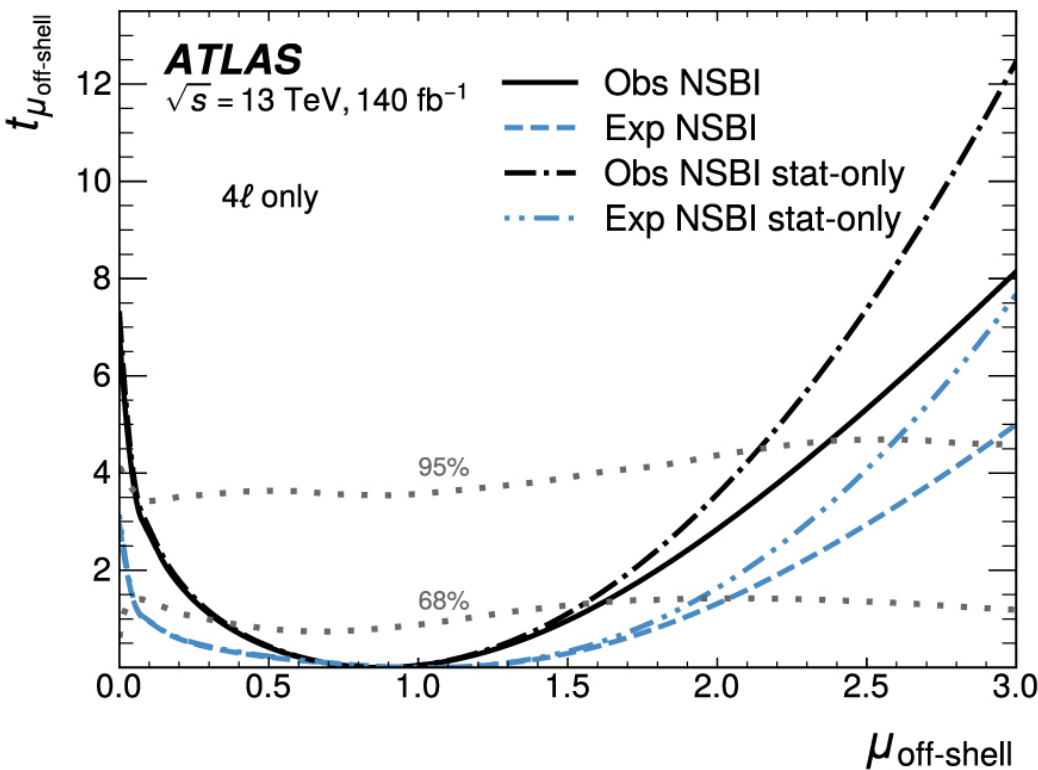
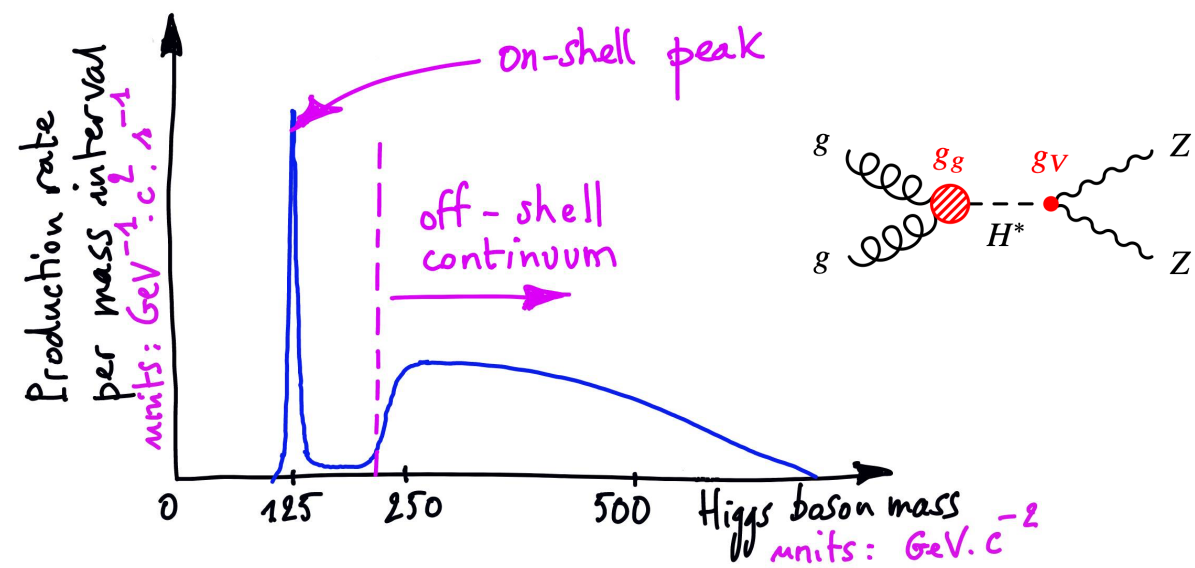


parametric ansatz for  $\mu, \nu$ -dependence

Ongoing effort mapping out the SBI landscape more generally (Cranmer & Sandesara)

|                                    | $x$            | $\mu$          | $\nu$          |
|------------------------------------|----------------|----------------|----------------|
| Traditional Statistical Model      | parametric     | parametric     | parametric     |
| Full Neural SBI                    | non-parametric | non-parametric | non-parametric |
| Semi-parametric SBI (POI&Nuisance) | non-parametric | parametric     | parametric     |

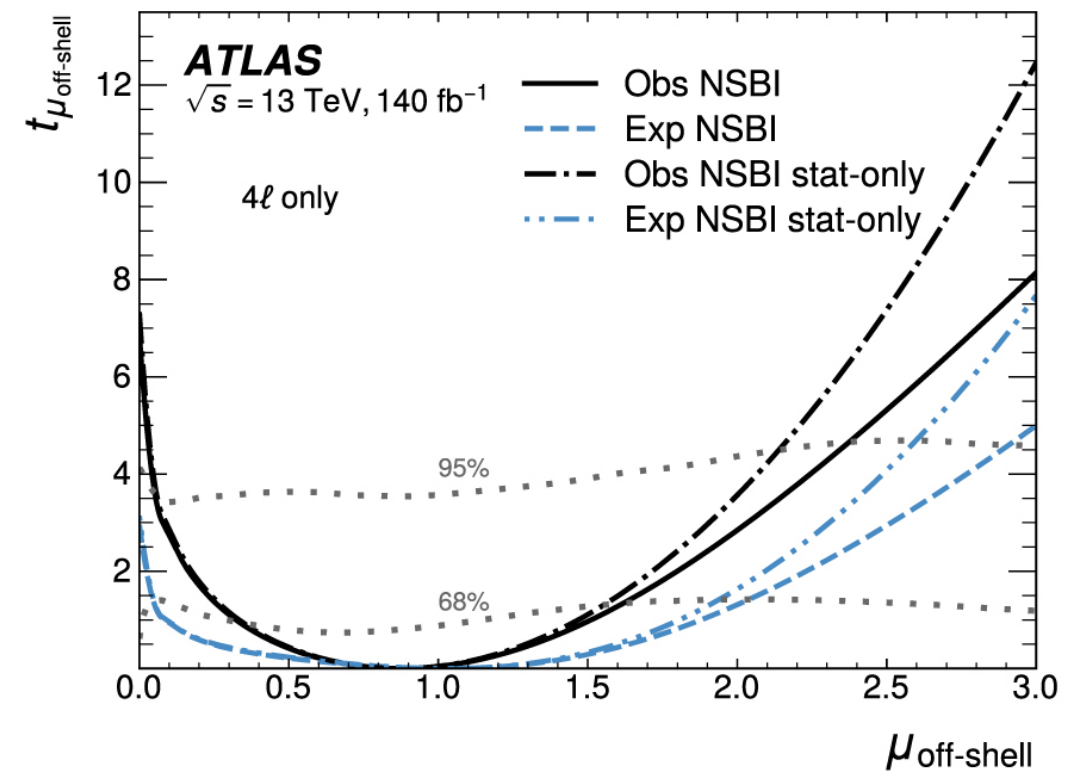
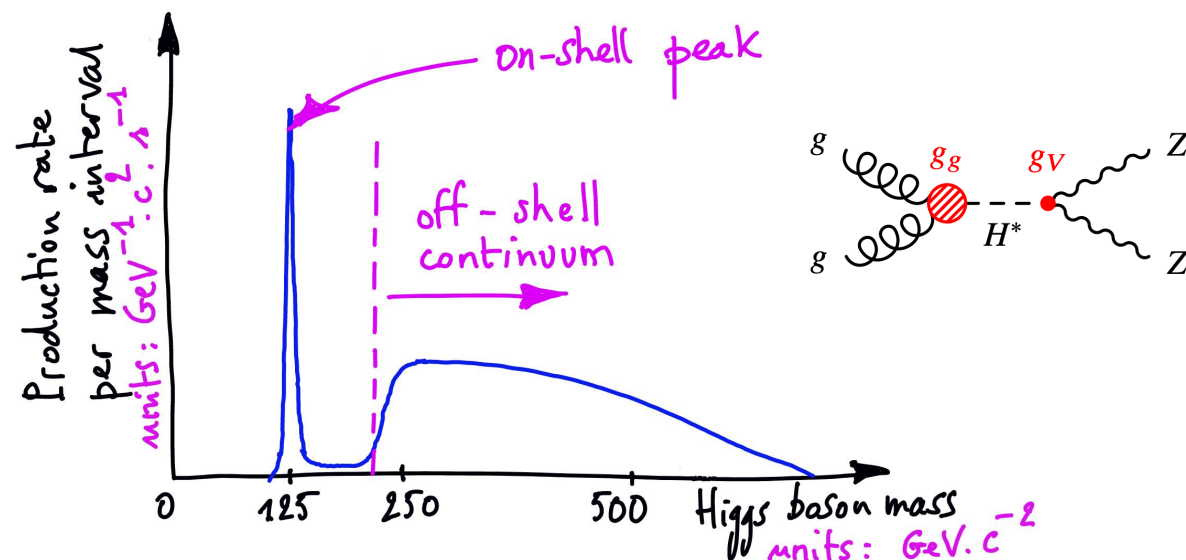
# Semi-parametric SBI



