



Reliable Likelihood Ratios for Neural Simulation-Based Inference in HEP

Nino Kovačić*, Stephen Jiggins & the NEEDLE team

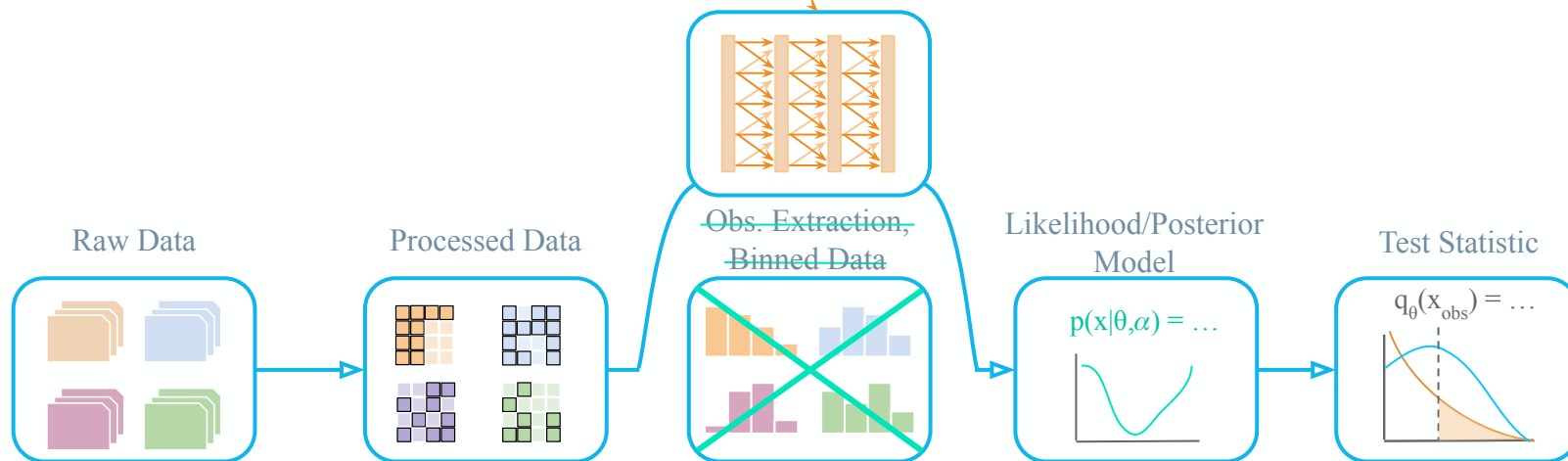
**nkovacic.phy@pmf.hr , <https://needle-sbi.github.io>*



A brief introduction to NSBI and NLRE

- What is Neural Simulation-Based Inference?^[1]

- The core idea is to utilize NNs as estimators (surrogates) for the likelihood/posterior
- How does that affect a typical workflow? (simplifying greatly)



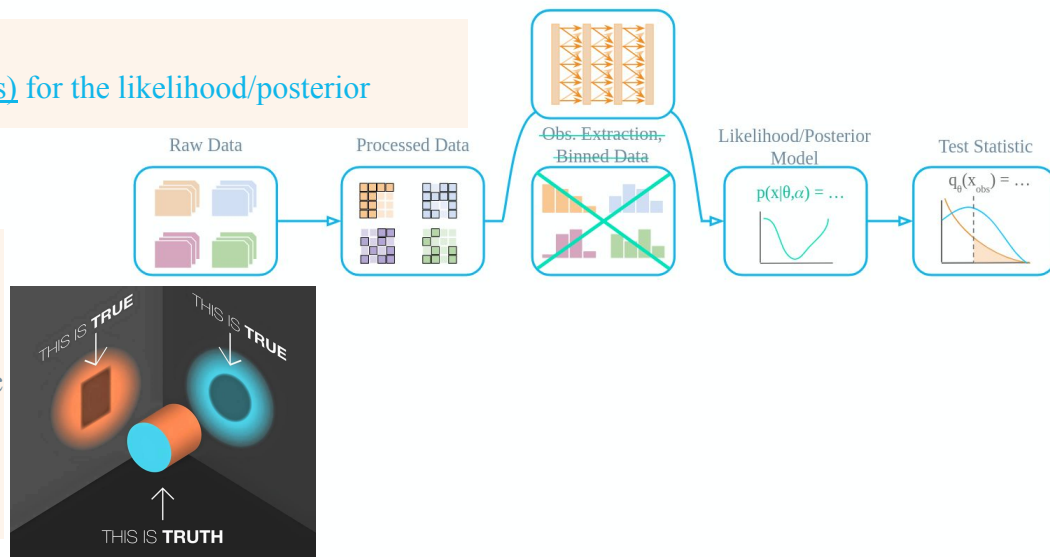
A brief introduction to NSBI and NLRE

- **What is Neural Simulation-Based Inference?**^[1]

- The core idea is to utilize NNs as estimators (surrogates) for the likelihood/posterior

- **What does NSBI offer?**

- Main point: increase in statistical power:
 - According to the Neyman–Pearson lemma, the likelihood ratio is the most powerful* test statistic (and that is what NSBI aims to obtain)
 - No dimensionality reduction of datasets required
 - No averaging of information across bins**



*Strictly speaking, we should be talking about sufficiency, this has been simplified in the interest of time

**If performing unbinned fits

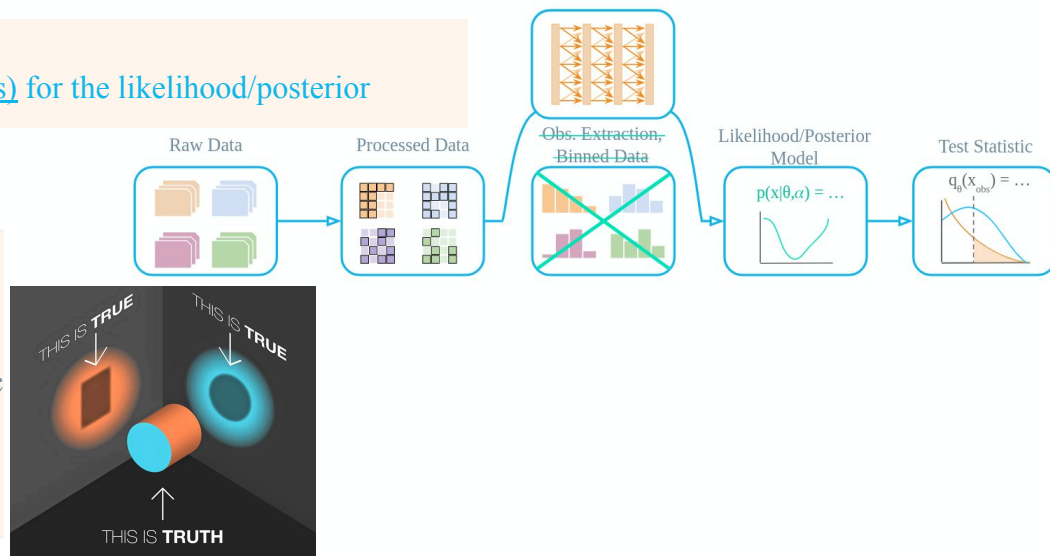
A brief introduction to NSBI and NLRE

• What is Neural Simulation-Based Inference?^[1]

- The core idea is to utilize NNs as estimators (surrogates) for the likelihood/posterior

• What does NSBI offer?

- Main point: increase in statistical power:
 - According to the Neyman–Pearson lemma, the likelihood ratio is the most powerful* test statistic (and that is what NSBI aims to obtain)
 - No dimensionality reduction of datasets required
 - No averaging of information across bins**



• What is Neural Likelihood Ratio Estimation (NLRE)?

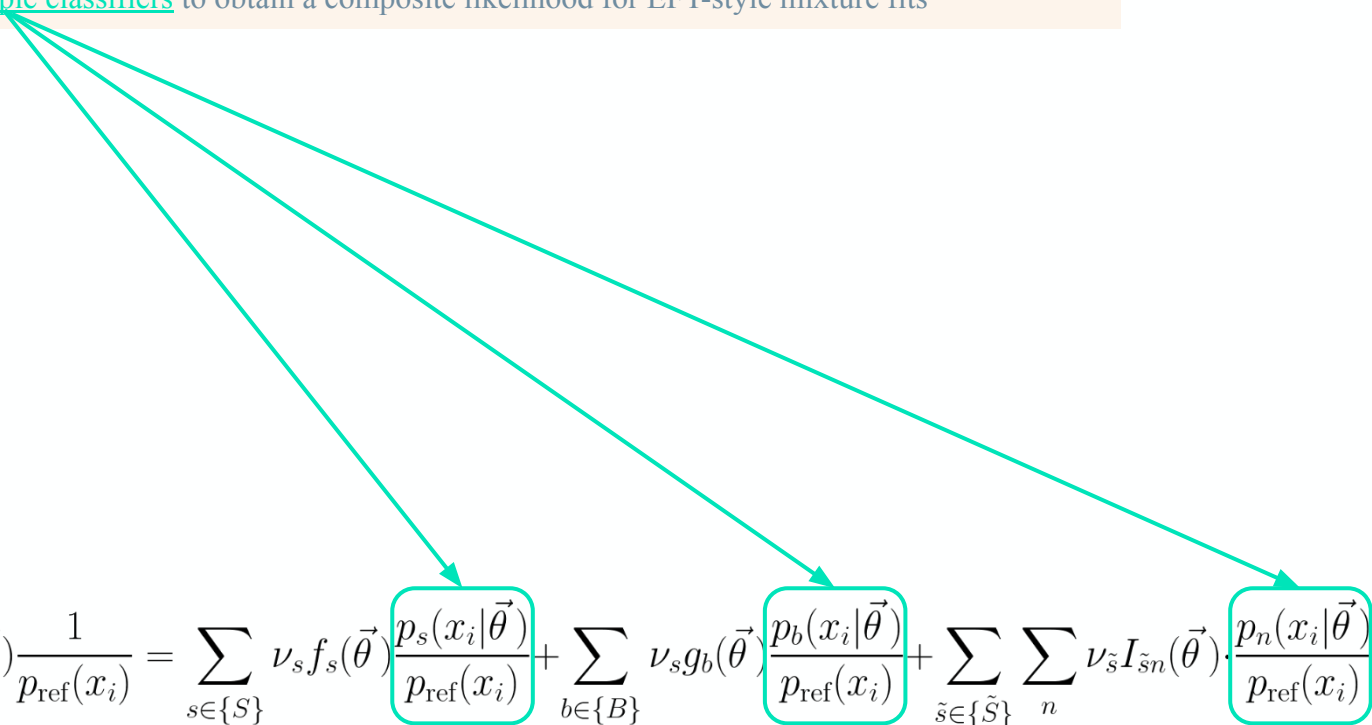
- A flavor of NSBI where the probability ratios are estimated via NNs and used to form the likelihood model
- In this talk we will refer to NLRE interchangeably with Calibrated Classifiers (CARL)^[2] method
- Here, NN binary classification scores are leveraged to extract probability ratios

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems

- E.g. combining multiple classifiers to obtain a composite likelihood for EFT-style mixture fits


$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \boxed{\frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}} + \sum_{b \in \{B\}} \nu_s g_b(\vec{\theta}) \boxed{\frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \cdot \boxed{\frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}}$$

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
 - E.g. [combining multiple classifiers](#) to obtain a composite likelihood for EFT-style mixture fits
- Well behaved training and an “easy” training objective:
 - Unlike NLE/NPE methods, binary classifiers do not need to learn a true density
 - Thus, there is no need to learn normalization which is the hardest component of density estimation

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_s g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
 - E.g. [combining multiple classifiers](#) to obtain a composite likelihood for EFT-style mixture fits
- Well behaved training and an “easy” training objective:
 - Unlike NLE/NPE methods, binary classifiers do not need to learn a true density
 - Thus, there is no need to learn normalization which is the hardest component of density estimation
 - Practically, NLRE is “promoted” to NLE by the [common reference](#) factoring out of the likelihood or its exact cancellation in profile likelihood fits

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_s g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \cdot \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**
 - Easily constructed likelihoods for HEP inference problems
 - Well behaved training and an “easy” training objective
- **What are the challenges with NLRE?**

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_s g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \cdot \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \cdot \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \cdot \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

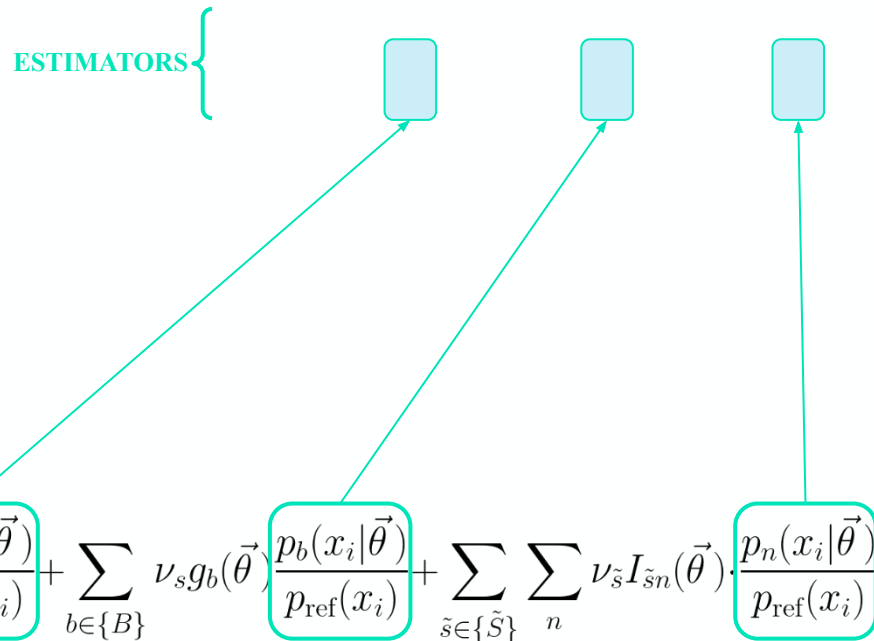
NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)
 - We need to train multiple estimators



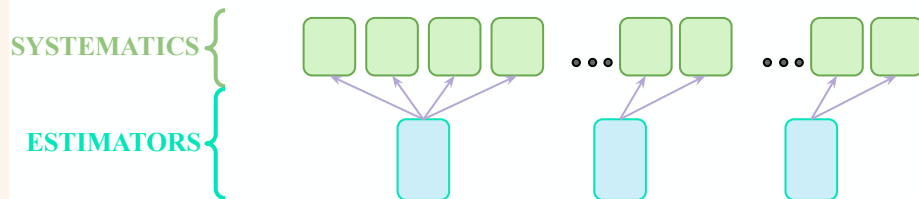
NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)
 - We need to train multiple estimators, learn the systematics



$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

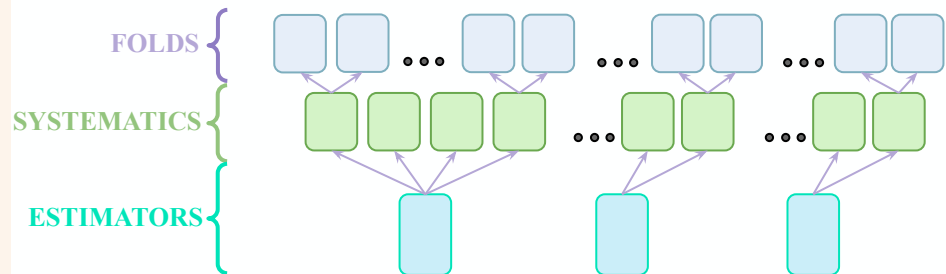
NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)
 - We need to train multiple estimators, learn the systematics, train via k-fold cross-validation



$$\mathcal{L}(\vec{\theta} | x_i) \frac{1}{\text{const.}} = p(x_i | \vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i | \vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i | \vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \frac{p_n(x_i | \vec{\theta})}{p_{\text{ref}}(x_i)}$$

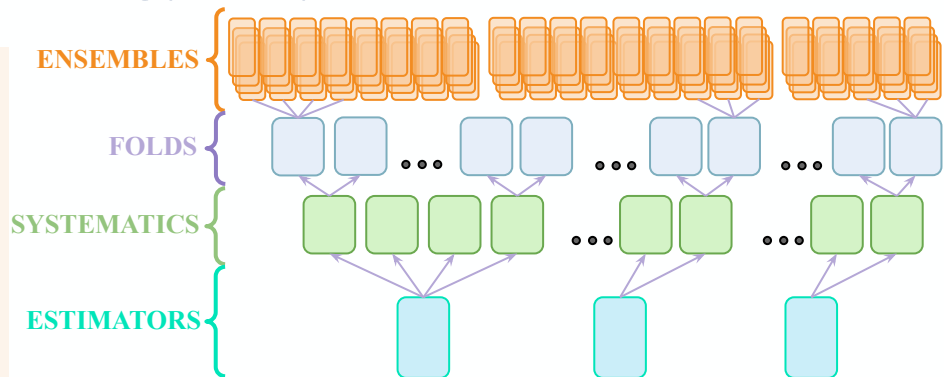
NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)
 - We need to train multiple estimators, learn the systematics, train via k-fold cross-validation, and ensemble



$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

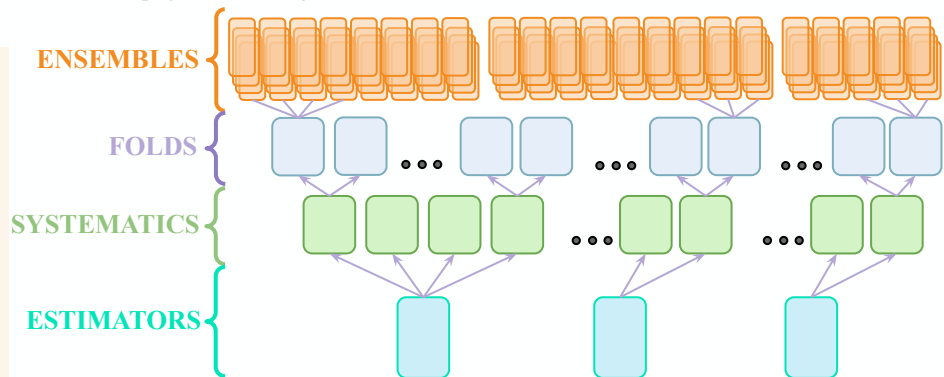
NLRE in High Energy Physics

- **Why is NLRE well-suited for HEP applications:**

- Easily constructed likelihoods for HEP inference problems
- Well behaved training and an “easy” training objective

- **What are the challenges with NLRE?**

- Exposed to any issues in estimator quality:
 - Derived ratios need to be of very high quality in order to perform an unbinned fit
 - Composite likelihoods mean that the errors compound!
- Handling training and bookkeeping:
 - Training and inference still need to happen on very large scales (of HEP datasets)
 - We need to train multiple estimators, learn the systematics, train via k-fold cross-validation, and ensemble
 - These costs are multiplicative, but parallelizable!