First ever measurement with real data using fully-differential high-dimensional information

|                                    | $\times$       | $\mu$          | $\nu$          |
|------------------------------------|----------------|----------------|----------------|
| Traditional Statistical Model      | parametric     | parametric     | parametric     |
| Full Neural SBI                    | non-parametric | non-parametric | non-parametric |
| Semi-parametric SBI (POI&Nuisance) | non-parametric | parametric     | parametric     |

# Semi-parametric SBI



First ever measurement with real data using fully-differential high-dimensional information

[Rep. Prog. Phys. 88 067801](#)

EUROPEAN ORGANISATION FOR NUCLEAR RESEARCH (CERN)

**ATLAS**  
EXPERIMENT

Rep. Prog. Phys. 88 (2025) 067801  
DOI: [10.1088/1361-6633/add370](#)

**An implementation of neural simulation-based inference for parameter estimation in ATLAS**

The ATLAS Collaboration

[Rep. Prog. Phys. 88 057803](#)

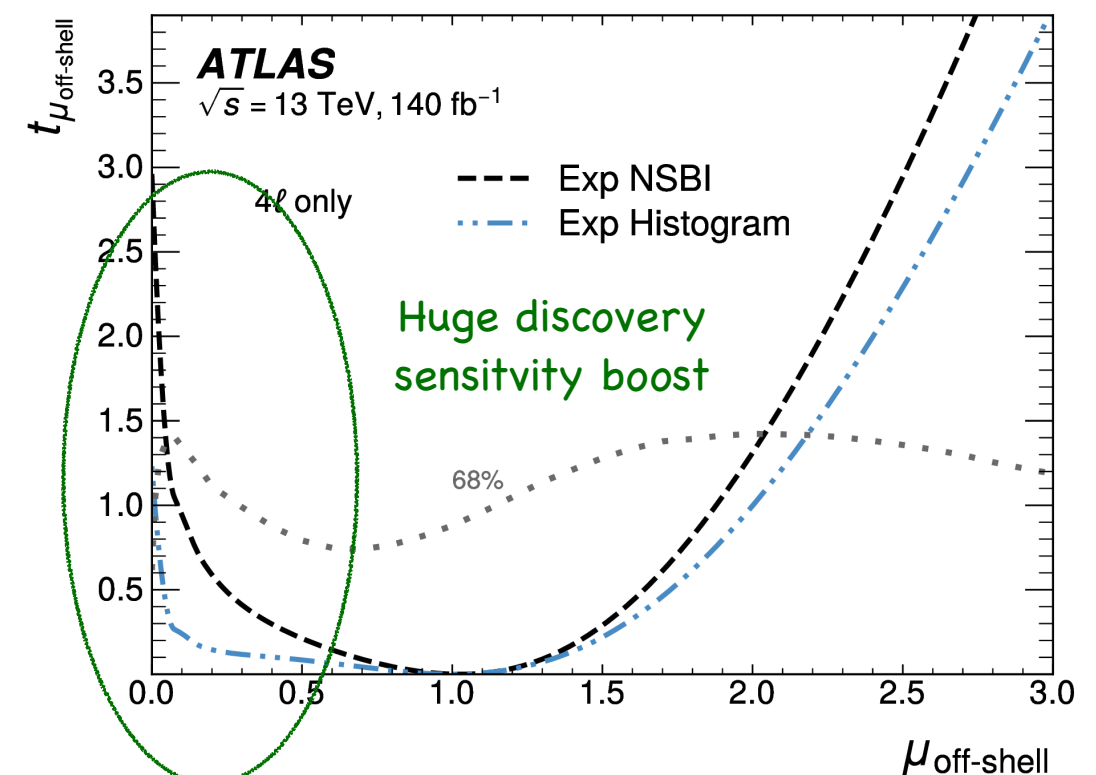
EUROPEAN ORGANISATION FOR NUCLEAR RESEARCH (CERN)

**ATLAS**  
EXPERIMENT

Rep. Prog. Phys. 88 (2025) 057803  
DOI: [DOI:10.1088/1361-6633/adcd9a](#)

**Measurement of off-shell Higgs boson production in the  $H^* \rightarrow ZZ \rightarrow 4\ell$  decay channel using a neural simulation-based inference technique in 13 TeV  $pp$  collisions with the ATLAS detector**

The ATLAS Collaboration



Both curves use same data, Histogram uses NN observable for  $x \rightarrow s(x)$

# Semi-parametric SBI

Refinable modeling for unbinned SMEFT analyses

Robert Schöfbeck

*Institute for High Energy Physics, Austrian Academy of Sciences,  
Dominikanerbastei 16, 1010 Vienna, Austria*

*Technische Universität Wien,  
Karlsplatz 13, 1040 Vienna*

[robert.schoefbeck@oeaw.ac.at](mailto:robert.schoefbeck@oeaw.ac.at)

Nuisance parameter dependent ratio:

$$\underbrace{\exp\left(\nu_A \hat{\Delta}_A(\mathbf{x})\right)}_{\text{Parametric ansatz}} \simeq \frac{d\sigma(\mathbf{x}|\boldsymbol{\nu})}{d\sigma(\mathbf{x}|\mathbf{0})} \quad \hat{\Delta}(x) \rightarrow \text{NN outputs}$$

Each  $x$ -dependent term is then fitted by minimizing the loss:

$$\begin{aligned} L[\hat{\Delta}_A(\mathbf{x})] &= \sum_{\nu \in \mathcal{V}} L_{\text{CE}}[\nu_A \hat{\Delta}_A(\mathbf{x})] \\ &= \sum_{\nu \in \mathcal{V}} \left( \left\langle \text{Soft}^+(\nu_A \hat{\Delta}_A(\mathbf{x})) \right\rangle_{\mathbf{x}, \mathbf{z}|\mathbf{0}} + \left\langle \text{Soft}^+(-\nu_A \hat{\Delta}_A(\mathbf{x})) \right\rangle_{\mathbf{x}, \mathbf{z}|\boldsymbol{\nu}} \right) \end{aligned}$$

Similar semi-parametric idea

Different implementation strategy

|                                    | $\mathbf{x}$   | $\mu$          | $\nu$          |
|------------------------------------|----------------|----------------|----------------|
| Traditional Statistical Model      | parametric     | parametric     | parametric     |
| Full Neural SBI                    | non-parametric | non-parametric | non-parametric |
| Semi-parametric SBI (POI&Nuisance) | non-parametric | parametric     | parametric     |



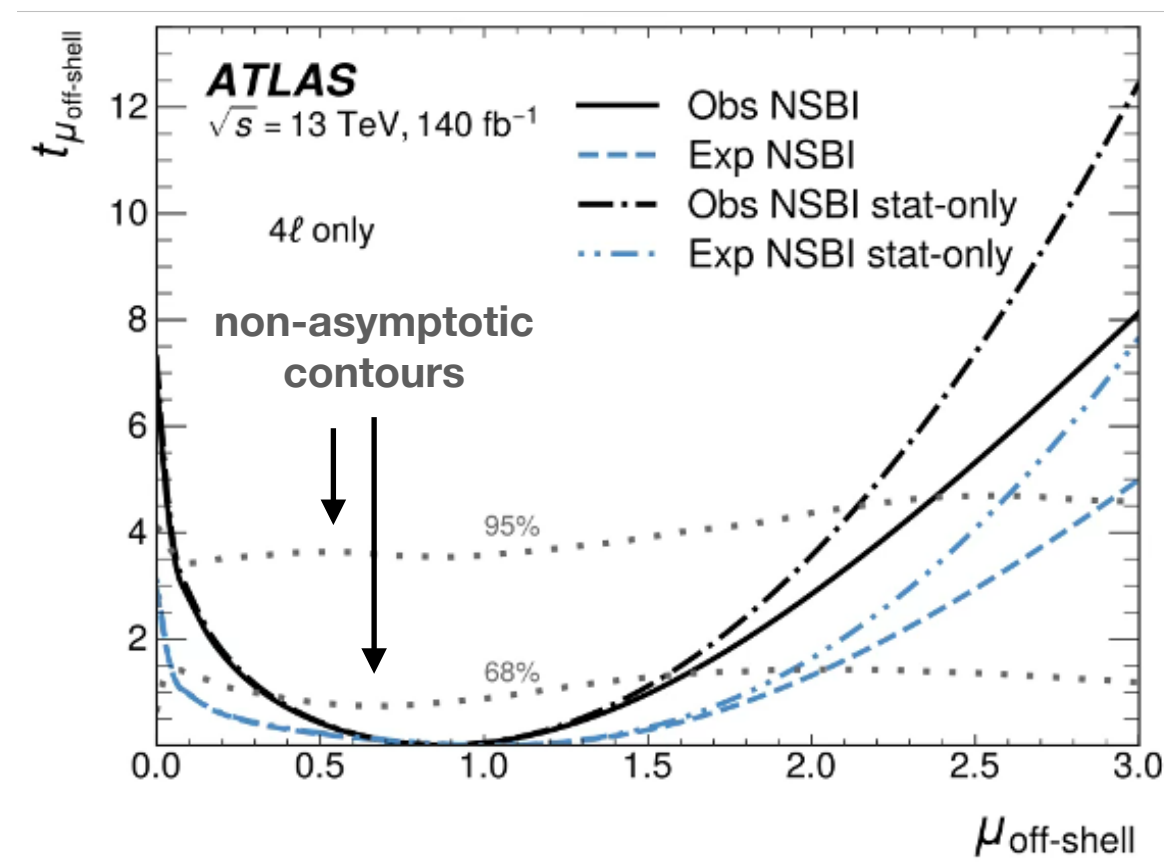
**We are now optimally constraining ALL parameters, with high-dimensional information from simulators.**

# Neyman Construction

## What if the neural networks are mis-modelled?

- The resulting test statistic we build does not retain the asymptotic properties typical of profile negative log-likelihood ratio test.
- Neyman Construction with simulator-sampled events more essential than before: build calibrated confidence intervals.

[Rep. Prog. Phys. 88 057803](#)



# Robust Model Building

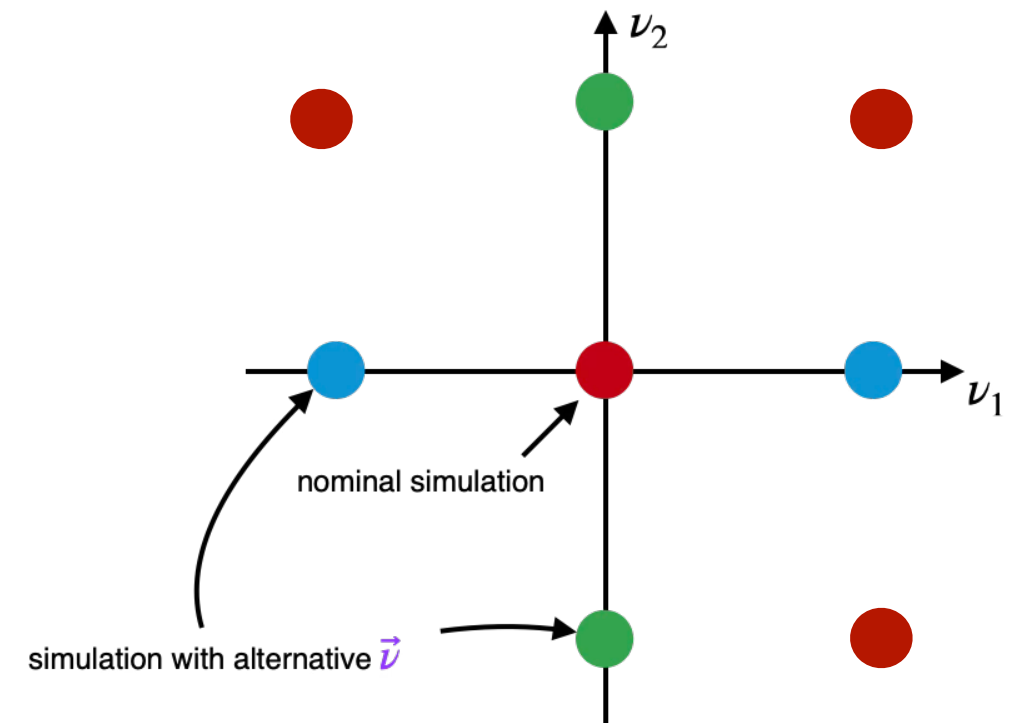
The HistFactory model for per-event density ratio:

$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \boxed{\frac{p_s(x_e | \nu)}{p_s(x_e)}} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

Can we do better than the  
HistFactory assumptions?

**Goal:**

modeling parameter space  
dependence more robustly



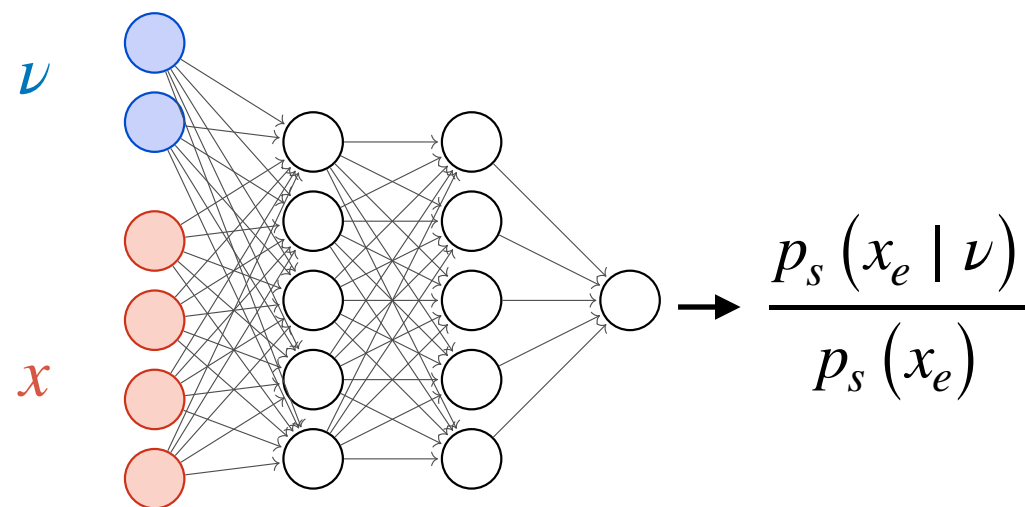
What if we generate simulations at  
more parameter points?

# Robust Model Building

The HistFactory model for per-event density ratio:

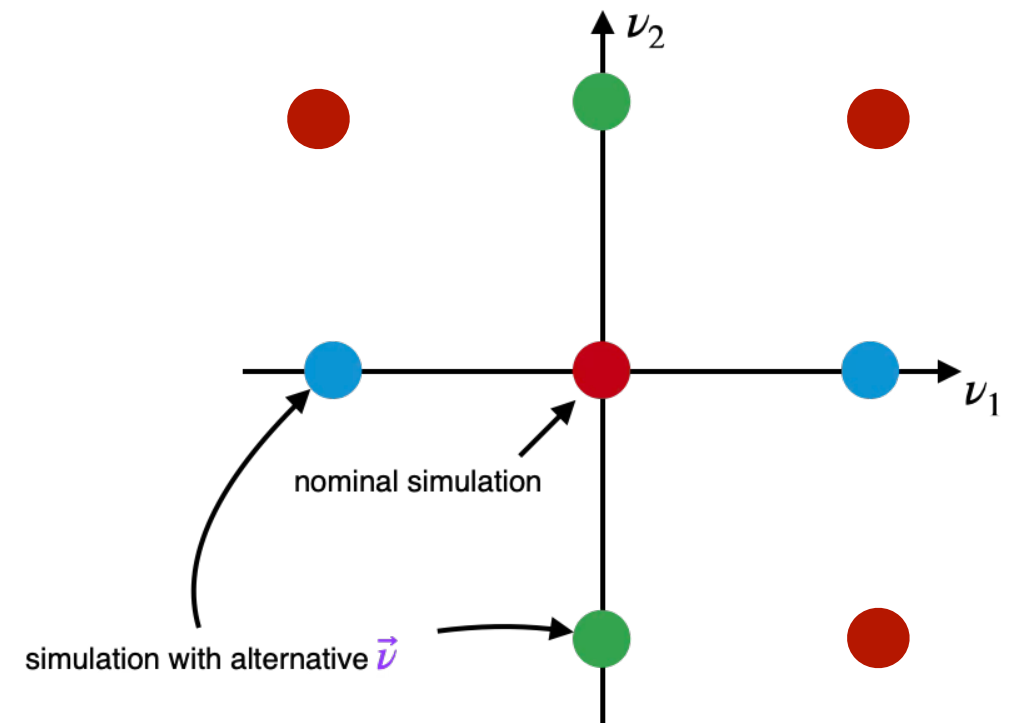
$$\frac{p(x_e | \mu, \nu)}{p_{\text{ref}}(x_e)} = \frac{1}{\sum_{s \in \text{samples}} \lambda_s(\mu, \nu)} \sum_{s \in \text{samples}} \lambda_s(\mu, \nu) \cdot \left[ \prod_{\mu} \kappa_s(\mu) \right] \cdot \boxed{\frac{p_s(x_e | \nu)}{p_s(x_e)}} \cdot \frac{p_s(x_e)}{p_{\text{ref}}(x_e)}$$