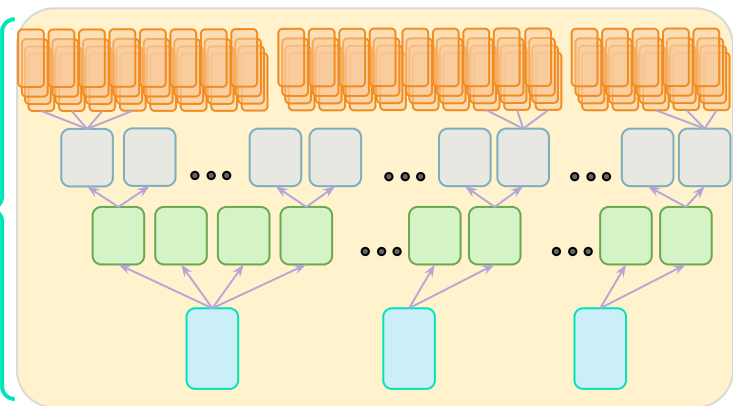


- Finally, for inference the flow through this DAG needs to “reverse” as it requires aggregation over all these NNs!

$$\mathcal{L}(\vec{\theta}|x_i) \frac{1}{\text{const.}} = p(x_i|\vec{\theta}) \frac{1}{p_{\text{ref}}(x_i)} = \sum_{s \in \{S\}} \nu_s f_s(\vec{\theta}) \frac{p_s(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{b \in \{B\}} \nu_b g_b(\vec{\theta}) \frac{p_b(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)} + \sum_{\tilde{s} \in \{\tilde{S}\}} \sum_n \nu_{\tilde{s}} I_{\tilde{s}n}(\vec{\theta}) \frac{p_n(x_i|\vec{\theta})}{p_{\text{ref}}(x_i)}$$

needle-sbi: solving the practical deployment of NSBI in HEP

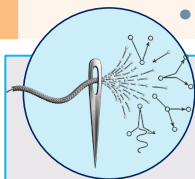
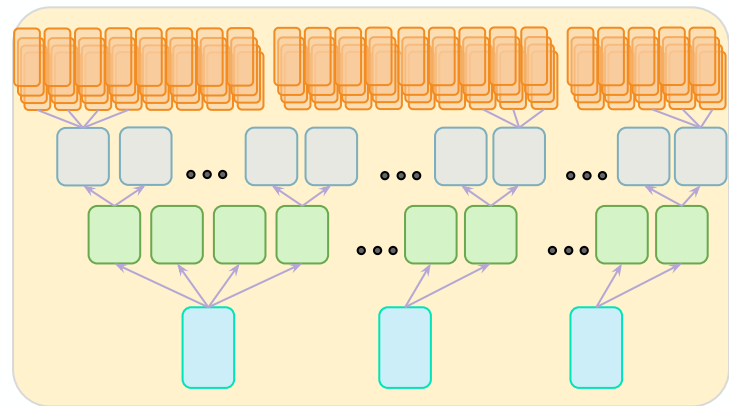
- **needle-sbi*** is a workflow orchestrator for large-scale NSBI in HEP:
 - Unified DAG execution with either LAW^[3] or b2luigi^[4] :
it stitches the models together!



needle-sbi: solving the practical deployment of NSBI in HEP

- **needle-sbi is a workflow orchestrator for large-scale NSBI in HEP:**

- Unified DAG execution with either LAW^[3] or b2luigi^[4]: it stitches the models together!
- It combines efficient memory usage and data loading with dask/PyTorch
- Allows for parallel and vectorized NN model evaluation
- Supports the common workload managers for HTC/HPC:
 - HTCondor, Slurm, LSF



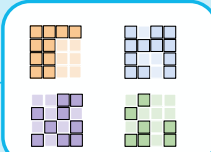
YAML HYDRA



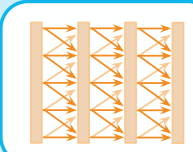
Raw Data



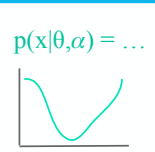
Processed Data



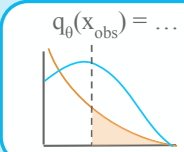
NN estimator(s)



Likelihood/Posterior Model



Test Statistic



Up to the user, but can be fully integrated into the needle-sbi DAG

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



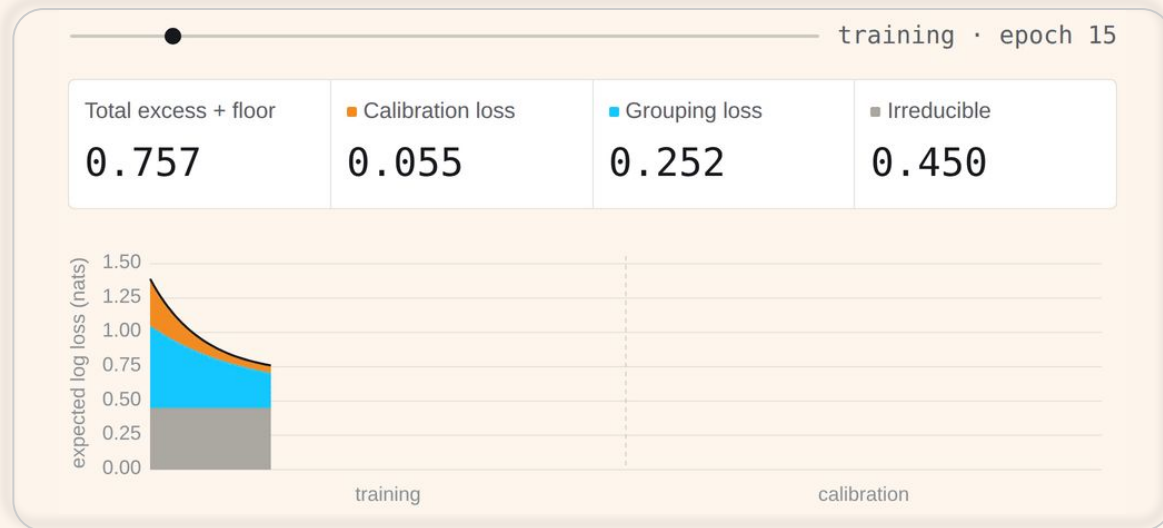
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$ the Bayes-optimal classifier and $\mathcal{L}^*$ its loss
- $S, f(x), s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
- $g(f(x)), \bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



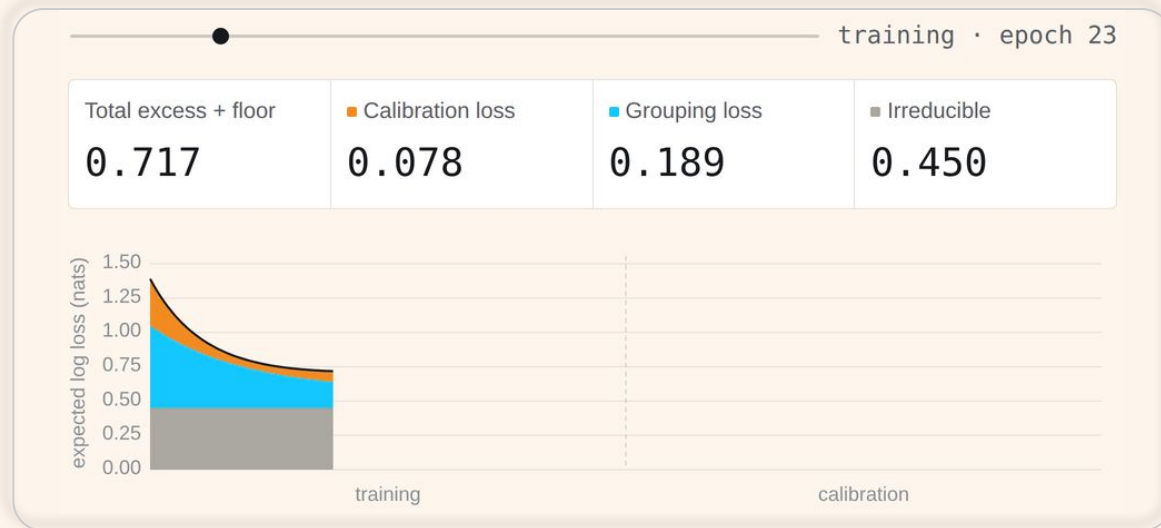
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



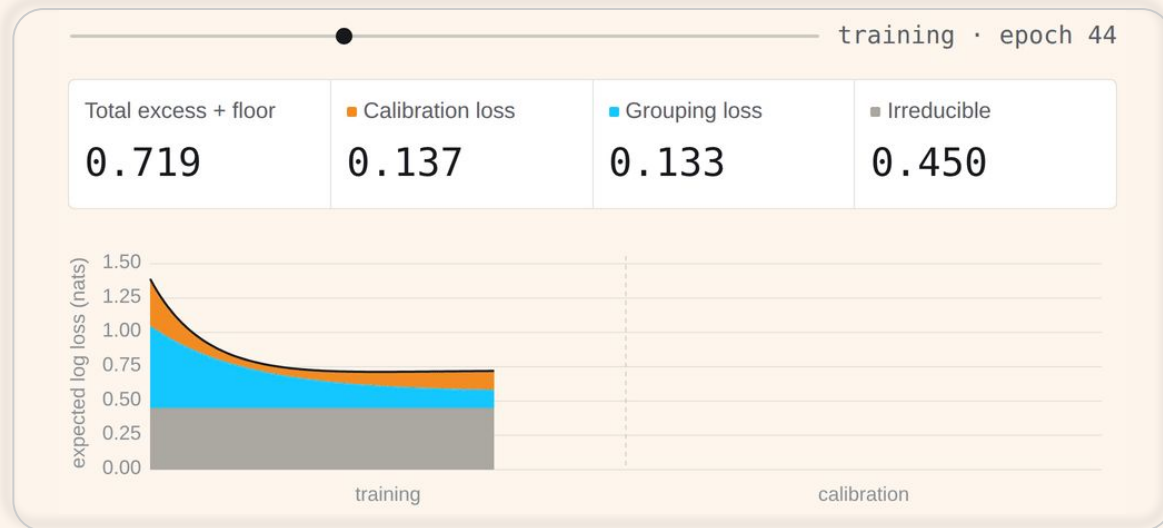
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



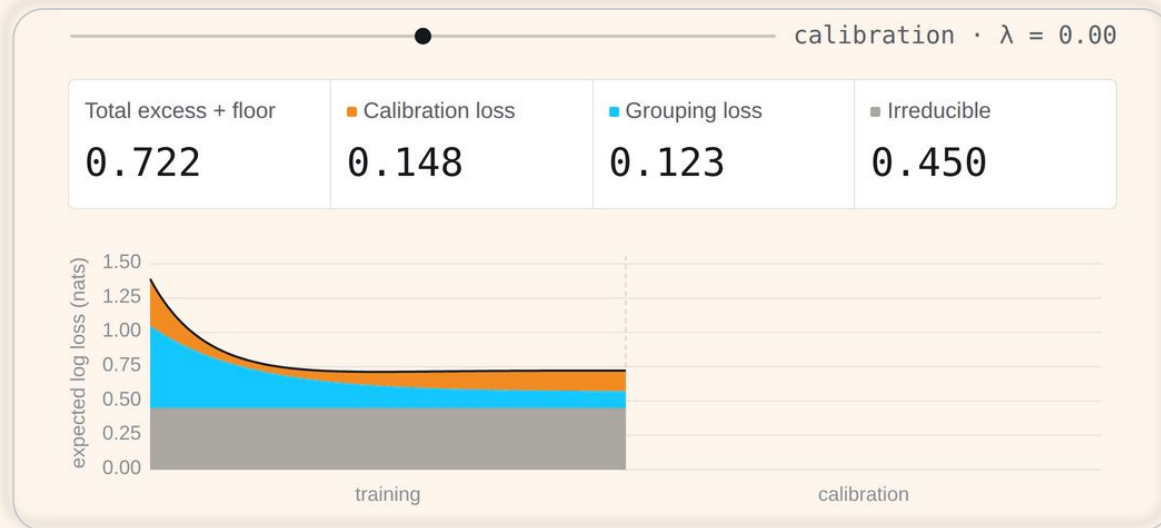
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



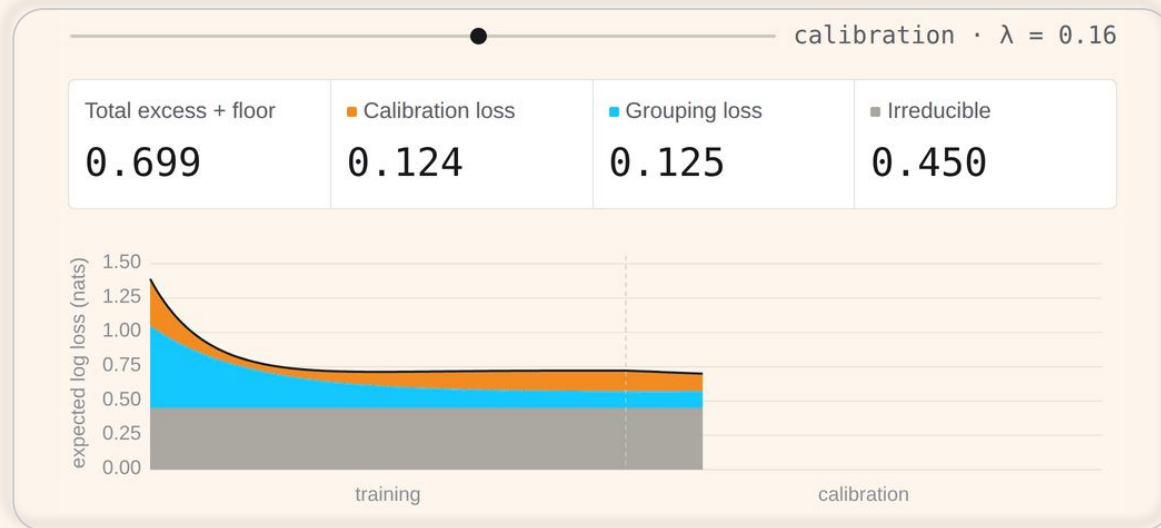
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



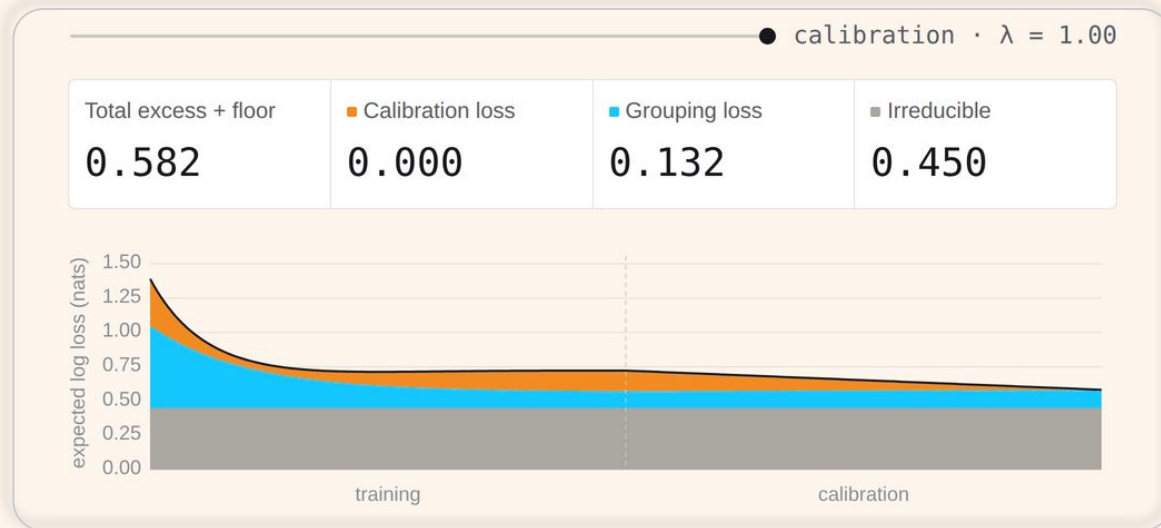
- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- Classifier error/loss decomposes into 3 terms:

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be



- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Calibration and classifier error decomposition

- **Classifier error/loss decomposes into 3 terms:**

$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be

- **For a given post-hoc calibration map g that acts on the classifier outputs alone:**

$$g \circ f \text{ is a function of } f \implies \text{GL}(g \circ f) \geq \text{GL}(f)$$

- Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

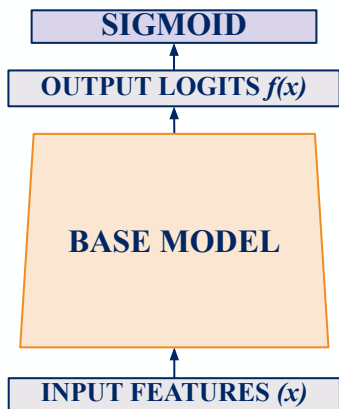
Calibration and classifier error decomposition

- **Classifier error/loss decomposes into 3 terms:**

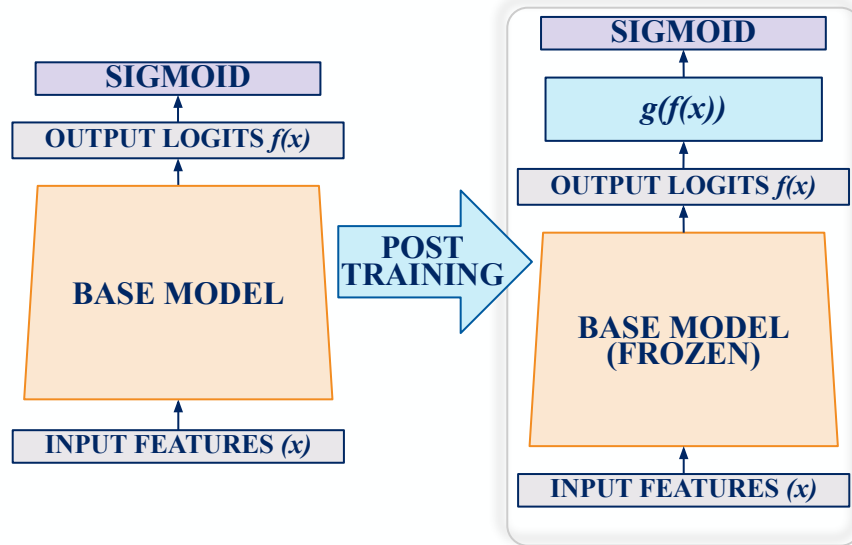
$$\mathbb{E}[\text{KL}(Y \parallel S)] = \underbrace{\mathbb{E}[\text{KL}(Y \parallel S^*)]}_{\text{irreducible}} + \underbrace{\mathbb{E}[\text{KL}(S^* \parallel C)]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[\text{KL}(C \parallel S)]}_{\text{calibration, CL}}$$

- Grouping loss $\text{GL}^{[6]}$ is the consequence of the classifier mapping two events epsilon-close in their scores when they shouldn't be
- **For a given post-hoc calibration map g that acts on the classifier outputs alone:**
 $g \circ f$ is a function of $f \implies \text{GL}(g \circ f) \geq \text{GL}(f)$
- **Even an ideal post-hoc calibration map that fully eliminates the calibration loss (CL), leaves GL is at best unchanged**
 $\text{CL} \rightarrow 0 \text{ (ideal post-hoc } g) \implies \mathcal{L}(g \circ f) - \mathcal{L}^* \longrightarrow \text{GL}(f)$
 - This means that if GL is present, the Neyman-Pearson lemma optimality argument (or sufficiency of the probability ratio) collapses and this classifier sub-optimality cannot be removed by post-hoc calibration
 - Above KL is the Kullback–Leibler divergence, $\mathcal{L}$ the loss (e.g. BCE or Brier), $C(x) = \mathbb{E}[Y|S(X)]$
 S , $f(x)$, $s(x) = \sigma[f(x)]$ the classifier, its logit and score output, S^* Bayes-optimal classifier and $\mathcal{L}^*$ its loss
 $g(f(x))$, $\bar{s}(x) = \sigma[g(f(x))]$ post-hoc (via calibration map g) calibrated classifier

Post-hoc Calibration vs Fine-tuning



Post-hoc Calibration vs Fine-tuning

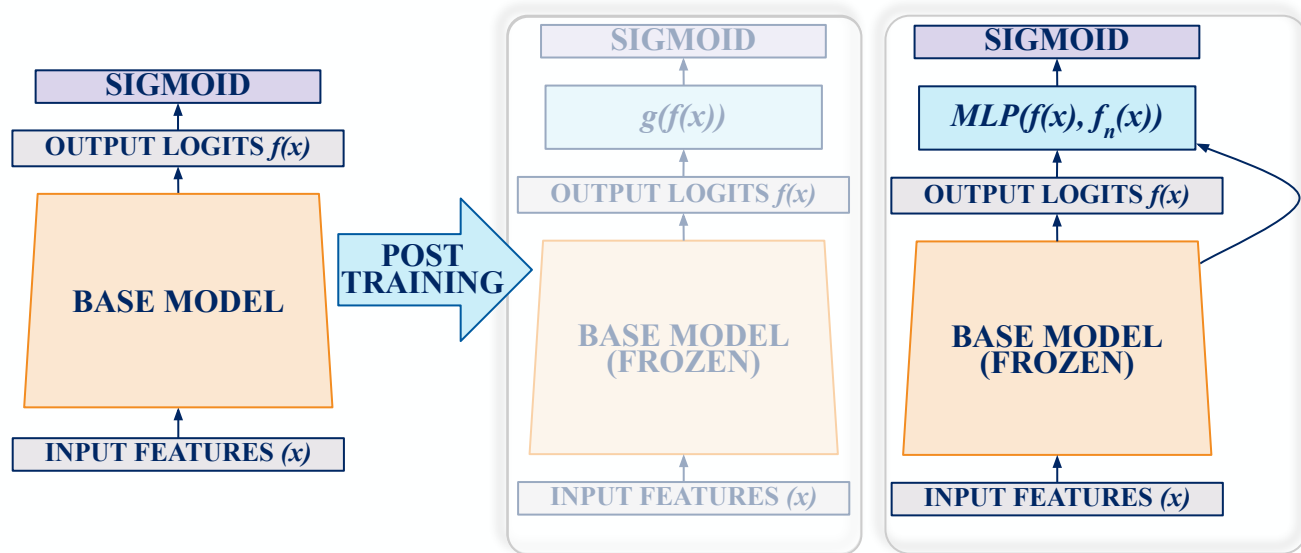


“Traditional” post-hoc methods:

- Platt scaling^[6]
- Isotonic Regression^[7]
- G-layers^[8]
- ...