Can we do better than the HistFactory assumptions?



Option: train parameterized neural network classifiers?

Challenge: hard to regularise



What if we generate simulations at more parameter points?

# Robust Model Building

Semi-parametric option:

more robust parametric approach with relaxed factorisation assumptions.

Example:

$$\begin{aligned} \nu_A \hat{\Delta}_A(\mathbf{x}) = & \nu_{\text{tes}} \hat{\Delta}_{\text{tes}}(\mathbf{x}) + \nu_{\text{jes}} \hat{\Delta}_{\text{jes}}(\mathbf{x}) + \nu_{\text{met}} \hat{\Delta}_{\text{met}}(\mathbf{x}) \\ & + \nu_{\text{tes}}^2 \hat{\Delta}_{\text{tes,tes}}(\mathbf{x}) + \nu_{\text{jes}}^2 \hat{\Delta}_{\text{jes,jes}}(\mathbf{x}) + \nu_{\text{met}}^2 \hat{\Delta}_{\text{met,met}}(\mathbf{x}) \\ & + \boxed{\nu_{\text{tes}} \nu_{\text{jes}} \hat{\Delta}_{\text{tes,jes}}(\mathbf{x})} + \boxed{\nu_{\text{jes}} \nu_{\text{met}} \hat{\Delta}_{\text{jes,met}}(\mathbf{x})} \\ & + \boxed{\nu_{\text{tes}} \nu_{\text{met}} \hat{\Delta}_{\text{tes,met}}(\mathbf{x})}. \end{aligned}$$

Example of simultaneous variations

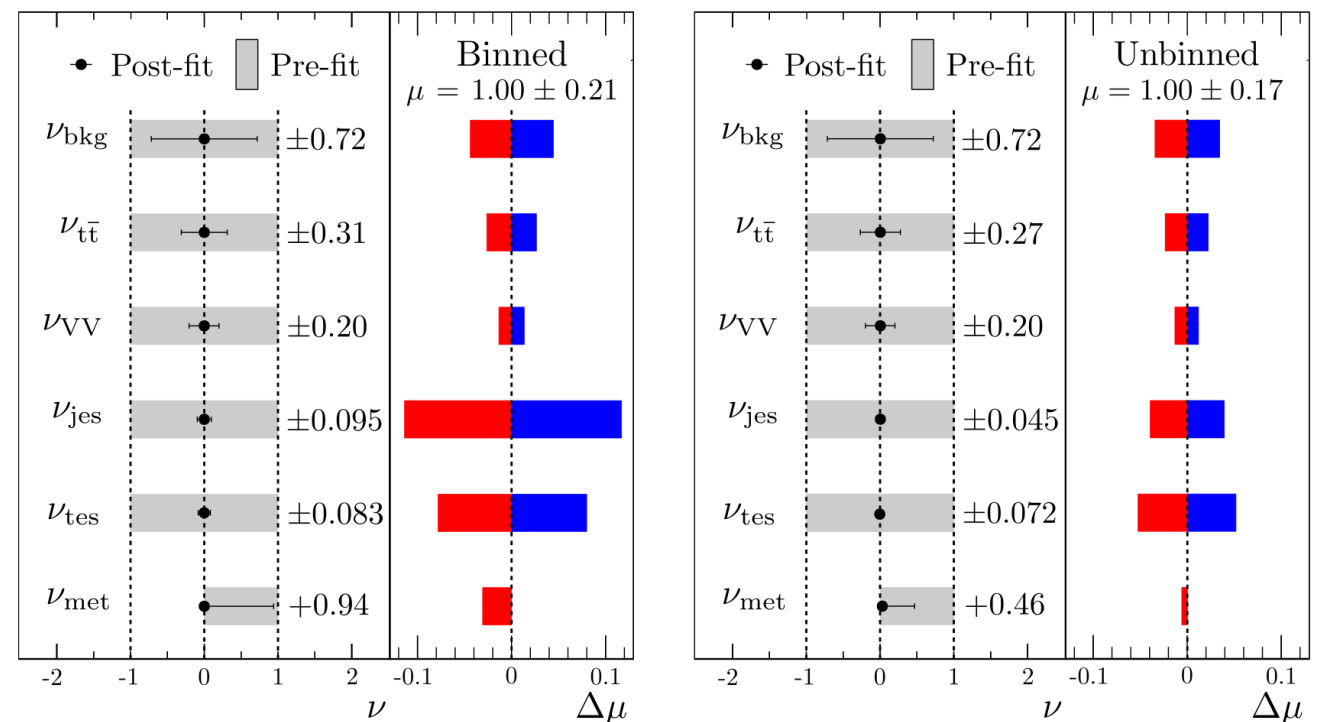
$\hat{\Delta}(x) \rightarrow$  NN outputs (unparameterized)

The polynomial ansatz can be made arbitrarily higher-order for expressivity

Unbinned inclusive cross-section measurements with machine-learned systematic uncertainties

Lisa Benato, Cristina Giordano, Claudius Krause, Ang Li,  
Robert Schöffbeck, Dennis Schwarz, Maryam Shooshtari, and Daohan Wang  
*Institute for High Energy Physics, Austrian Academy of Sciences Dominikanerbastei 16, 1010 Vienna, Austria*  
(Dated: August 26, 2025)

Reduced impact on sensitivity from NPs when fitting with SBI



low-dimensional summary

NSBI likelihood fit

# Robust Model Building

|                                        | $x$                    | $\mu$          | $\nu$          |
|----------------------------------------|------------------------|----------------|----------------|
| Traditional Statistical Model          | parametric             | parametric     | parametric     |
| Full Neural SBI                        | non-parametric         | non-parametric | non-parametric |
| Semi-parametric SBI (POI&Nuisance)     | non-parametric         | parametric     | parametric     |
| Semi-parametric SBI (POI& GP Nuisance) | non-parametric with NN | parametric     | GP not NN      |

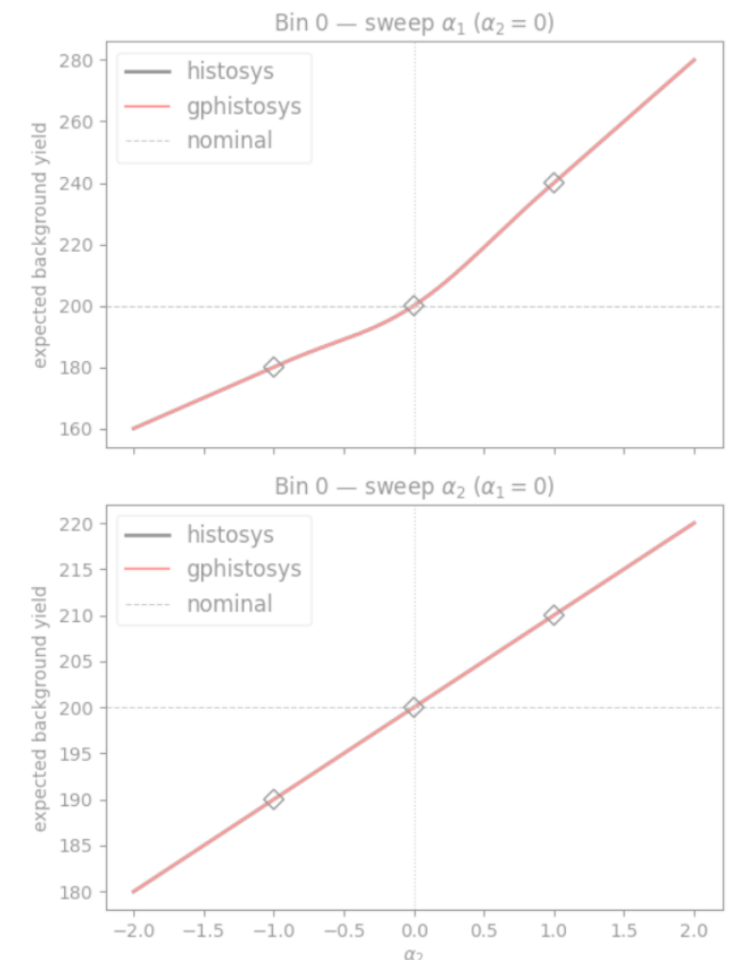
**Idea: flexibility of non-parametric model with control of parametric one**

Gaussian Processes are an excellent non-parametric replacement.