They only transform models' outputs

Post-hoc Calibration vs Fine-tuning



“Traditional” post-hoc methods:

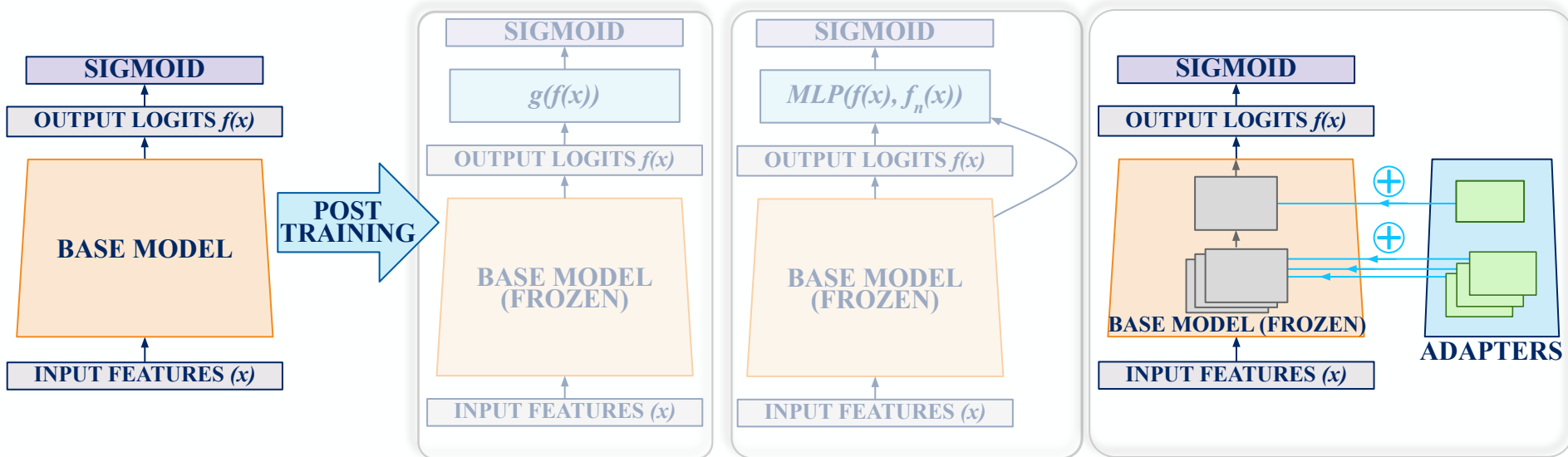
- Platt scaling^[6]
- Isotonic Regression^[7]
- G-layers^[8]
- ...

They only transform models’ outputs

“Generalized” g-layers:

MLP now also hooks into base model’s classification head

Post-hoc Calibration vs Fine-tuning



“Traditional” post-hoc methods:

- Platt scaling^[6]
- Isotonic Regression^[7]
- G-layers^[8]
- ...

They only transform models’ outputs

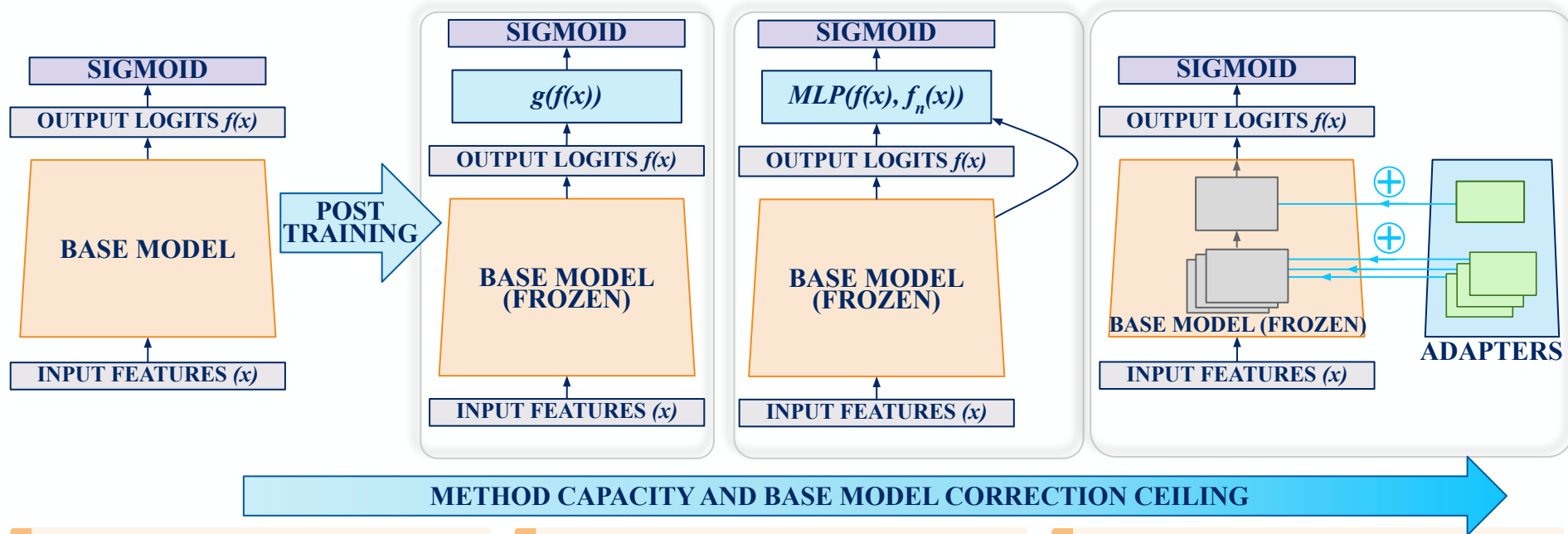
“Generalized” g-layers:

MLP now also hooks into base model’s classification head

Adapter-based model modification:

Low-rank adapters^[9] (LoRA) are injected into base model’s linear transformations’ matrices

Post-hoc Calibration vs Fine-tuning



“Traditional” post-hoc methods:

- Platt scaling^[6]
- Isotonic Regression^[7]
- G-layers^[8]
- ...

They only transform models’ outputs

“Generalized” g-layers:

MLP now also hooks into base model’s classification head

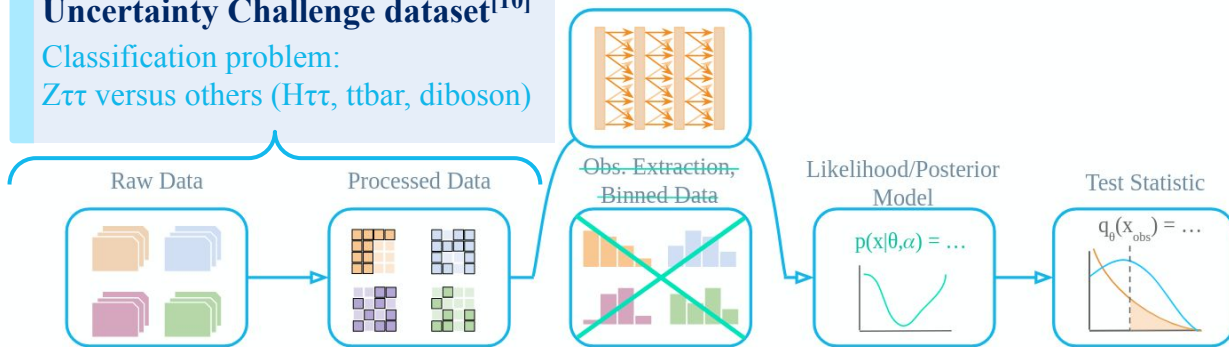
Adapter-based model modification:

Low-rank adapters^[9] (LoRA) are injected into base model’s linear transformations’ matrices

Experiment setup

We used FAIR Universe HiggsML
Uncertainty Challenge dataset^[10]

Classification problem:
 $Z\tau\tau$ versus others ($H\tau\tau$, $t\bar{t}$, diboson)



Experiment setup

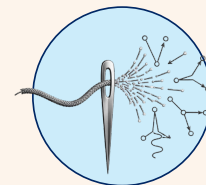
We used FAIR Universe HiggsML
Uncertainty Challenge dataset^[10]

Classification problem:
 $Z\tau\tau$ versus others ($H\tau\tau$, $t\bar{t}$, diboson)

Base model setups:

- Permutationally equivariant transformer-based models w.r.t particle ordering
- Use only kinematic features as inputs

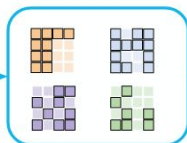
We run everything fully automated by needle-sbi!



Raw Data



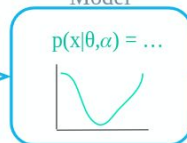
Processed Data



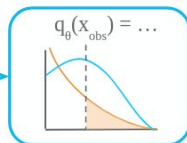
Obs. Extraction,
Binned Data



Likelihood/Posterior
Model



Test Statistic



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

Classifier quality: **Brier Score**

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

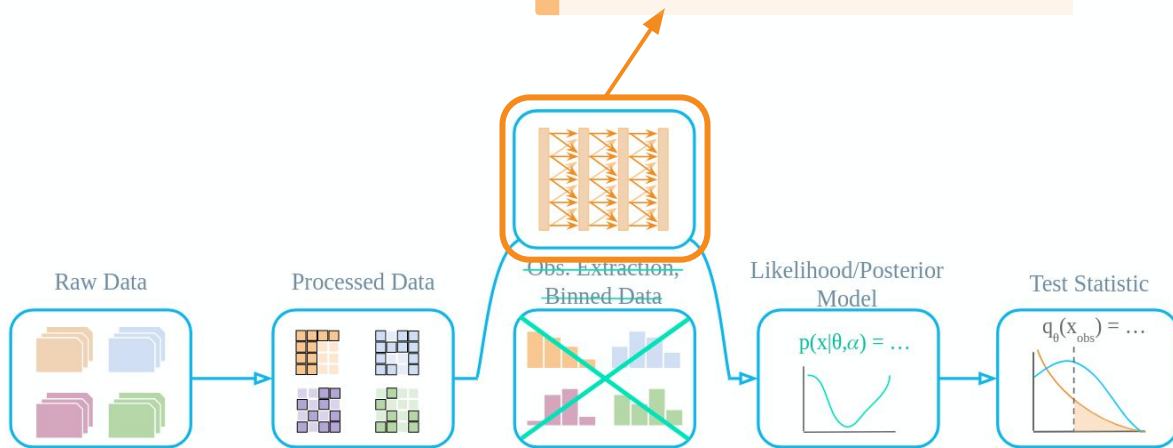
Classifier quality: Brier Score

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$

Calibration quality:

(Adaptive/Averaged) Expected calibration error

$$ECE = \frac{\sum_{b \in \text{bins}} w_b |y_b^{\text{total}} - s_b^{\text{total}}|}{\sum_{b \in \text{bins}} w_b}$$



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

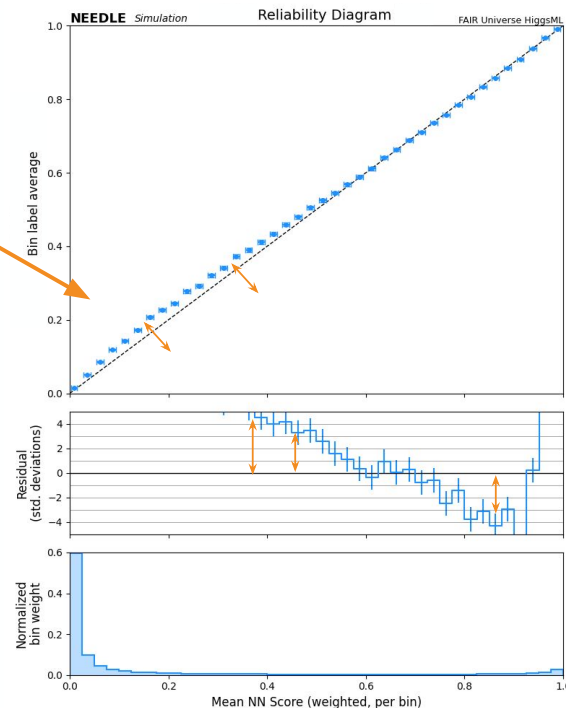
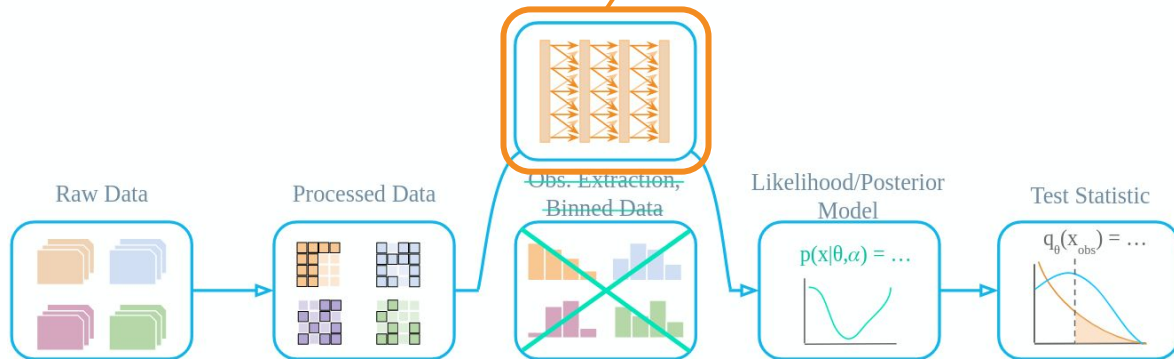
Classifier quality: Brier Score

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$

Calibration quality:

(Adaptive/Averaged) Expected calibration error

$$ECE = \frac{\sum_{b \in \text{bins}} w_b |y_b^{\text{total}} - s_b^{\text{total}}|}{\sum_{b \in \text{bins}} w_b}$$



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

Classifier quality: Brier Score

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$

Calibration quality:

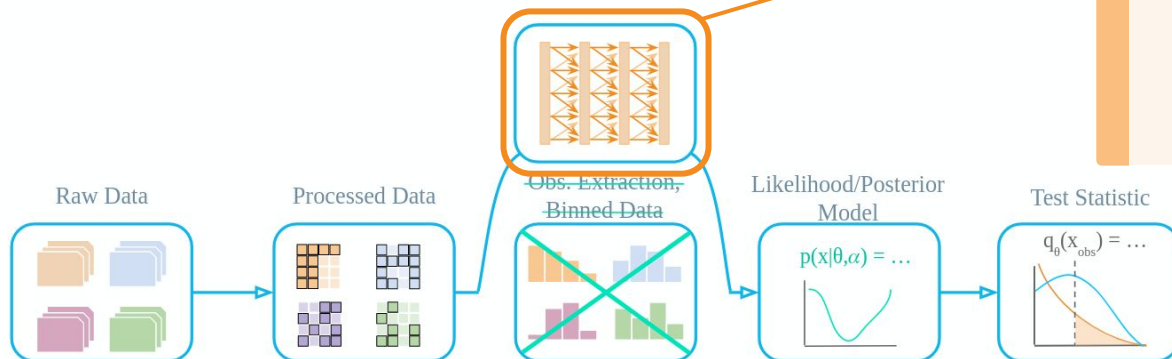
(Adaptive/Averaged) Expected calibration error

$$ECE = \frac{\sum_{b \in \text{bins}} w_b |y_b^{\text{total}} - s_b^{\text{total}}|}{\sum_{b \in \text{bins}} w_b}$$

Ratio Estimation quality:

Reweighting Tests

For each feature: Wasserstein-1 distance (EMD) between distributions of the target and source after reweighting by the estimated prob. ratio



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

Classifier quality: Brier Score

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$

Calibration quality:

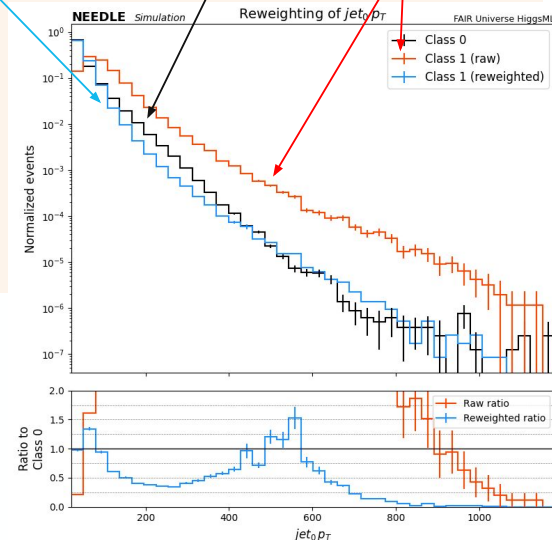
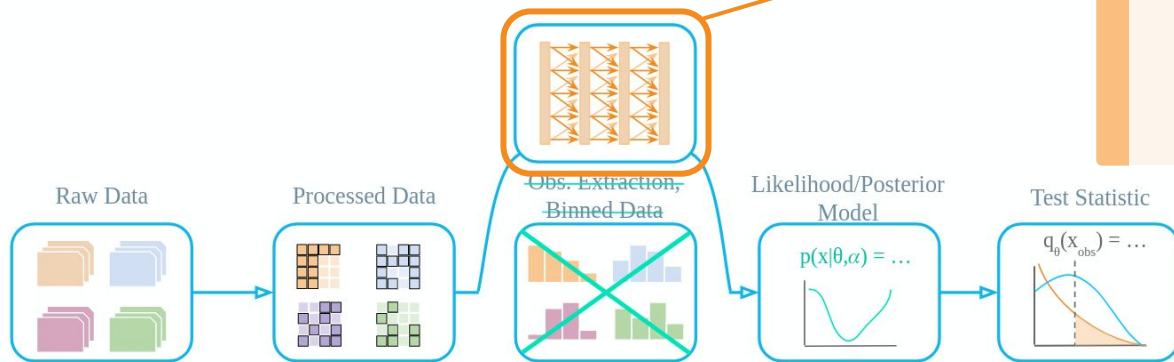
(Adaptive/Averaged) Expected calibration error

$$ECE = \frac{\sum_{b \in \text{bins}} w_b |y_b^{\text{total}} - s_b^{\text{total}}|}{\sum_{b \in \text{bins}} w_b}$$

Ratio Estimation quality:

Reweighting Tests

For each feature: Wasserstein-1 distance (EMD) between distributions of the **target** and **source** after **reweighting** by the estimated prob. ratio



Experiment setup

- We track a large number of relevant metrics^{*,[11]} for each NN and each calibration method:

Classifier quality: Brier Score

$$BS = \frac{1}{\sum_i w_i} \sum_i w_i (s(x_i) - y_i)^2$$

Calibration quality:

(Adaptive/Averaged) Expected calibration error

$$ECE = \frac{\sum_{b \in \text{bins}} w_b |y_b^{\text{total}} - s_b^{\text{total}}|}{\sum_{b \in \text{bins}} w_b}$$

Ratio Estimation quality:

Reweighting Tests

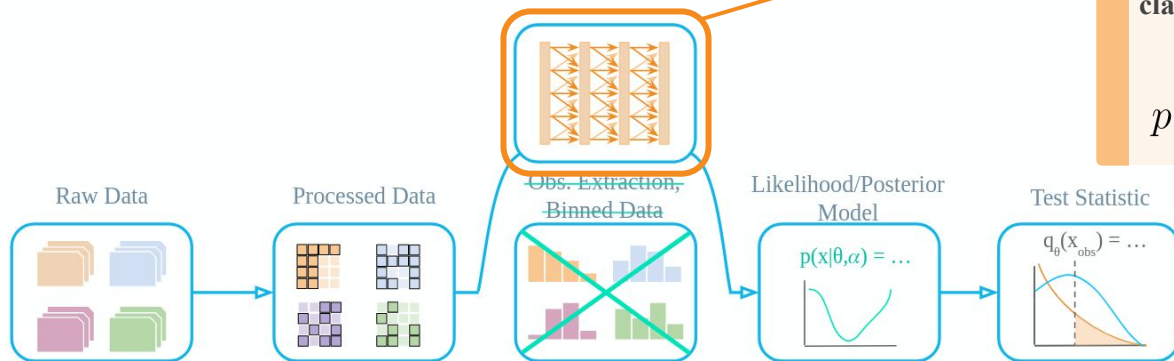
For each feature: Wasserstein-1 distance (EMD) between distributions of the target and source after reweighting by the estimated prob. ratio