Regularisable with parametric priors, more expressive, and come with robust uncertainty quantification.

Ongoing work on GP modeling (Cranmer, Heinrich & Sandesara)

Easy to regularise with choice of prior



Gaussian Process = HistFactory

when densely sampled simulations are not available

(HistFactory used as prior)

# Semi-regularised Approach

|                                        | $x$                    | $\mu$          | $\nu$          |
|----------------------------------------|------------------------|----------------|----------------|
| Traditional Statistical Model          | parametric             | parametric     | parametric     |
| Full Neural SBI                        | non-parametric         | non-parametric | non-parametric |
| Semi-parametric SBI (POI&Nuisance)     | non-parametric         | parametric     | parametric     |
| Semi-parametric SBI (POI& GP Nuisance) | non-parametric with NN | parametric     | GP not NN      |

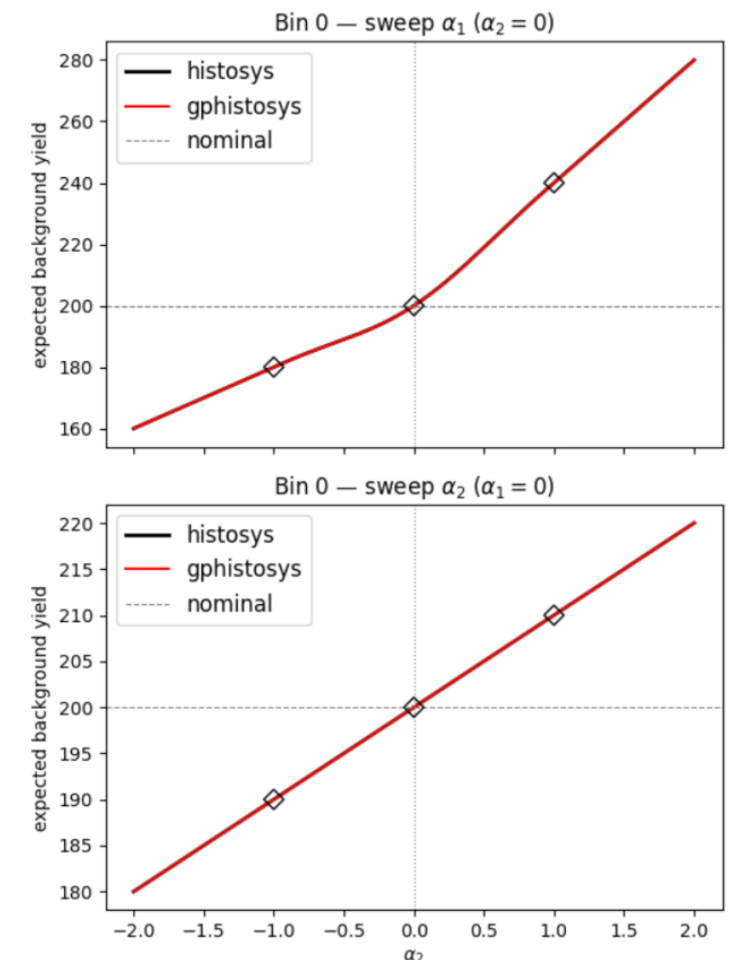
**Idea: flexibility of non-parametric model with control of parametric one**

Gaussian Processes are an excellent non-parametric replacement.

Regularisable with parametric priors, more expressive, and come with robust uncertainty quantification.

Ongoing work on GP modeling (Cranmer, Heinrich & Sandesara)

**Easy to regularise with choice of prior**



**Gaussian Process** = HistFactory

when densely sampled simulations are **not** available

(HistFactory used as prior)

# What is Simulation-Based Inference ?

|      |   | $x$                                    | $\mu$                  | $\nu$          |                |              |
|------|---|----------------------------------------|------------------------|----------------|----------------|--------------|
| SBI  | { | Traditional Statistical Model          | parametric             | parametric     | parametric     | NSBI for LHC |
|      |   | Full Neural SBI                        | non-parametric         | non-parametric | non-parametric |              |
| NSBI | { | Semi-parametric SBI (POI&Nuisance)     | non-parametric         | parametric     | parametric     |              |
|      |   | Semi-parametric SBI (POI& GP Nuisance) | non-parametric with NN | parametric     | GP not NN      |              |

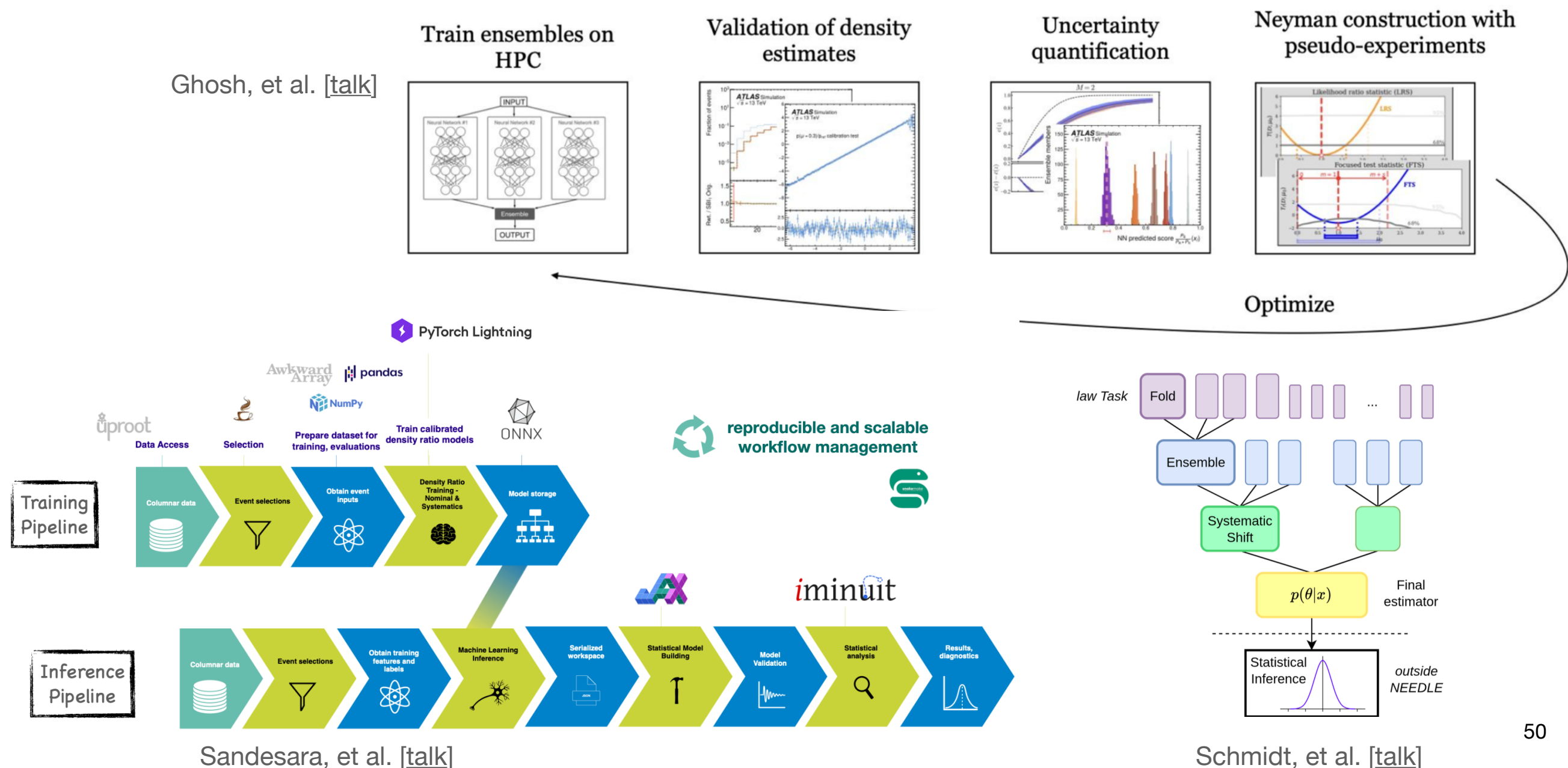
# What is Neural Simulation-Based Inference ?

# What is the future of NSBI?



# Building a Roadmap for wider adoption

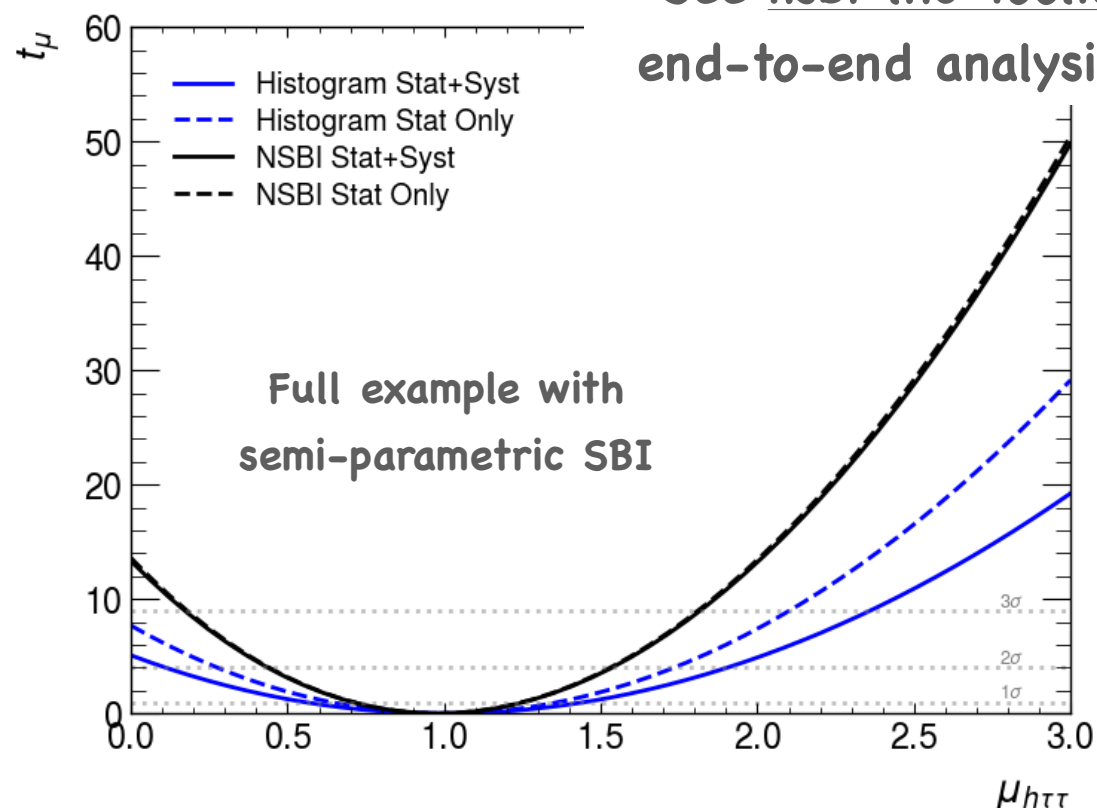
- Organised a ["blueprint" workshop](#) earlier this year, aiming to build a roadmap towards developing an ecosystem of tools that enables broad adoption of neural SBI at the LHC.
- Community whitepaper being written now outlining the plans and discussions.



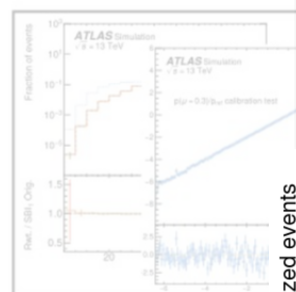
# Building a Roadmap for wider adoption

Source:

See [nsbi-lhc-toolkit](#) for an end-to-end analysis example



Validation of density estimates

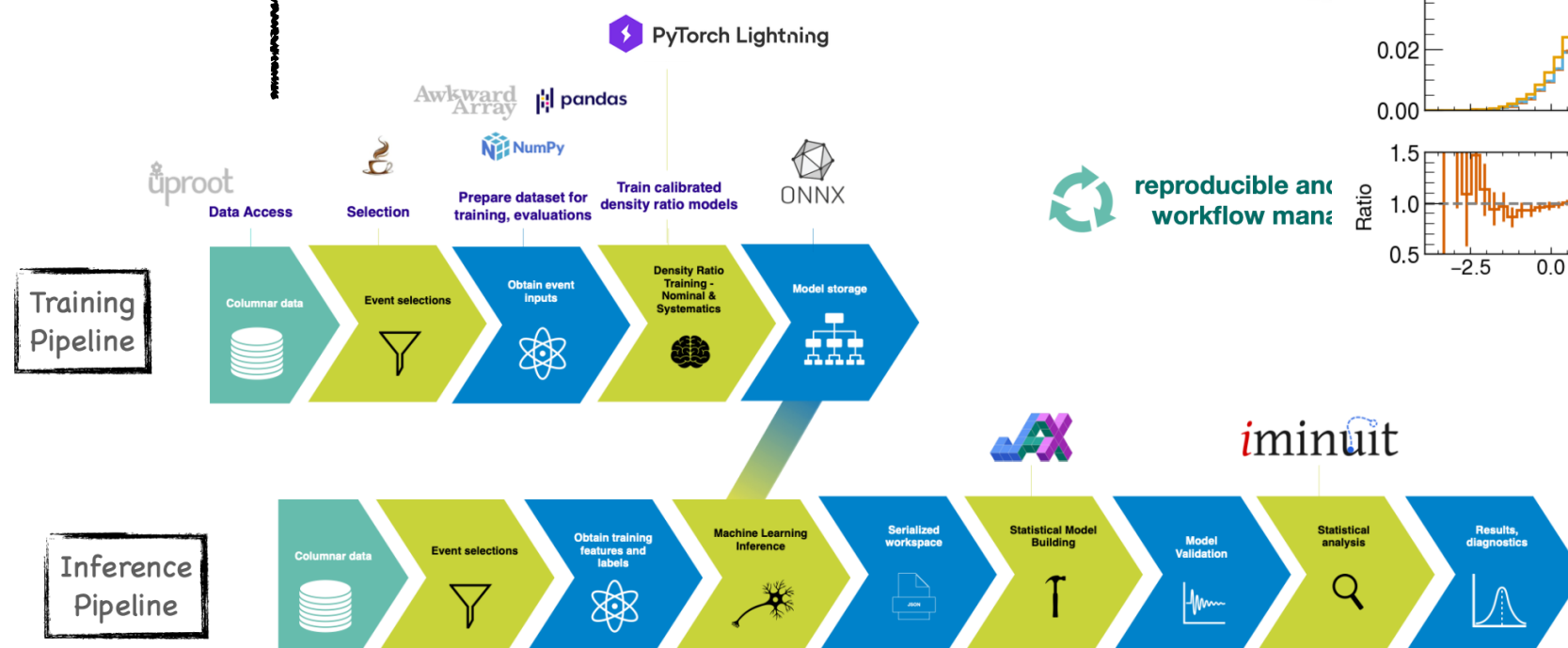
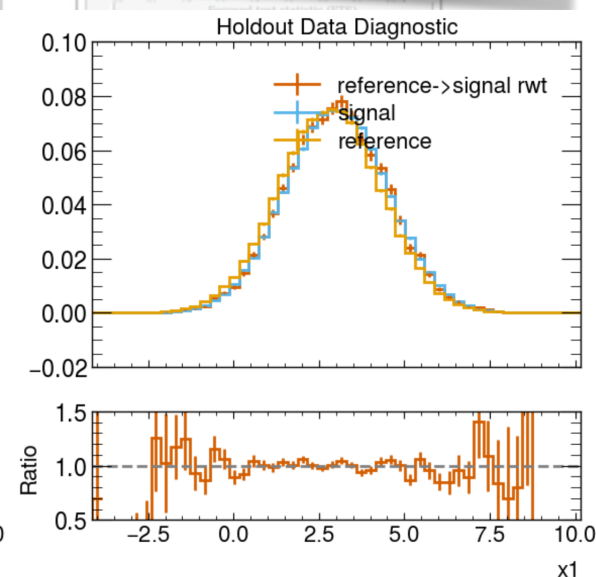
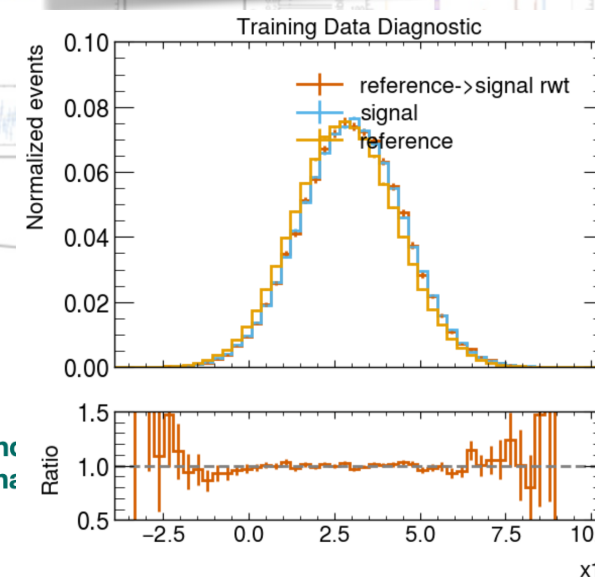