Asimov Fit Tests:

Simple 1-parameter signal injection test between 2 classes:

$$\nu_{\text{total}} = c \cdot \nu_1 + \nu_0$$

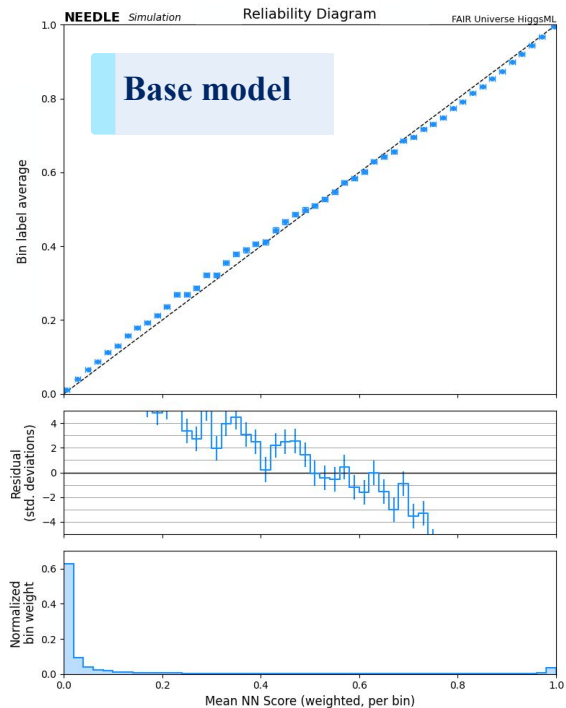
$$p(x|c) = c \cdot p(x|y = 1) + p(x|y = 0)$$



Results: calibration versus asimov fits

We compare: **base model** versus several calibrated/fine tuned variants

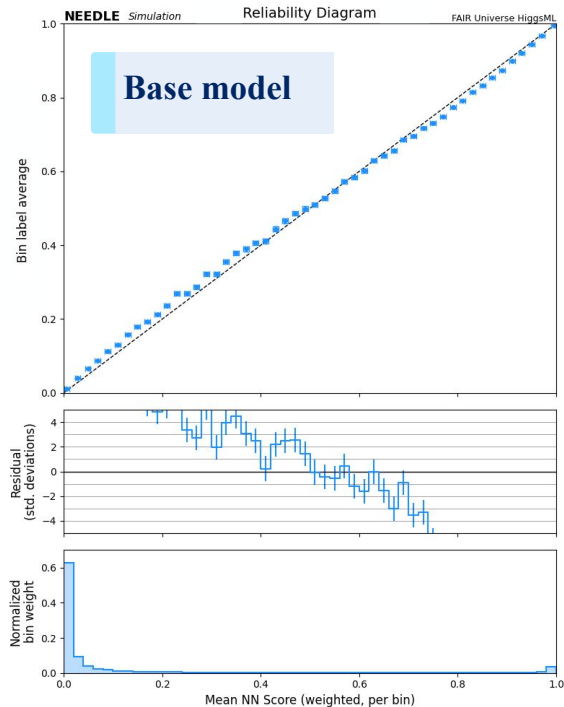
Shown are aggregated results over ensembles x folds



Results: calibration versus asimov fits

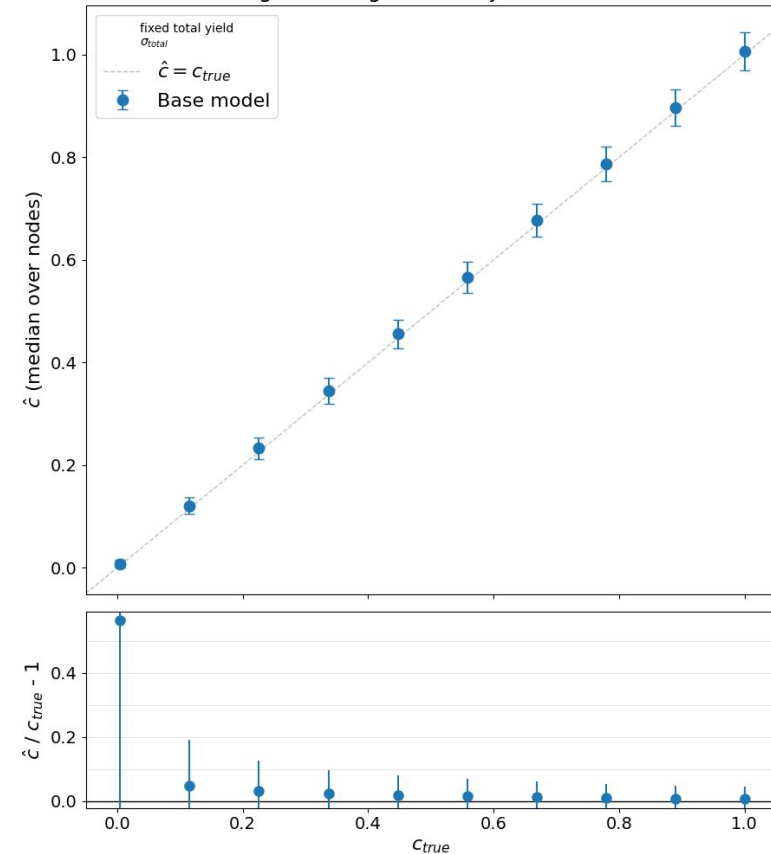
We compare: **base model** versus several calibrated/fine tuned variants

Shown are aggregated results over ensembles x folds



NEEDLE Simulation
Fair Universe HiggsML

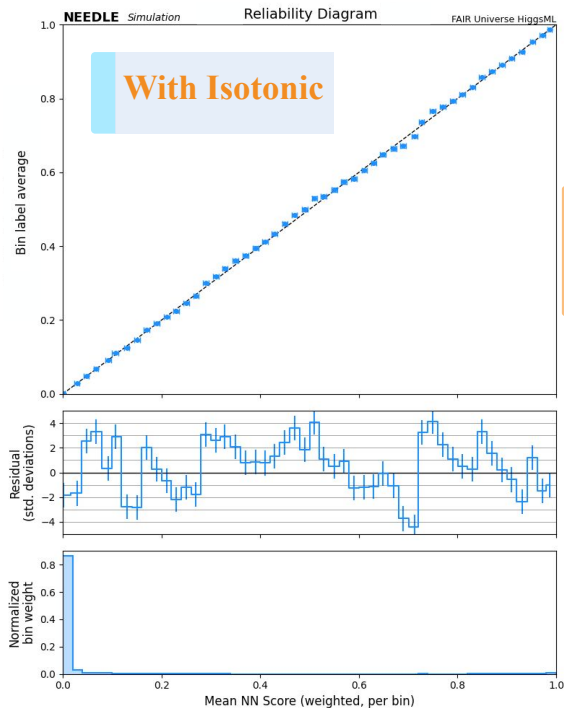
Signal strength fit vs injected value



Results: calibration versus asimov fits

We compare: **base model** versus **several calibrated/fine tuned variants**

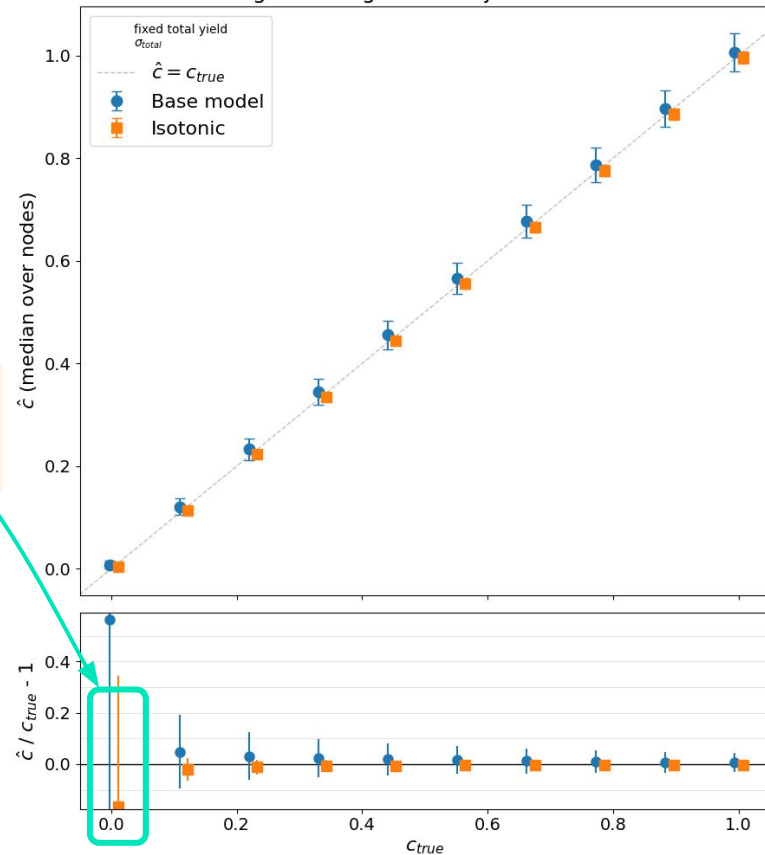
Shown are aggregated results over ensembles x folds



~17% **residue** at the lowest signal strength
Can this residue be eliminated?

NEEDLE Simulation
Fair Universe HiggsML

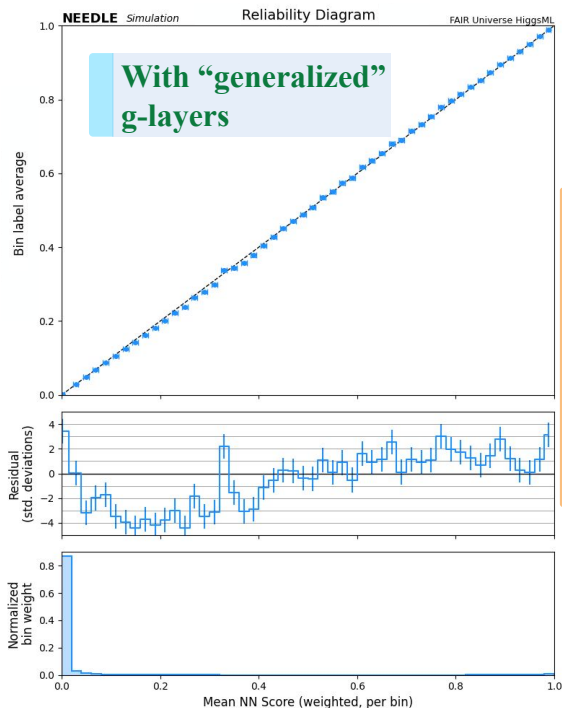
Signal strength fit vs injected value



Results: calibration versus asimov fits

We compare: **base model** versus **several calibrated/fine tuned variants**

Shown are aggregated results over **ensembles x folds**

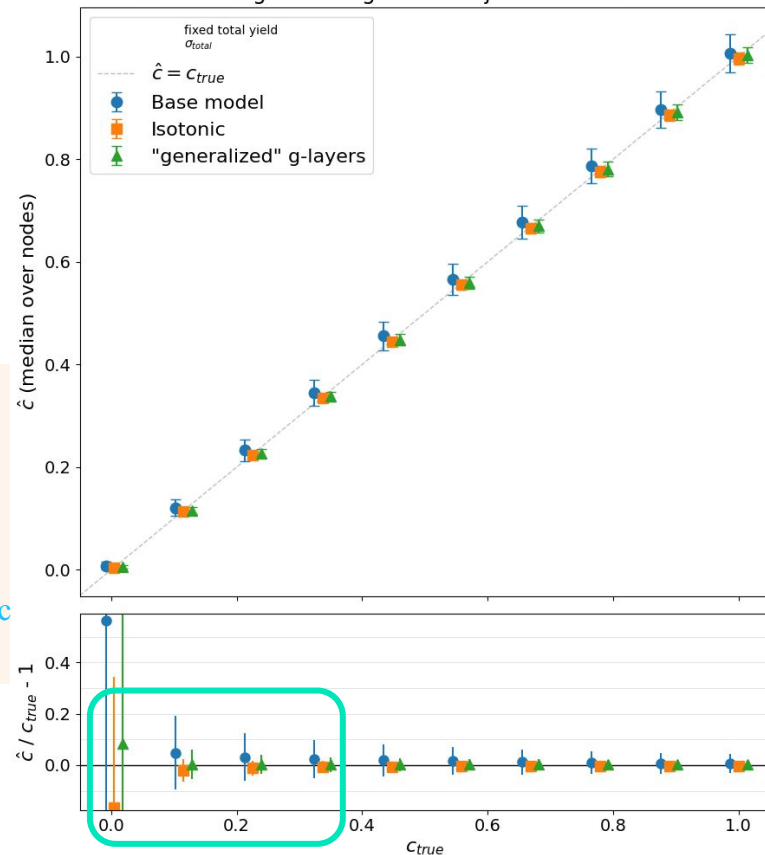


~17% residue at the lowest signal strength (with Isotonic)
Can this residue be eliminated?

For low values of injected signal, the fit model becomes sensitive to the high-values of the probability ratio, which is exactly the place where post-hoc calibration won't help

NEEDLE Simulation
Fair Universe HiggsML

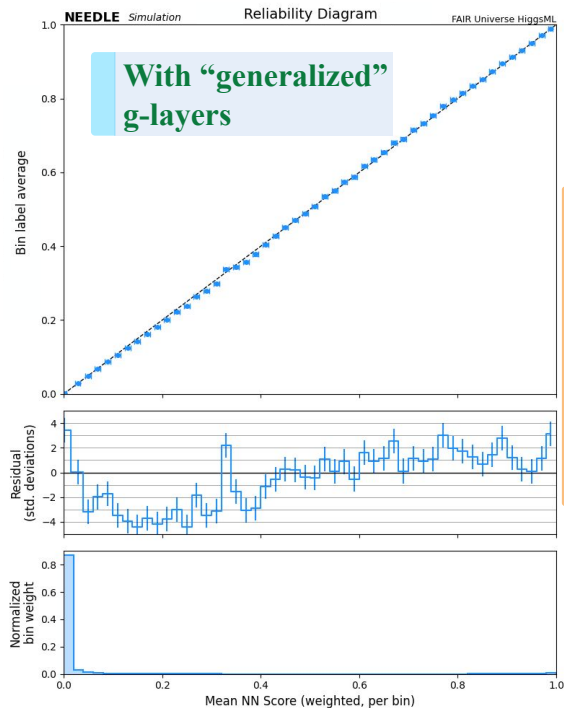
Signal strength fit vs injected value



Results: calibration versus asimov fits

We compare: base model versus several calibrated/fine tuned variants

Shown are aggregated results over ensembles x folds

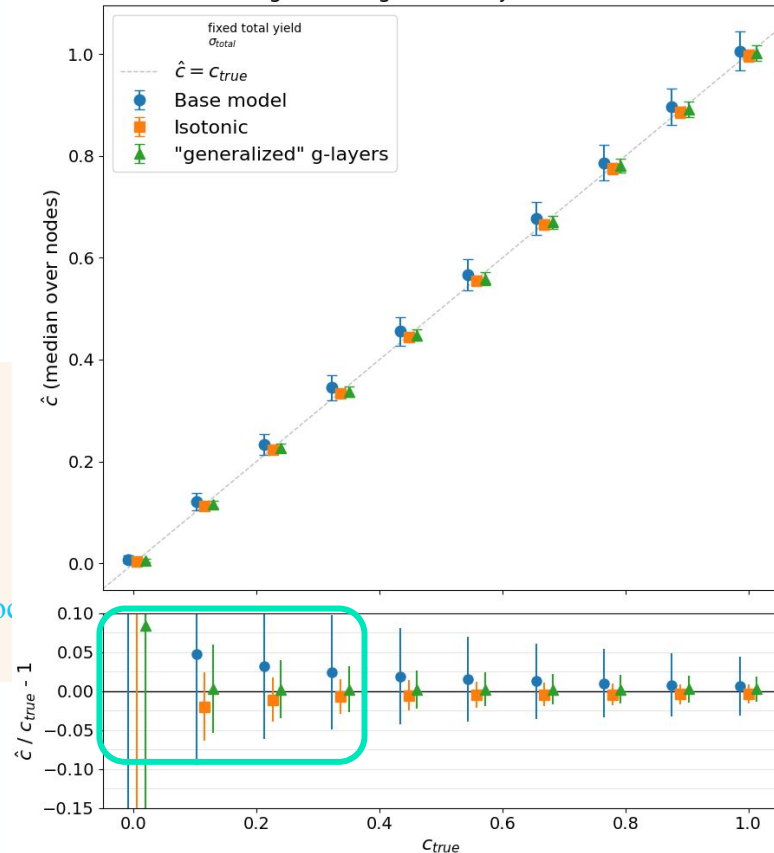


~17% residue at the lowest signal strength (with Isotonic)
Can this residue be eliminated?

For low values of injected signal, the fit model becomes sensitive to the high-values of the probability ratio, which is exactly the place where post-hoc calibration won't help

NEEDLE Simulation
Fair Universe HiggsML

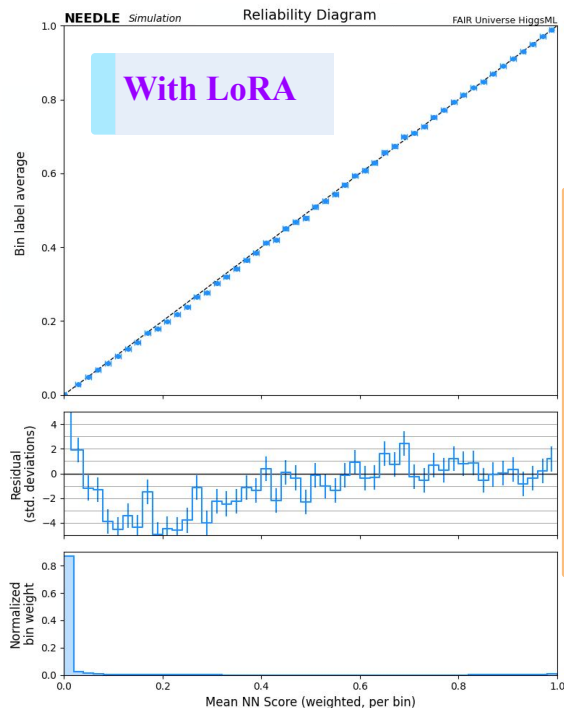
Signal strength fit vs injected value



Results: calibration versus asimov fits

We compare: base model versus several calibrated/fine tuned variants

Shown are aggregated results over ensembles x folds



~17% residue at the lowest signal strength (with Isotonic)

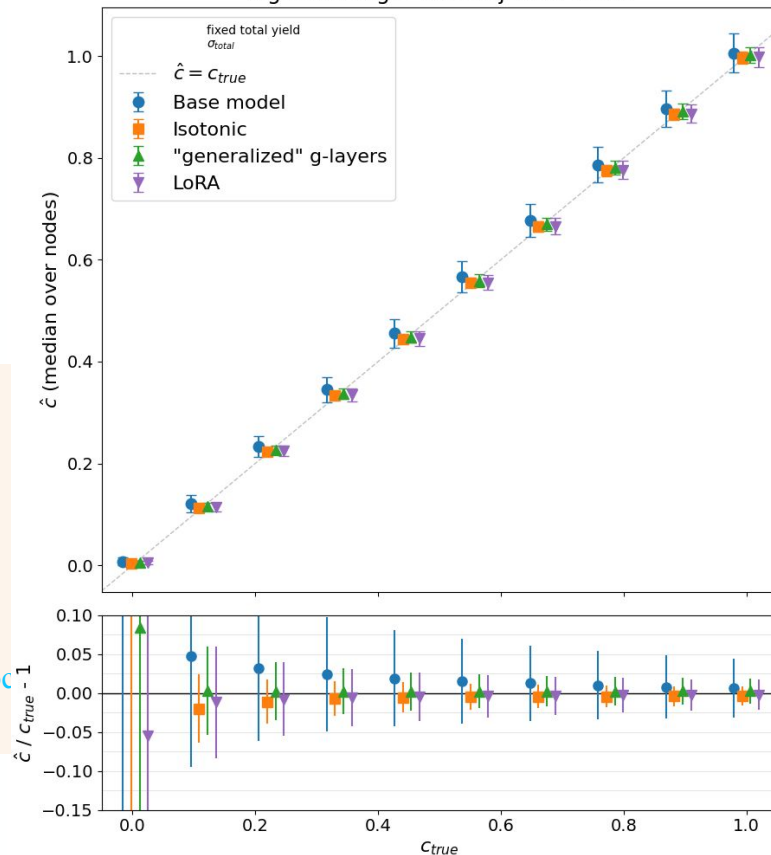
Can this residue be eliminated?

Yes. And at no cost to the calibration performance

For low values of injected signal, the fit model becomes sensitive to the high-values of the probability ratio, which is exactly the place where post-hoc calibration won't help

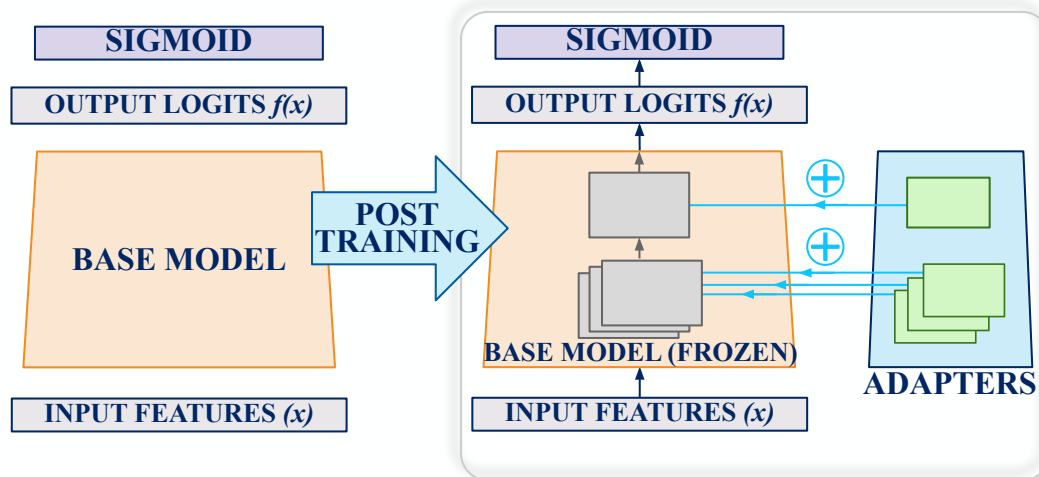
NEEDLE Simulation
Fair Universe HiggsML

Signal strength fit vs injected value



In-depth exploration of the adapter-based method

If this method can improve models after training, can it also make general classifier training more practical?

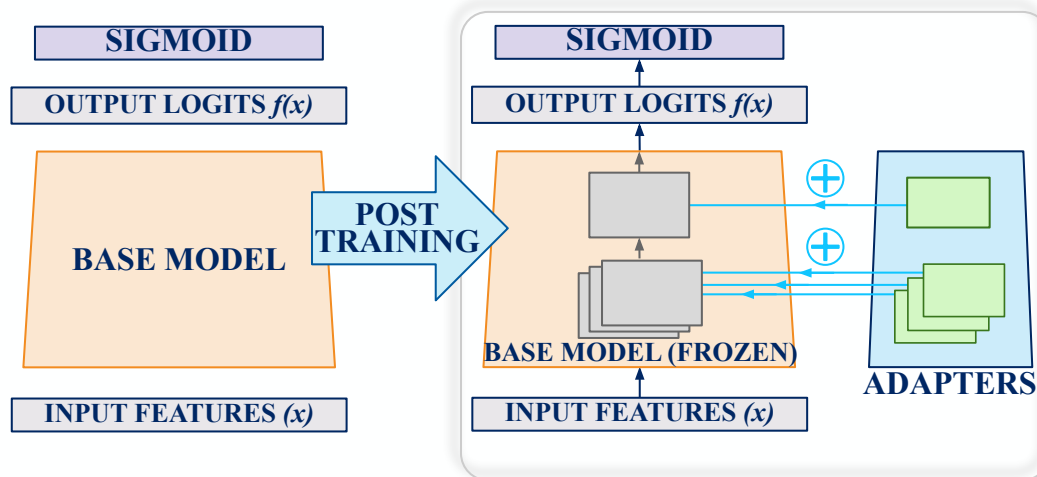


In-depth exploration of the adapter-based method

If this method can improve models after training, can it also make general classifier training more practical?

Test setup:

- Base models trained on different sized datasets, keeping their architecture constant
- For each NN, a set of adapters was fitted to it
 - We kept the adapter architecture and total training cost constant to better control the test
 - The total dataset size adapters were allowed to train was gradually increased
- As control we calibrated each model with isotonic regression
 - Isotonic regression was the best performing of the post-hoc methods we explored
 - The dataset size isotonic was fitted with is the same as the lowest-dataset size adapter
- The plots that follow show the results of aggregated metrics across all NNs (over folds x ensemble members)

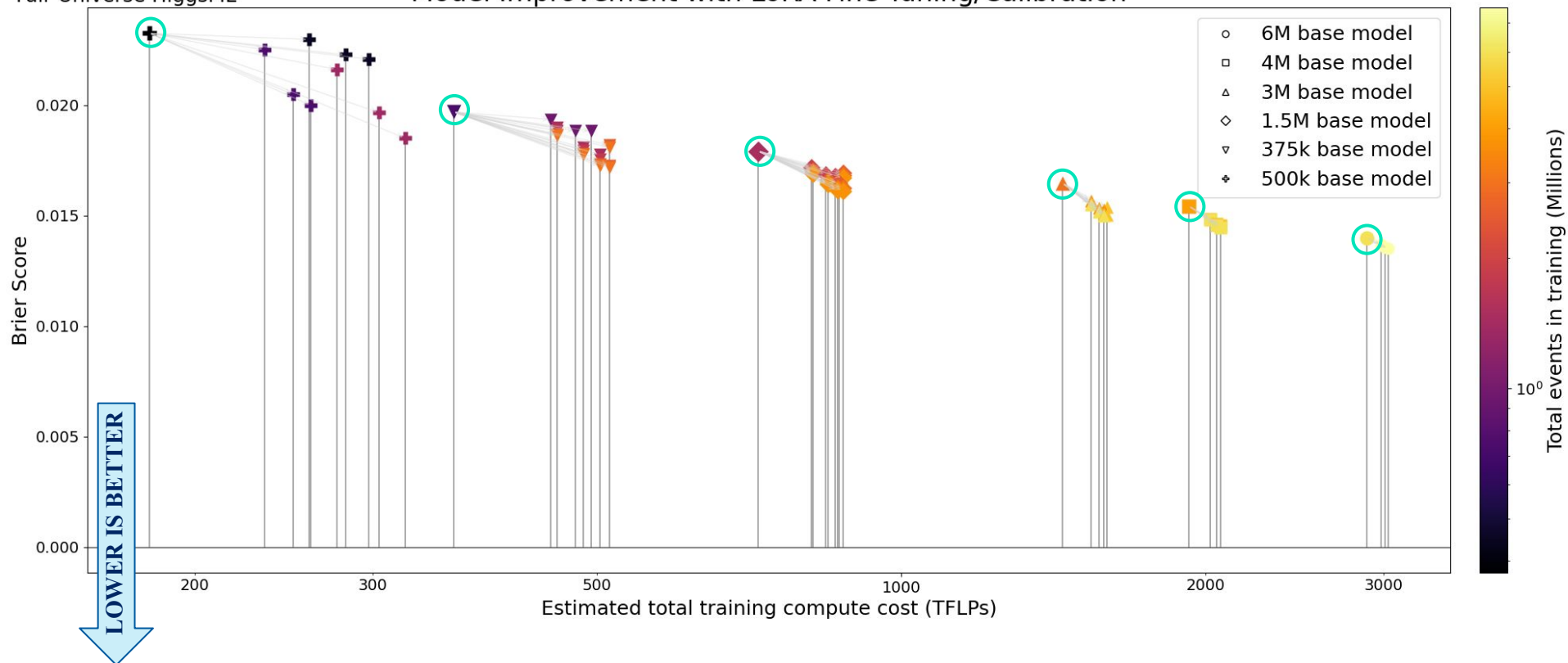


Results: base models versus calibrated/refined variants (Brier score)

True x-axis positions

NEEDLE Simulation
Fair Universe HiggsML

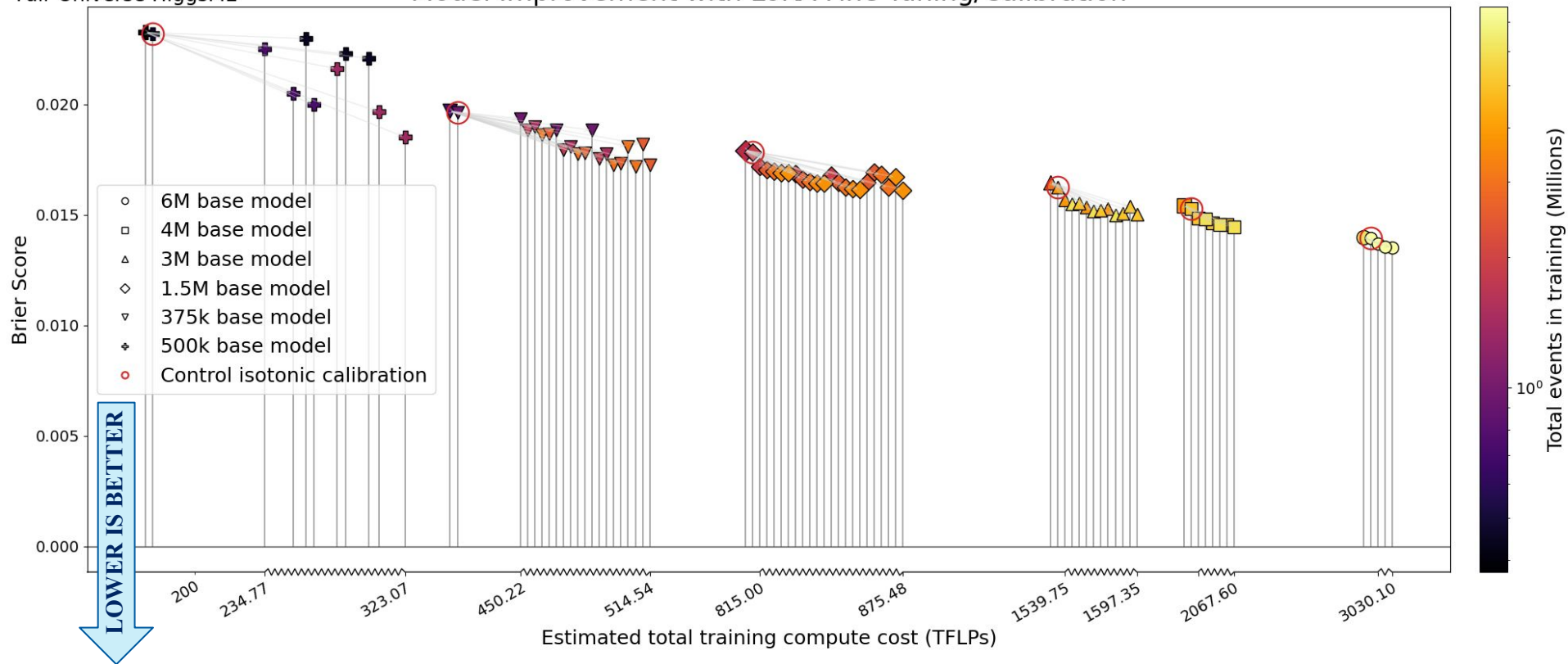
Model Improvement with LoRA Fine Tuning/Calibration



Results: base models versus calibrated/refined variants (Brier score)

NEEDLE Simulation
Fair Universe HiggsML

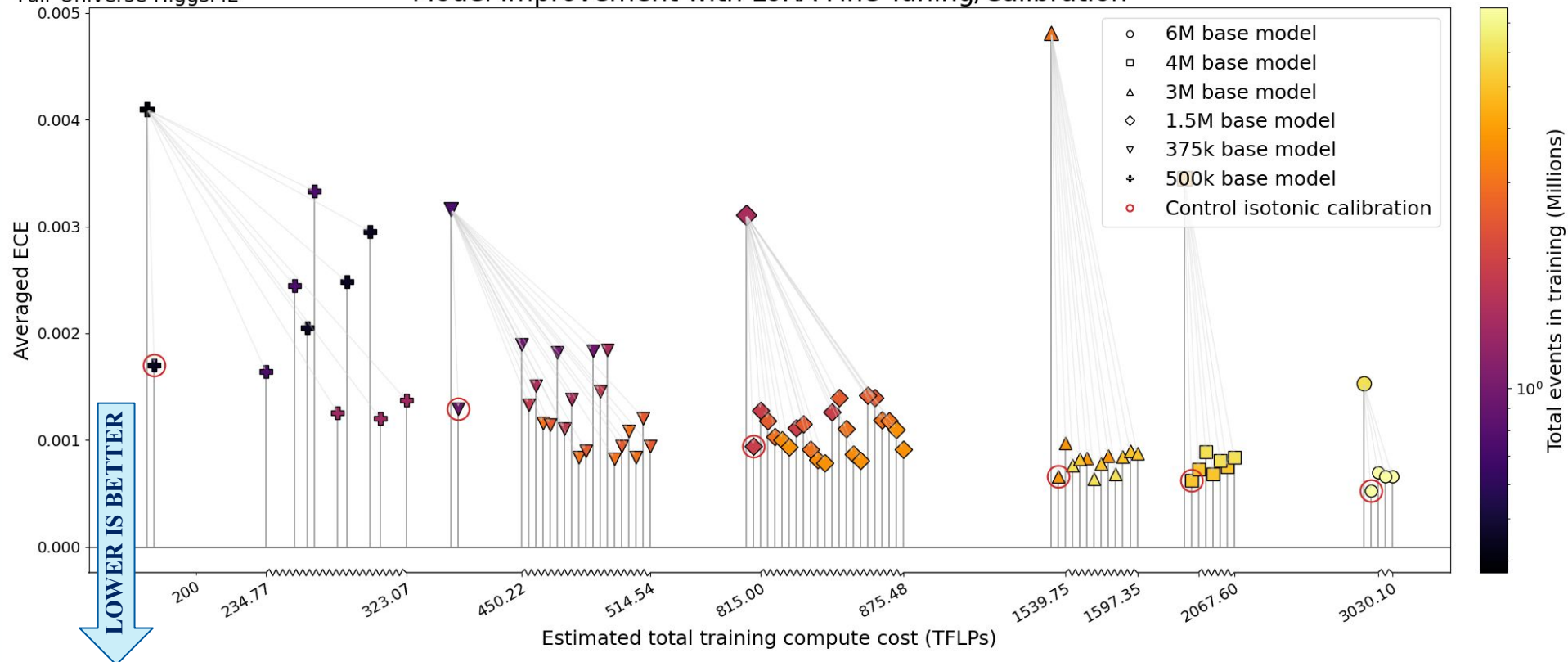
Model Improvement with LoRA Fine Tuning/Calibration



Results: base models versus calibrated/refined variants (AECE)

NEEDLE Simulation
Fair Universe HiggsML

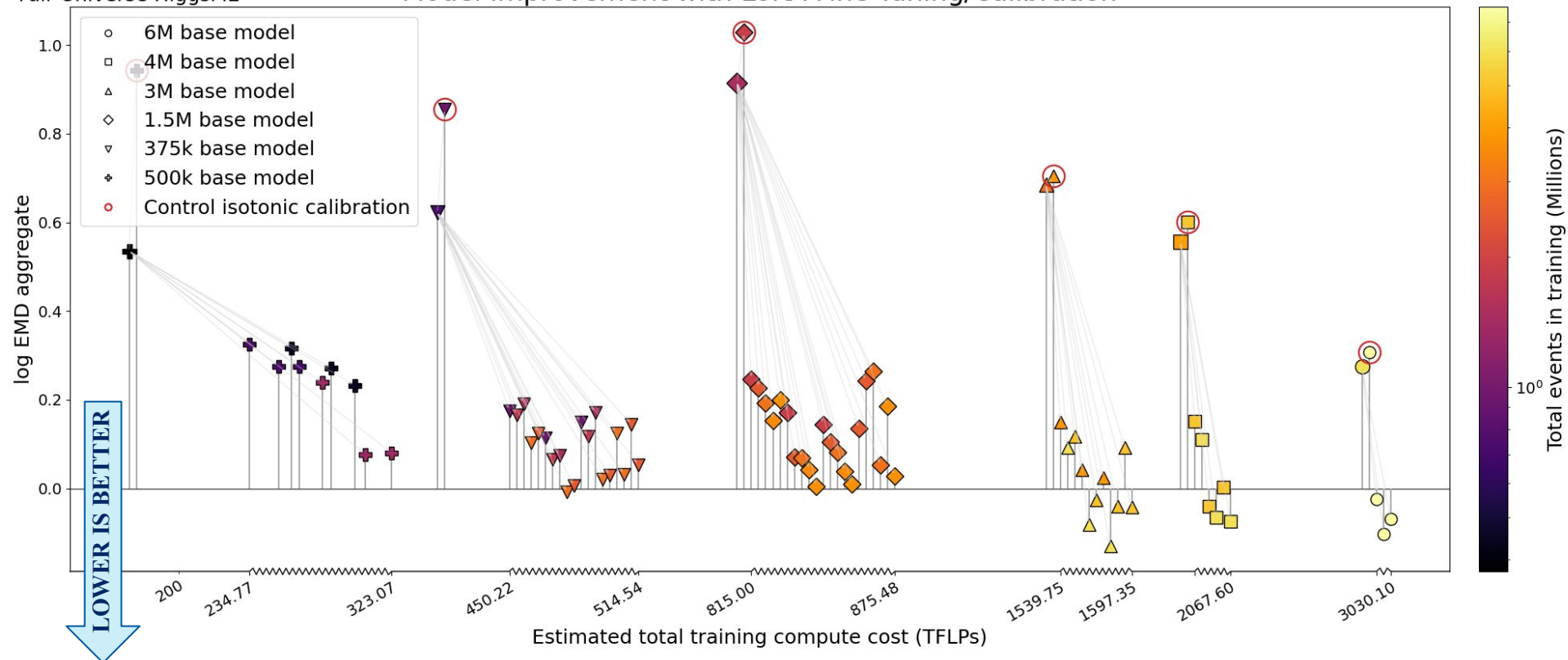
Model Improvement with LoRA Fine Tuning/Calibration



Results: base models versus calibrated/refined variants (reweighting EMD aggregates, minority→majority class)