## Systematic uncertainties

Shape + normalization uncertainties point to the up/down varied ROOT files:

```
Systematics:
- Name: JES
  Type: NormPlusShape
  Nominal: 0
  Samples: [htautau, ztautau, ttbar]
  Up:
    - SampleName: htautau
      Path: ./saved_datasets/dataset_JES_up.root
      Tree: tree_htautau
      Weight: weights
  Dn:
    - SampleName: htautau
      Path: ./saved_datasets/dataset_JES_dn.root
      Tree: tree_htautau
      Weight: weights
```

- **Nominal: 0** — the nuisance parameter starts at zero (the Gaussian constraint is centred here).
- The workspace builder computes variation ratios (varied / nominal) automatically from the histograms.

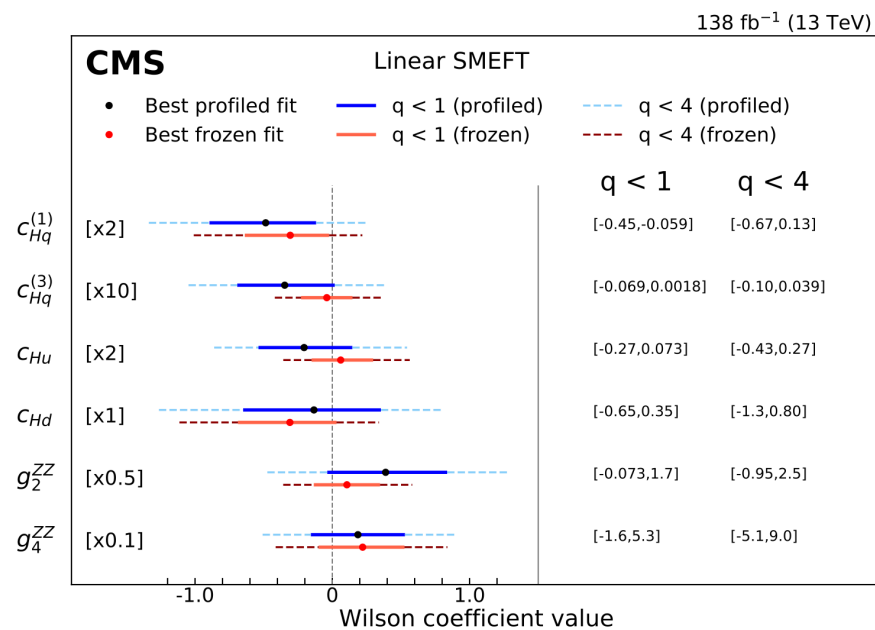
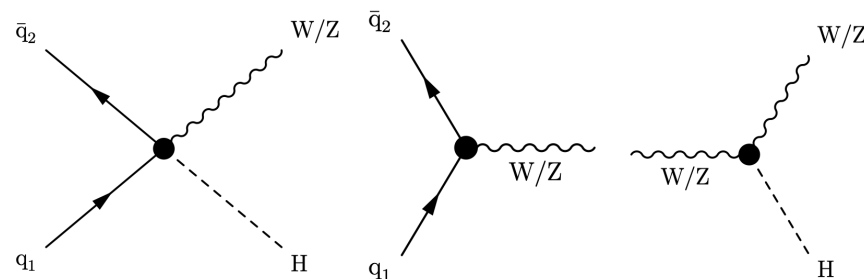


Sandesara, et al. [talk]

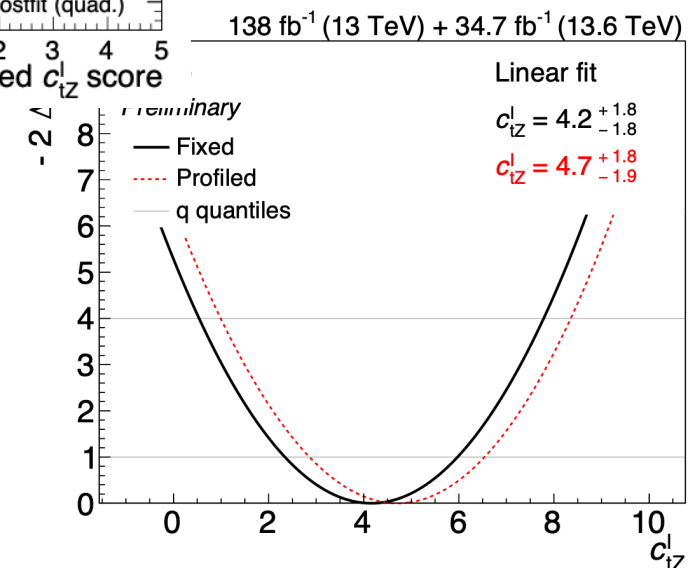
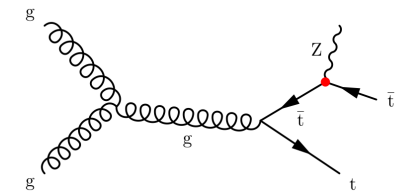
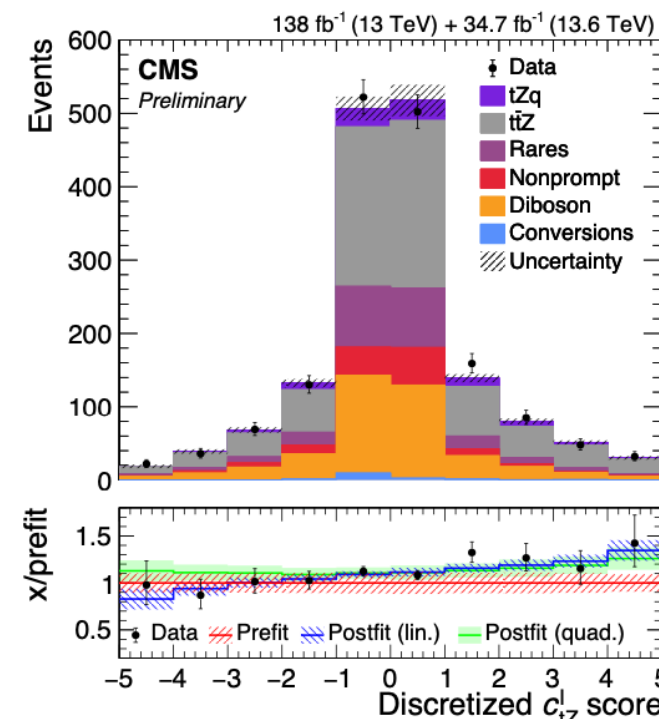
Schmidt, et al. [talk]

# Summary and Outlook

- Sufficient summary-based SBI not covered in this talk – can achieve dimensionality reduction  $x \rightarrow s(x)$  without losing sensitivity to parameters.



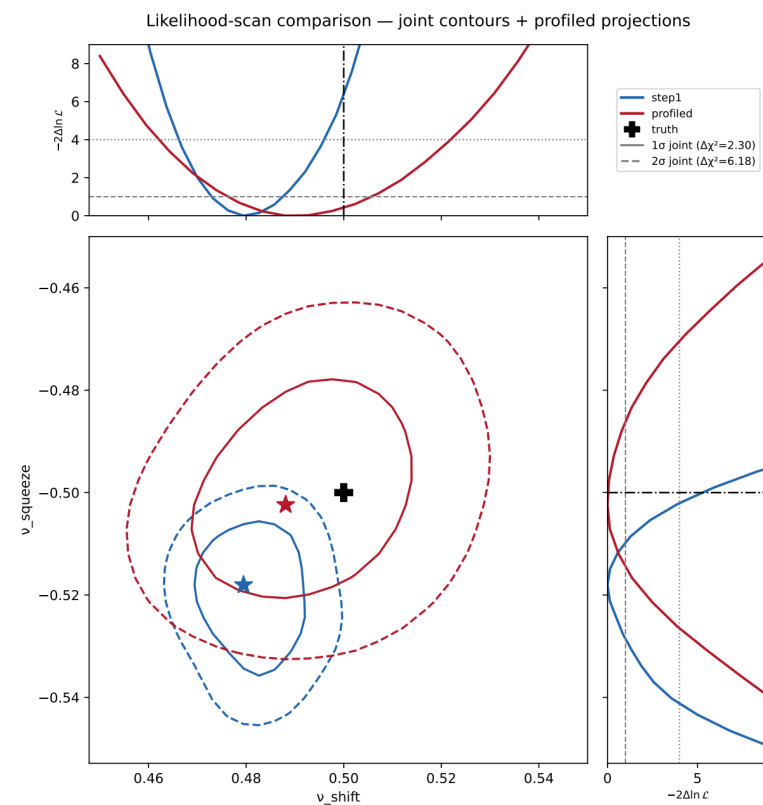
[2411.16907]