NEEDLE Simulation
Fair Universe HiggsML

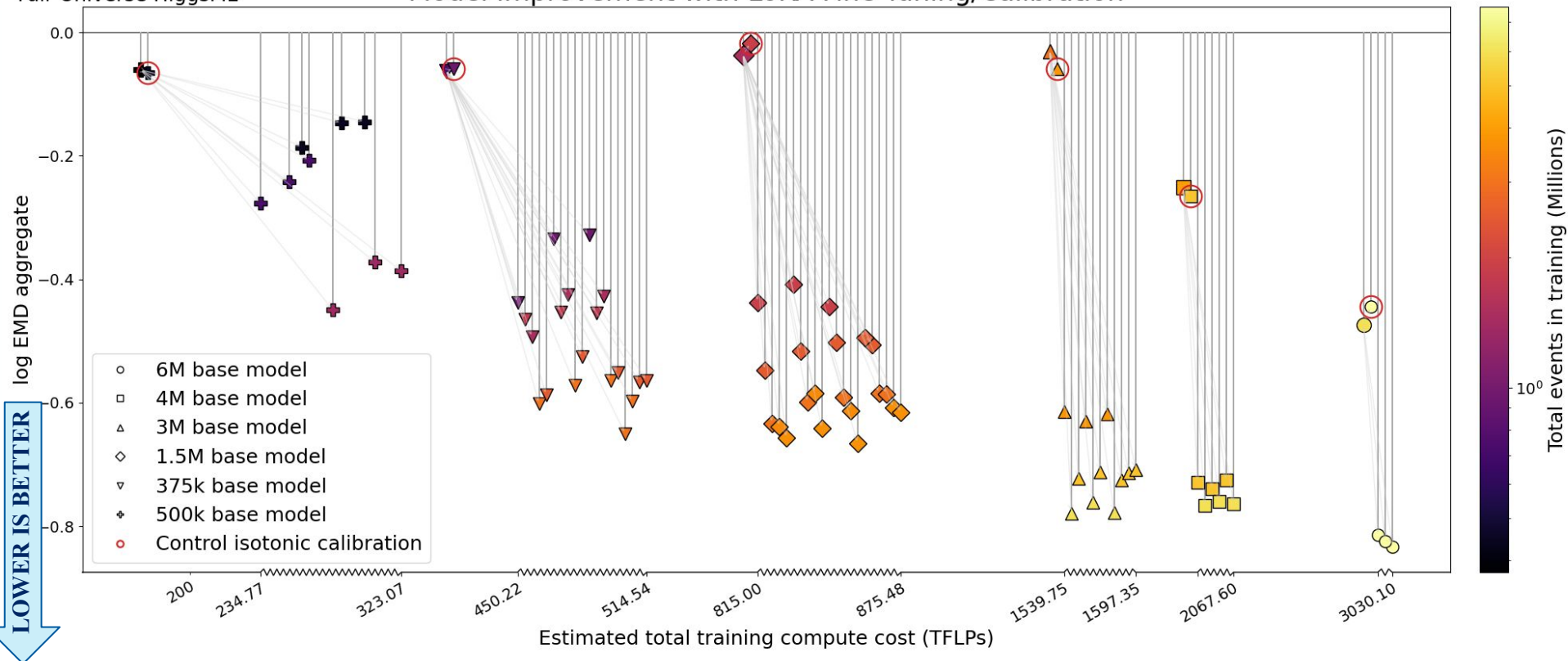
Model Improvement with LoRA Fine Tuning/Calibration



Results: base models versus calibrated/refined variants (reweighting EMD aggregates, majority→minority class)

NEEDLE Simulation
Fair Universe HiggsML

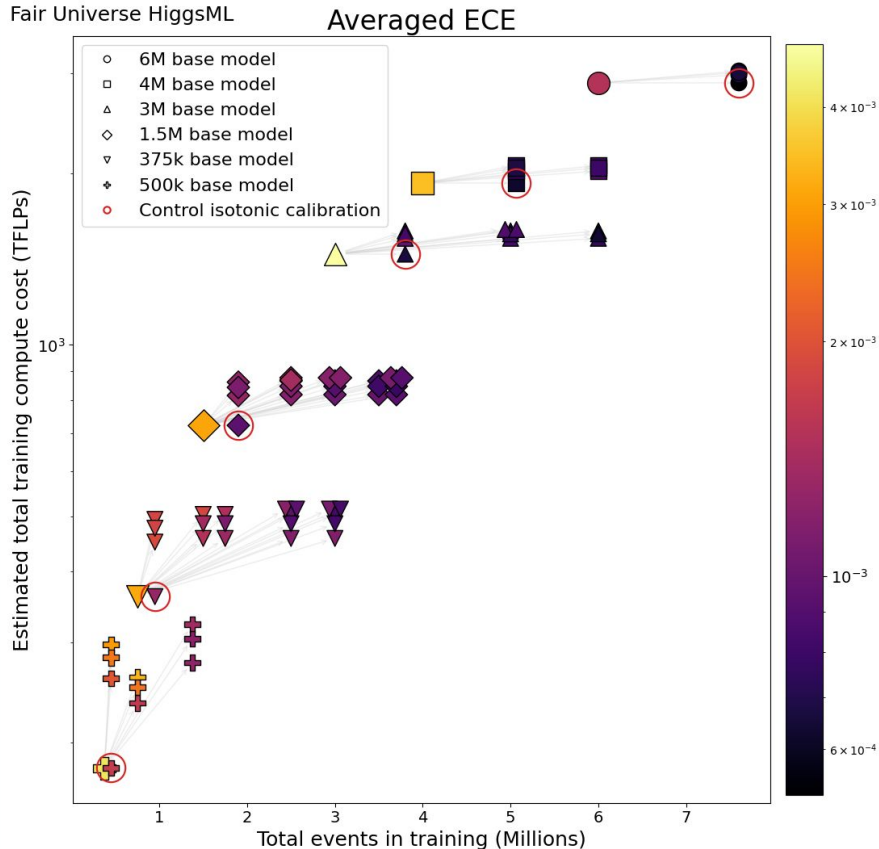
Model Improvement with LoRA Fine Tuning/Calibration



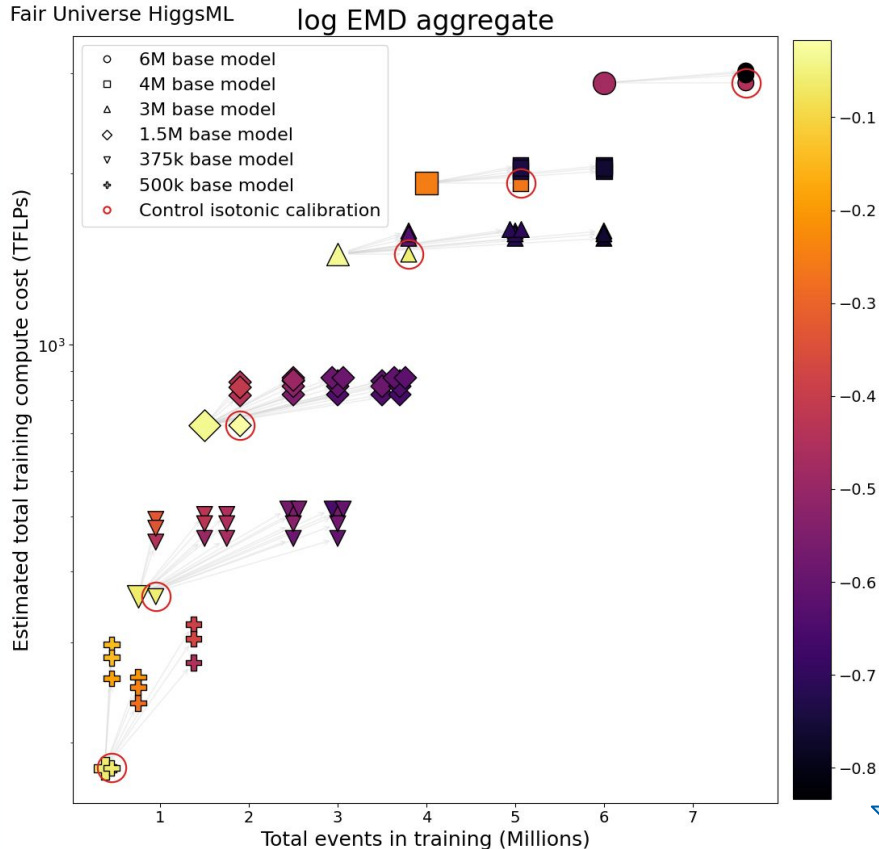
Results: base models versus calibrated/refined variants

A better view of the comp. costs and training events

NEEDLE Simulation
Fair Universe HiggsML



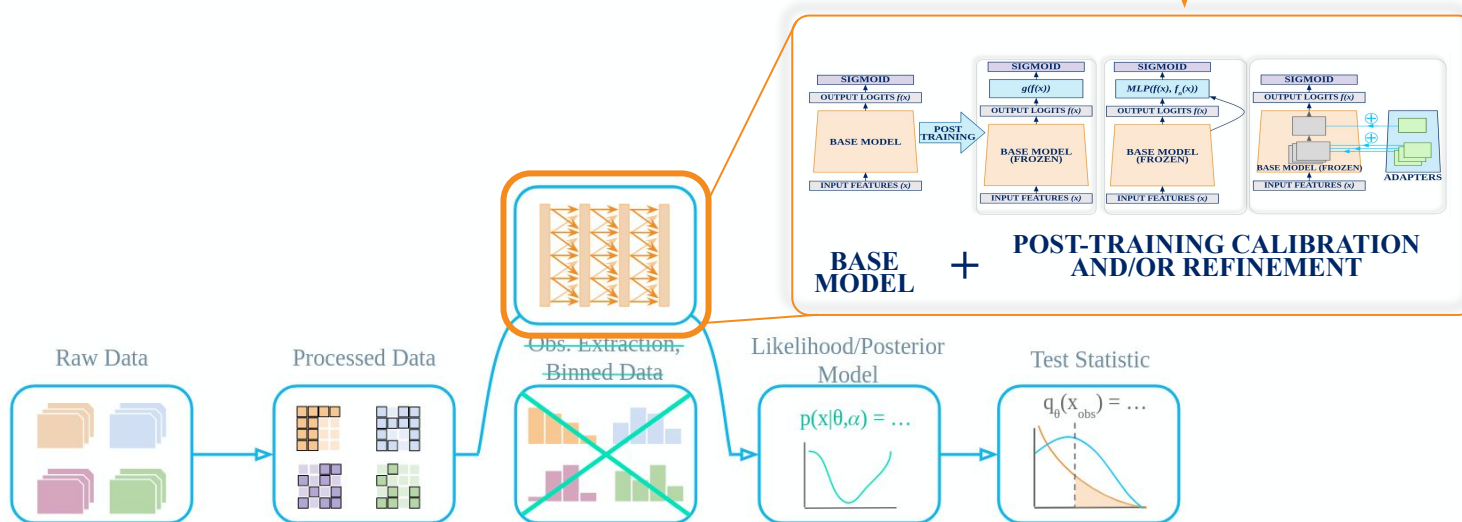
NEEDLE Simulation
Fair Universe HiggsML



“DARKER” IS BETTER

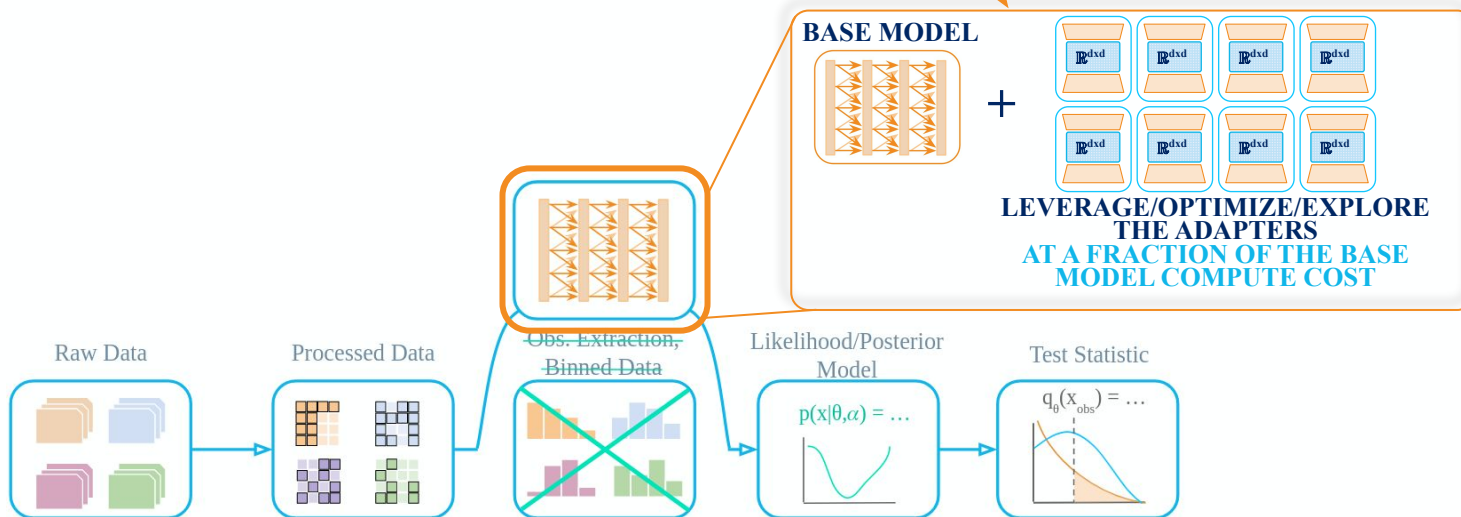
Conclusions

- Obtaining high-quality, reliable probability ratios can require **more than** just removing the **calibration error**
 - Post-hoc methods cannot reduce grouping loss, unlike methods which have access to the base model
 - One way to determine the effect of grouping loss is to compare model improvement based on post-hoc and refinement methods



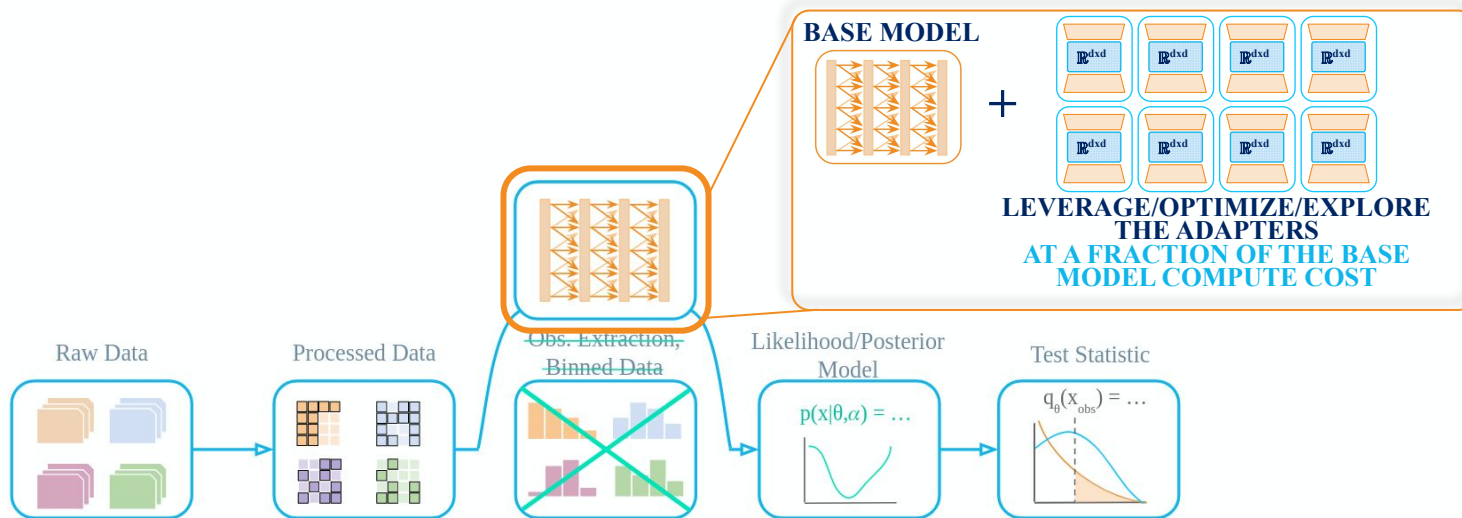
Conclusions

- **Obtaining high-quality, reliable probability ratios can require more than just removing the calibration error**
 - Post-hoc methods cannot reduce grouping loss, unlike methods which have access to the base model
 - One way to determine the effect of grouping loss is to compare model improvement based on post-hoc and refinement methods
- **Adapter-based model modification can also be used to improve a typical NLRE workflow**
 - Observed behaviour also strongly indicates base model training + fine-tuning performs better than one-shot model training



Conclusions

- **Obtaining high-quality, reliable probability ratios can require more than just removing the calibration error**
 - Post-hoc methods cannot reduce grouping loss, unlike methods which have access to the base model
 - One way to determine the effect of grouping loss is to compare model improvement based on post-hoc and refinement methods
- **Adapter-based model modification can also be used to improve a typical NLRE workflow**
 - Observed behaviour also strongly indicates base model training + fine-tuning performs better than one-shot model training
- **Probability ratio estimation also allows us to obtain likelihoods directly via DRDPM**
 - Fast and works even in cases when there's a lack of support between the classes



Stay tuned: many updates and a preprint coming out soon!

Calibration paper preprint

Refining Reliable Likelihood Ratio Estimators: From Calibration To Grouping Loss

Nino Kovačić*, Stephen Jiggins†, Judith Katzy‡

*University of Zagreb Faculty of Science,
nkovic.phy@pmf.hr

†Deutsches Elektronen-Synchrotron, DESY

needle-sbi CHEP 2026 in-proceedings

NEEDLE: A Columnar Workflow Orchestrator for Large-Scale Neural Simulation-Based Inference in HEP

Kylian Schmidt^{1*}, Felix Kahlhöfer¹, Judith Katzy², Levi Evans², Nicolò Trevisani¹, Nino Kovicic³, Stephen Jiggins², and Ulrich Husemann²

¹Karlsruhe Institute of Technology (KIT)

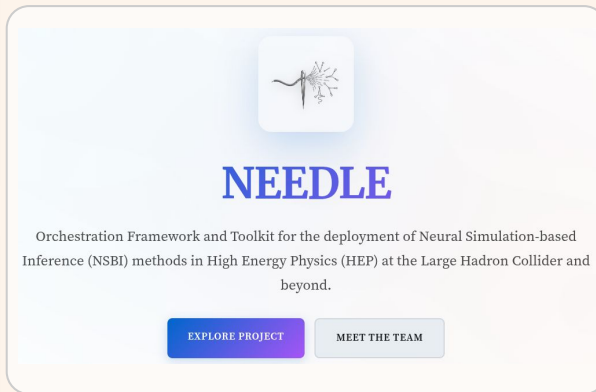
²Deutsches Elektronen-Synchrotron (DESY)

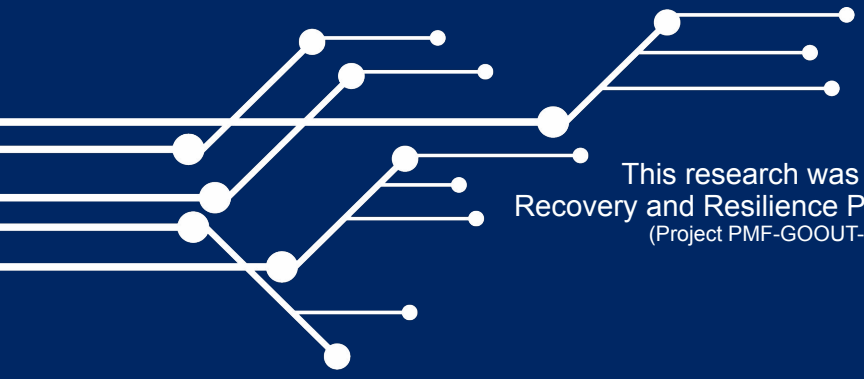
³University of Zagreb

Updates to needle-sbi public release in the pipeline:

- b2luigi and python API for needle-sbi
- NLRE Calibration and validation package
- DRPDM package

Check out our GitHub and Website!





Thank you for your attention!

Acknowledgement:
This research was supported by the European Union – NextGenerationEU through the National Recovery and Resilience Plan 2021–2026 Institutional grants of University of Zagreb Faculty of Science (Project PMF-GOOUT-2026-03 “Calibrated Density Ratio Estimation Approach to Diffusion Probabilistic Models as a Fast Neural Simulation-based Inferencing Method for Large Hadron Collider Physics” as well as RP3 “quark-gluon plasma research, hardware development at CERN and bioelectronics”)



References

- [1] For a general overview of NSBI see Cranmer et al. *The frontier of simulation-based inference*, [ArXiv link](#)
- [2] Cranmer et al. *Approximating Likelihood Ratios with Calibrated Discriminative Classifiers*, [ArXiv link](#)
- [3] Luigi Analysis Workflow, [ArXiv link](#)
- [4] batch 2 luigi, [docs link](#)
- [5] Perez-Lebel et al. *Beyond calibration: estimating the grouping loss of modern neural networks*, [ArXiv link](#)
- [6] Platt *Probabilistic Outputs for Support Vector Machines and Comparisons to Regularized Likelihood Methods*, [ResearchGate link](#)
- [7] Zadrozny and Elkan *Transforming Classifier Scores into Accurate Multiclass Probability Estimates*, [ResearchGate link](#), [ACM Digital Library link](#)
- [8] Rahimi et al. *Post-hoc Calibration of Neural Networks by g-Layers*, [ArXiv link](#)
- [9] Hu et al. LoRA: Low-Rank Adaptation of Large Language Models, [ArXiv link](#)
- [10] Benato et al. *FAIR Universe HiggsML Uncertainty Dataset and Competition*, [Codabench link](#), [ArXiv link](#)
- [11] For a general overview of classifier performance and calibration metrics see, e.g., Lane *A comprehensive review of classifier probability calibration metrics*, [ArXiv link](#)

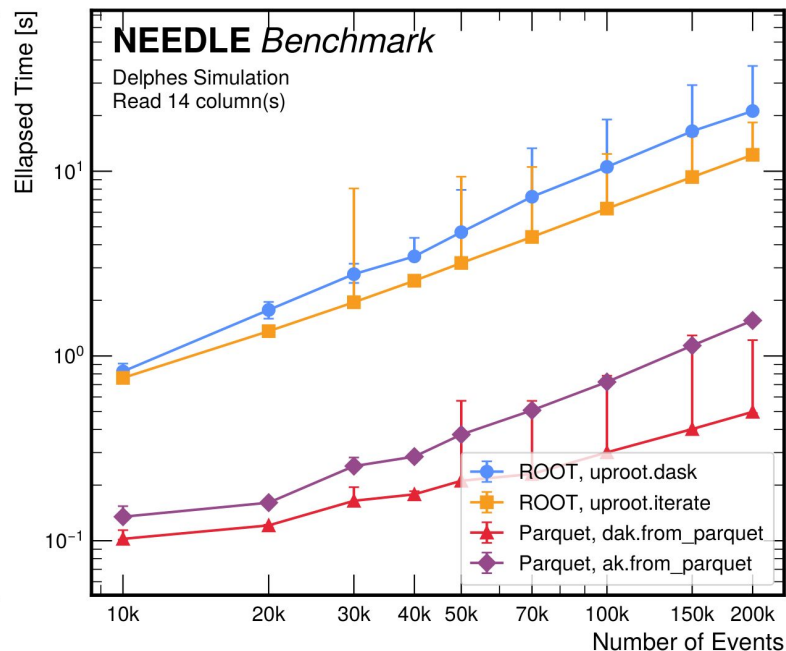
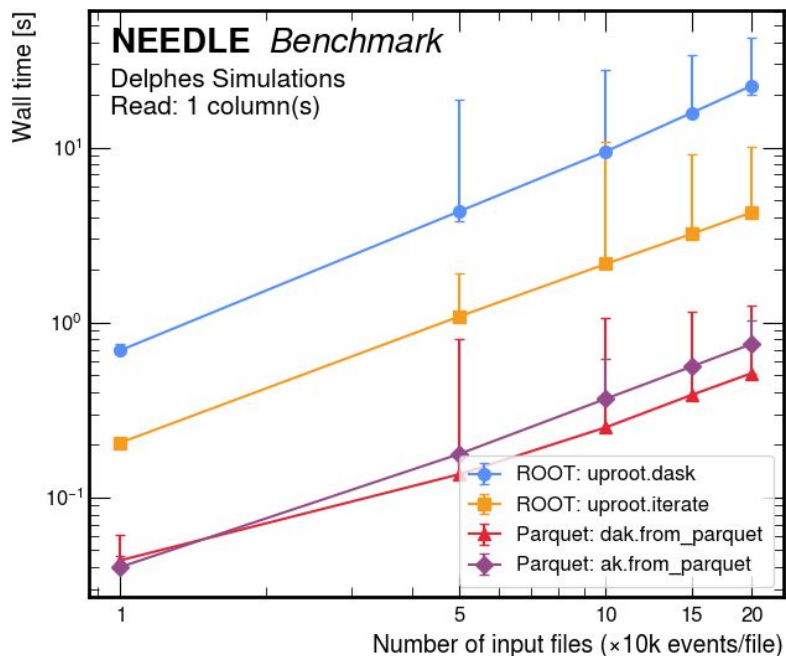


Backup

needle-sbi: data loading speed benchmark

needle-sbi data loading speed benchmarks:

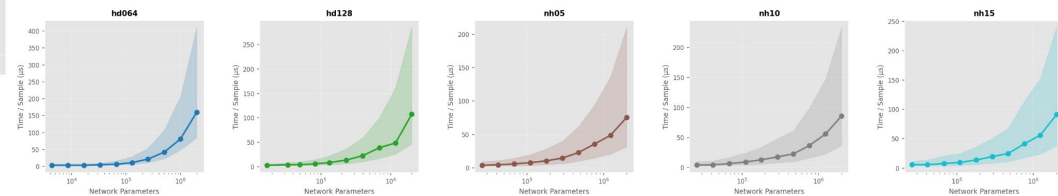
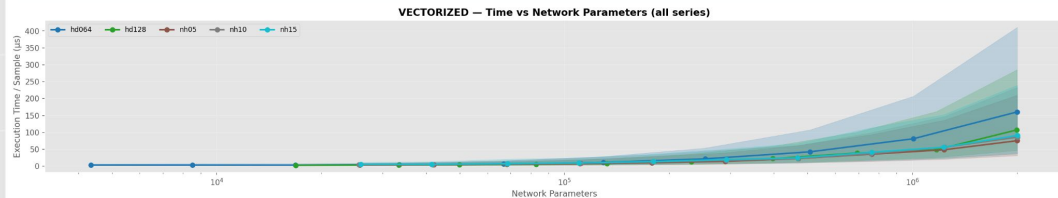
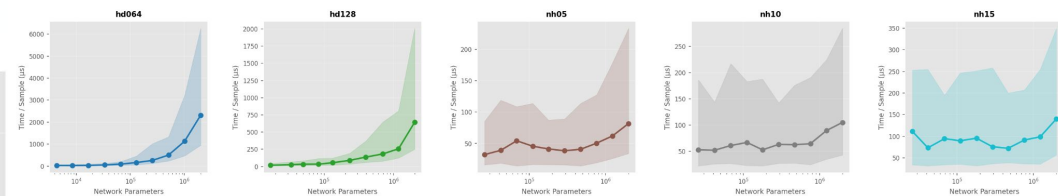
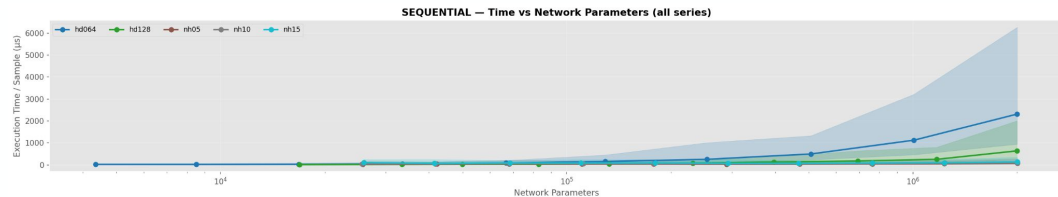
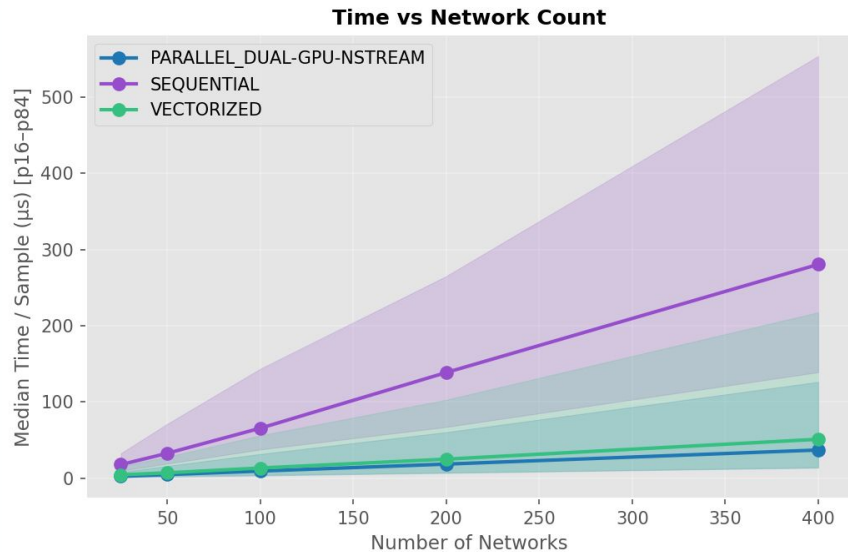
- Reading from ROOT files and from parquet files
- With two variants for parquet: one with and one without parsing in the array shape before data loading (ak-vs-dak*)



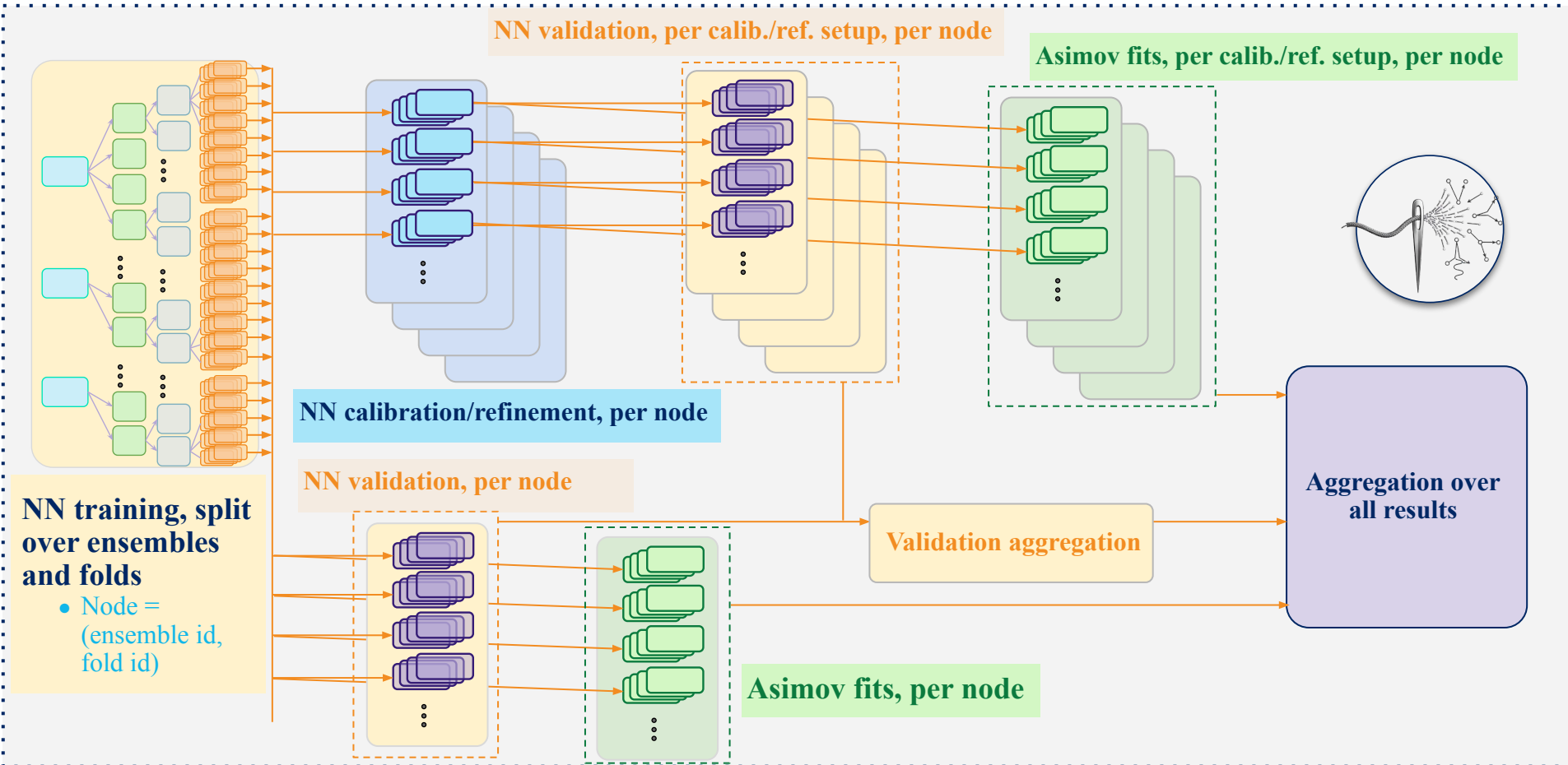
needle-sbi: parallel and vectorized model evaluation

needle-sbi NN evaluation benchmarks:

- Benchmarks performed bare metal on Tesla-V100-SXM2 (32GB)



DAG scheme* for NN training and calibration downstream tasks



*Technically what is shown in a sketch of a flow diagram, the DAG structure is inverted compared to this



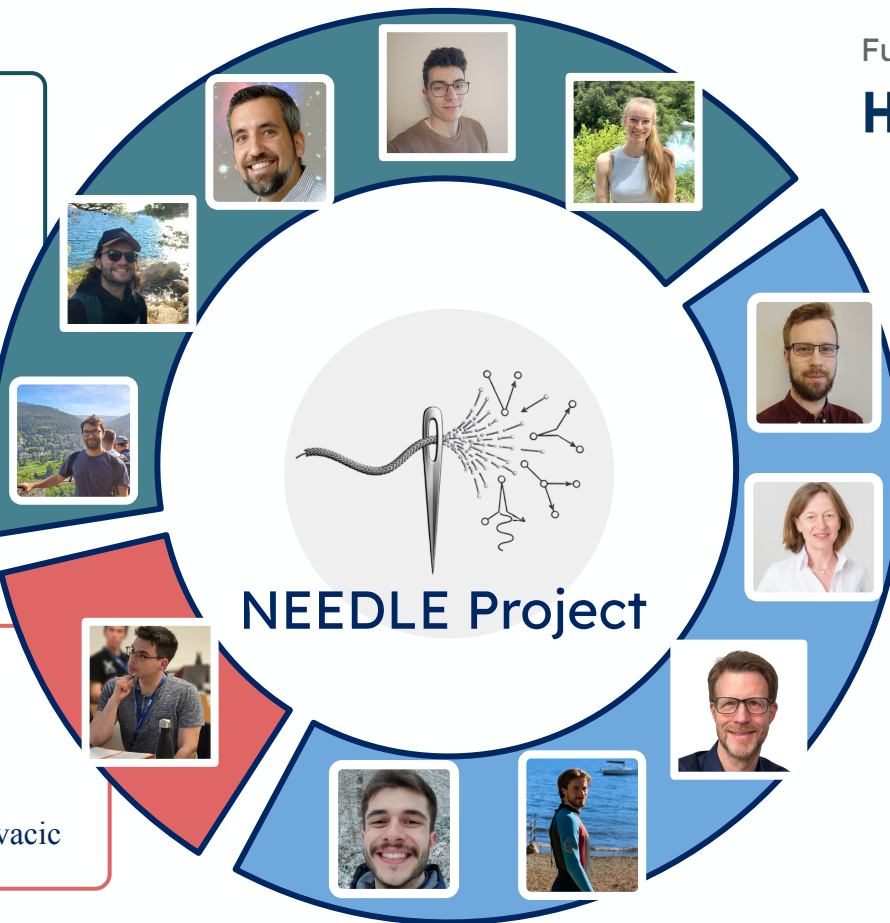
Lena Rathmann
Niklas Reus
Felix Kahlhöfer
Kylían Schmidt
Nicolò Trevisani



ALICE



Nino Kovacic



NEEDLE Project

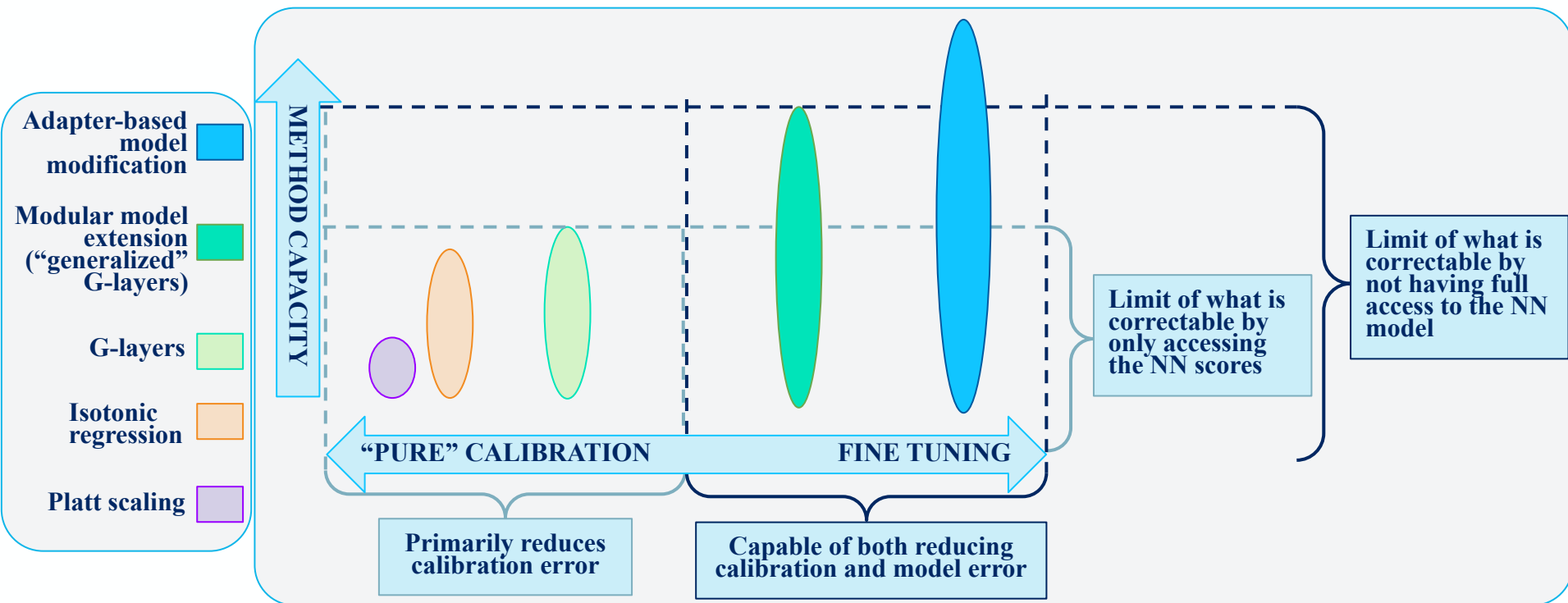
Funded by
HELMHOLTZAI



Stephen Jiggins
Judith Katzy
Ulrich Husemann
Levi Evans
Stefan Katsarov

The “Buffet” of Calibration (Model Improvement) Methods

Sketch of different methods’ hypothetical performance improvement ceilings and base model modification:



Grouping loss, example with Brier score

$$\begin{aligned}
 \mathbb{E}[(S-Y)^2] &= \mathbb{E}[(S-S^*) + (S^*-Y)]^2 = \mathbb{E}[(S-S^*)^2] + \mathbb{E}[(S^*-Y)^2] + 2\mathbb{E}[(S-S^*)(S^*-Y)] \\
 \mathbb{E}[(S^*-Y)^2] &= \mathbb{E}[\text{Var}(Y|X)] = \mathbb{E}[S^*(1-S^*)] \\
 \mathbb{E}[(S-S^*)(S^*-Y)] &= \mathbb{E}[(S-S^*) \underbrace{\mathbb{E}[S^*-Y|X]}_{=0}] = 0 \\
 &= \mathbb{E}[(S-S^*)^2] + \mathbb{E}[S^*(1-S^*)] \\
 \mathbb{E}[(S-S^*)^2] &= \mathbb{E}[(S-\bar{S})^2] + \mathbb{E}[(\bar{S}-S^*)^2] + 2\mathbb{E}[(S-\bar{S})(\bar{S}-S^*)] \\
 \mathbb{E}[(S-\bar{S})(\bar{S}-S^*)] &= \mathbb{E}\left[\mathbb{E}[(S-\bar{S})(\bar{S}-S^*)|S]\right] \\
 &= \mathbb{E}\left[(S-\bar{S})\mathbb{E}[\bar{S}-S^*|S]\right] \\
 &= \mathbb{E}\left[(S-\bar{S})(\bar{S}-\mathbb{E}[S^*|S])\right] \\
 &= \mathbb{E}\left[(S-\bar{S})(\bar{S}-\bar{S})\right] = 0 \\
 \Rightarrow \underbrace{\mathbb{E}[(S-Y)^2]}_{\text{Brier}} &= \underbrace{\mathbb{E}[S^*(1-S^*)]}_{\text{irreducible, IL}} + \underbrace{\mathbb{E}[(S^*-\bar{S})^2]}_{\text{grouping, GL}} + \underbrace{\mathbb{E}[(\bar{S}-S)^2]}_{\text{calibration, CL}}
 \end{aligned}$$