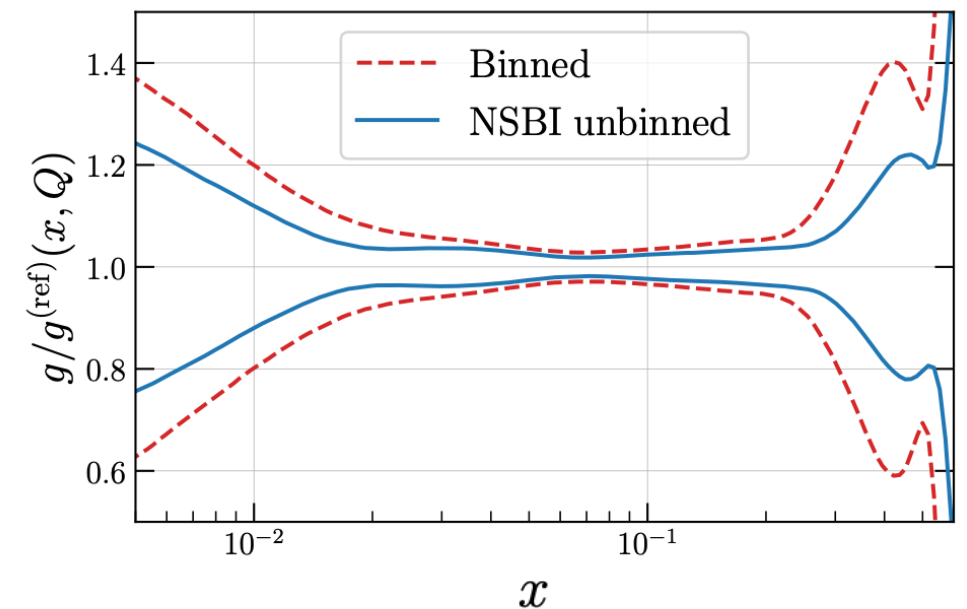
[2411.16907, 2405.13524]

# Summary and Outlook

- Sufficient summary-based SBI not covered in this talk – can achieve dimensionality reduction  $x \rightarrow s(x)$  without losing sensitivity to parameters.
- A lot of studies and work are being done to make NSBI robust and scalable in real experiments.



Systematics aware semi-parametric training but with distributions directly, via Normalizing Flows  
[\[2602.13184\]](#)



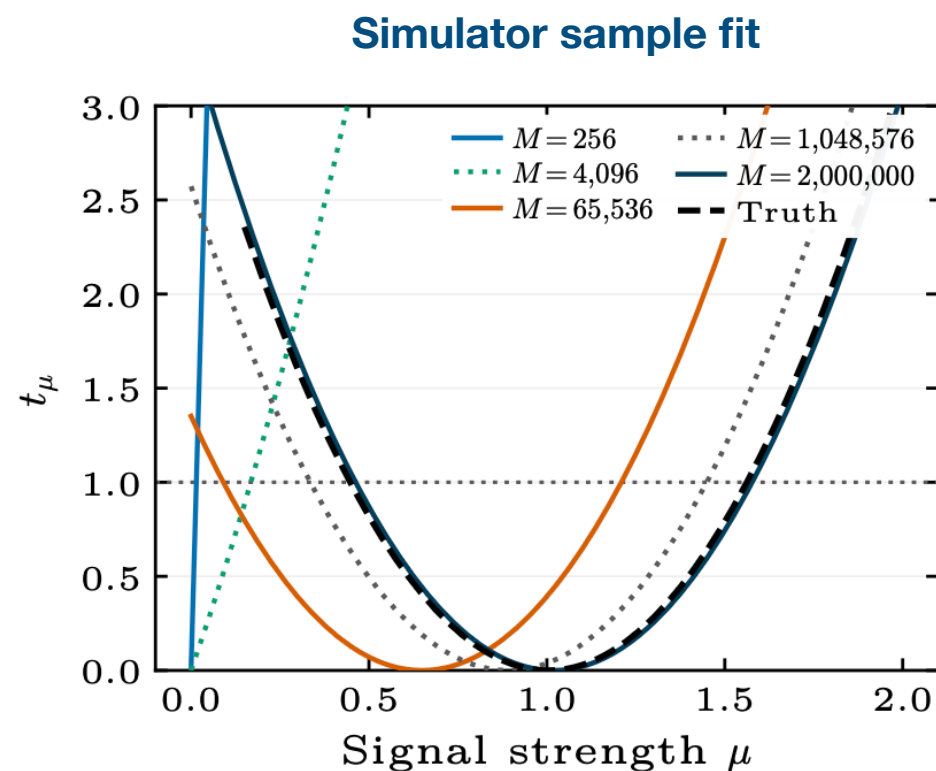
Reduce the reliance on coarse approximations of uncertainties and their correlations, which are propagated as nuisance parameters.

[\[2604.13157\]](#)

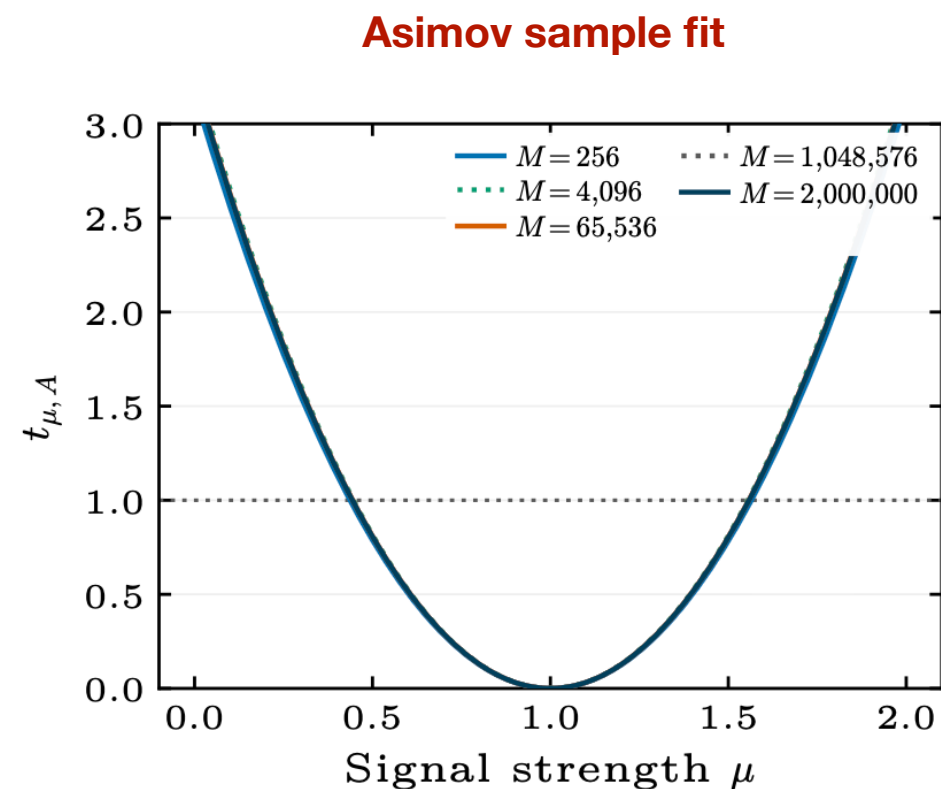
# Summary and Outlook

- Sufficient summary-based SBI not covered in this talk – can achieve dimensionality reduction  $x \rightarrow s(x)$  without losing sensitivity to parameters.
- A lot of studies and work are being done to make NSBI robust and scalable in real experiments.

Extend the Asimov idea now to the unbinned approach [2609.14136]



**2M events to reproduce unbiased truth**



**256 events reproduce the unbiased truth**

# Summary and Outlook

- Sufficient summary-based SBI not covered in this talk – can achieve dimensionality reduction  $x \rightarrow s(x)$  without losing sensitivity to parameters.
- A lot of studies and work are being done to make NSBI robust and scalable in real experiments.
- We are entering the era of precision: NSBI is an integral technology for accelerating the discovery potential of our very expensive experiment.

**Continue to beat projections!**

example figure from [here](#)