Where S is the classifier under scrutiny

$S^* \equiv \mathbb{E}[Y|X]$ is the Bayes-optimal classifier,

$\bar{S} \equiv \mathbb{E}[Y|S]$ is the ideally calibrated classifier

$$\underbrace{Y}_{\text{label}} \rightarrow \underbrace{\mathbb{E}[Y|X] = S^*}_{\text{best possible}} \rightarrow \underbrace{\mathbb{E}[Y|S] = \bar{S}}_{\text{best from the score}} \rightarrow \underbrace{S}_{\text{what the network outputs}}$$

$$\bar{S} = \mathbb{E}[S^*|S] \Rightarrow \text{GL}_B = \mathbb{E}[\text{Var}(S^*|S)]$$

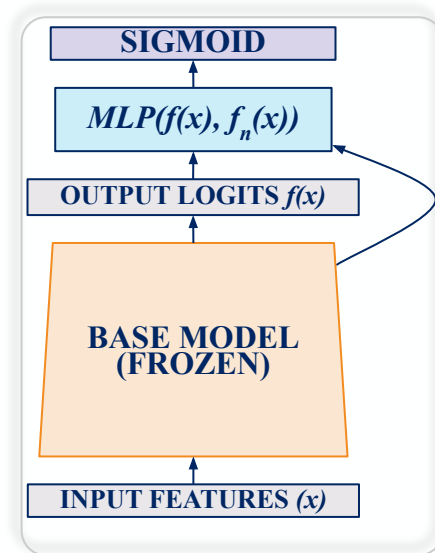
A bit more on “generalized” g-layers

- **Upsides:**

- Cheaper to train than LoRAs, but also has the ability to reduce grouping error
- Due to their implementation they can “more sharply” adjust the base model behavior
- Not in competition with LoRA, but rather serves as complementary method:
 - When a large, but smooth correction is needed, LoRAs shine
 - When a sharp refinement is needed, generalized g-layers are better (and more data efficient)

- **Downsides:**

- Prone to overfitting, needs careful control of parameters
- If the damage to the model is done “early on”, i.e. if the errors stem from the attention blocks, then this method can’t correct much



Asimov fit setup in bit more detail

Model and likelihood:

$$\nu_c(x) = c \nu_1(x) + \nu_0(x)$$

$$\frac{\nu_c(x)}{\nu_0(x)} = 1 + c r(x), \quad r(x) = \frac{\nu_1(x)}{\nu_0(x)} = \frac{s(x)}{1 - s(x)} = e^{f(x)}$$


$$\ln L(c) = c_{\text{true}} \sum_{i \in \text{class1}} w_i \ln(1 + c r_i) + \sum_{i \in \text{class0}} w_i \ln(1 + c r_i) - c W_1 + \text{const}$$

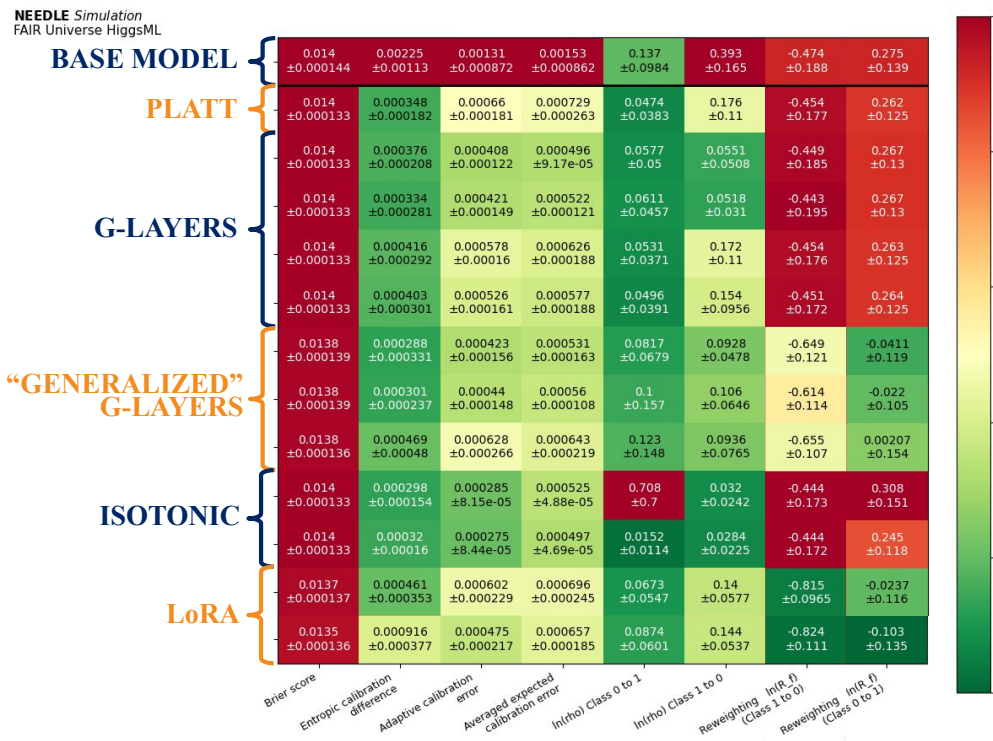
**Saturation into a “counting experiment”,
simplified explanation**

$$\frac{\partial}{\partial c} \ln(1 + c r) = \frac{r}{1 + c r} \xrightarrow{c r \gg 1} \frac{1}{c}$$

Connection to Fisher information

$$I(c_{\text{true}}) = \int \nu_0 \frac{r^2}{1 + c_{\text{true}} r} \xrightarrow{c_{\text{true}} \rightarrow 0} \int \nu_0 r^2$$

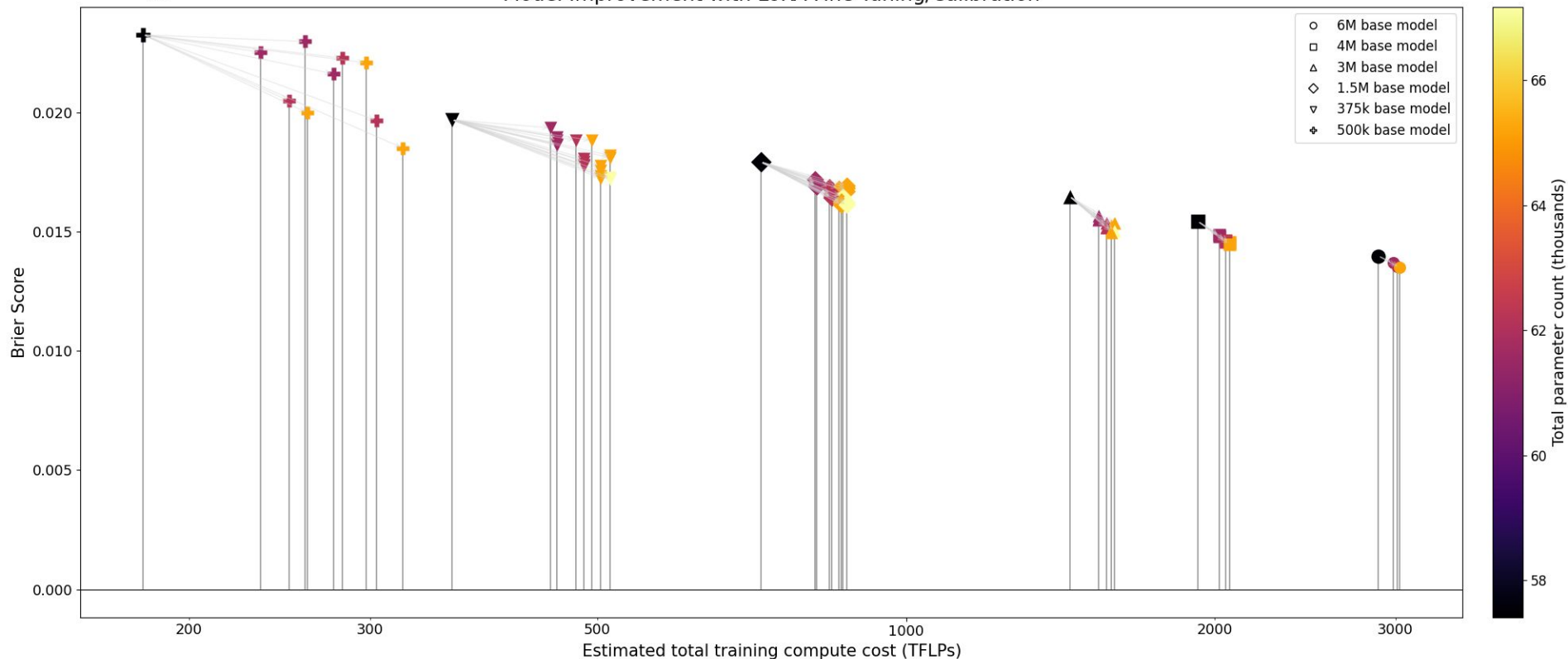
Bonus: diagnosing your model with post-calibration performance checks (calibration/refinement as a form of validation!)



Results: base models versus calibrated/refined variants (Brier score), true x-axis, with total model parameter counts

NEEDLE Simulation
Fair Universe HiggsML

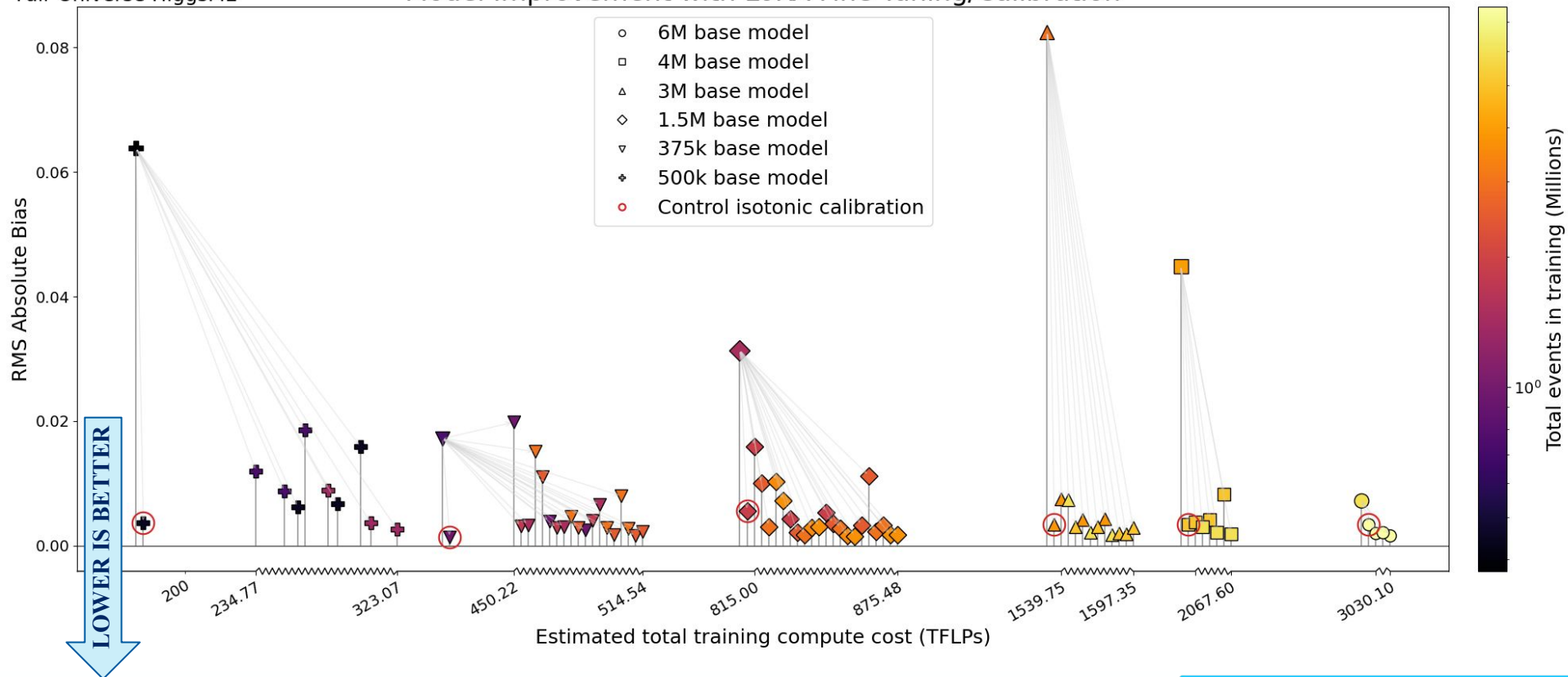
Model Improvement with LoRA Fine Tuning/Calibration



Results: base models versus calibrated/refined variants (asimov fit aggregate)

NEEDLE Simulation
Fair Universe HiggsML

Model Improvement with LoRA Fine Tuning/Calibration





More on NLRE: Density Ratio Approach to Diffusion Probabilistic Models

Density Ratio Approach to Diffusion Probabilistic Models: bridging the gap

Where NLRE falls short:

- If there is a clear separation boundary between the two classes (in some phase-space sub-region)*
- You have to perform tricks to get to the likelihood (“to get rid of the ratio”)

But, NLRE has very good properties:

- Overall good training+inference speed
- “Simple” training task

*This “lack of support” issue is very well known in NLRE. We’re simplifying the discussion on this point (also in relation to NLE) in the interest of time

Density Ratio Approach to Diffusion Probabilistic Models: bridging the gap

Where NLRE falls short:

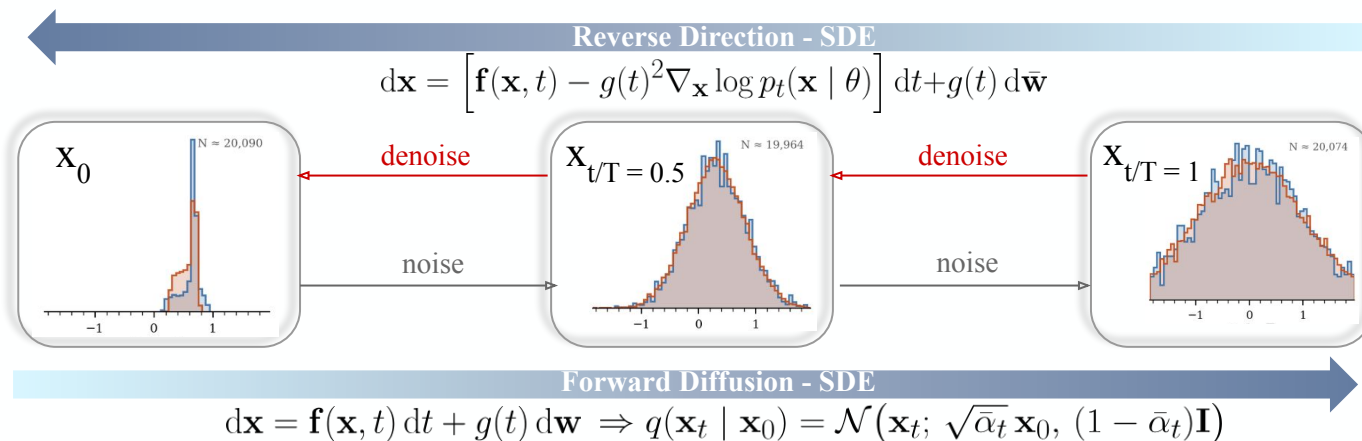
- If there is a clear separation boundary between the two classes (in some phase-space sub-region)
- You have to perform tricks to get to the likelihood (“to get rid of the ratio”)

But, NLRE has very good properties:

- Overall good training+inference speed
- “Simple” training task

Where generative NLE approaches fall short:

- Either large total training+inference compute cost*



*Here we give a diffusion probabilistic model as an example. The compute cost/speed argument is very nuanced, and we simplify in the interest of time.

Density Ratio Approach to Diffusion Probabilistic Models: bridging the gap

Where NLRE falls short:

- If there is a clear separation boundary between the two classes (in some phase-space sub-region)
- You have to perform tricks to get to the likelihood (“to get rid of the ratio”)

But, NLRE has very good properties:

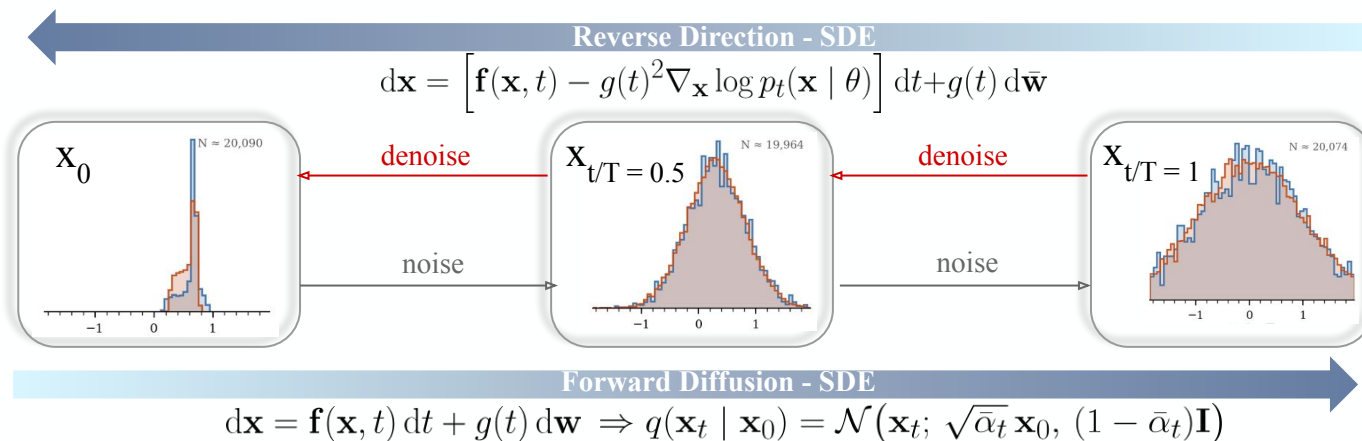
- Overall good training+inference speed
- “Simple” training task

Where generative NLE approaches fall short:

- Either large total training+inference compute cost* or
- Architecture limitations

But, generative NLE:

- Has no problems if an inter-sample separation boundary exists
- Gives you the likelihood outright



*Here we give a diffusion probabilistic model as an example. The compute cost/speed argument is very nuanced, and we simplify in the interest of time.

Density Ratio Approach to Diffusion Probabilistic Models: bridging the gap

Where NLRE falls short:

- If there is a clear separation boundary between the two classes (in some phase-space sub-region)
- You have to perform tricks to get to the likelihood (“to get rid of the ratio”)

But, NLRE has very good properties:

- Overall good training+inference speed
- “Simple” training task

Where generative NLE approaches fall short:

- Either large total training+inference compute cost* or
- Architecture limitations

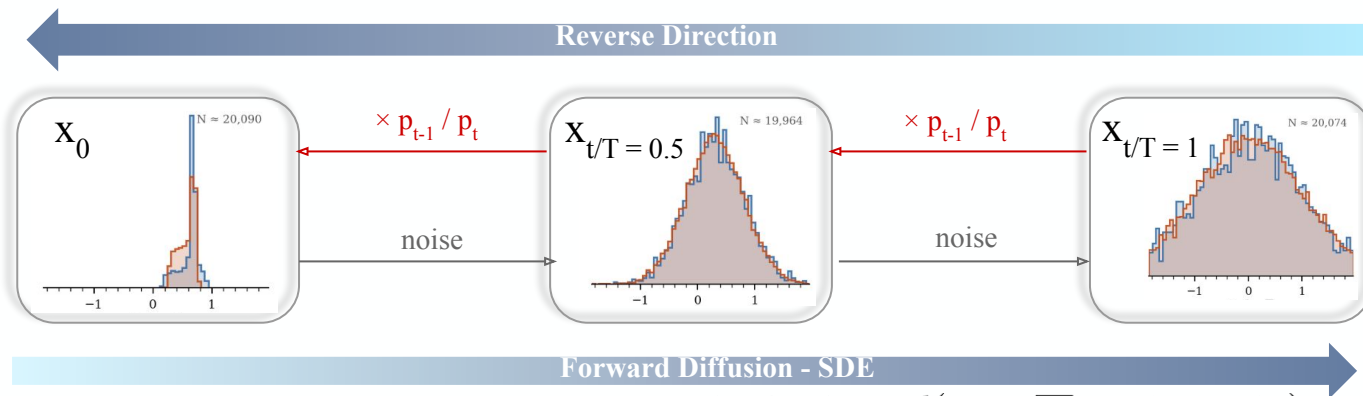
But, generative NLE:

- Has no problems if an inter-sample separation boundary exists
- Gives you the likelihood outright

What DRDPM (Density Ratio Approach to Diffusion Probabilistic Models) achieves?

- Takes “best of both worlds”: uses probability ratios to directly compute the likelihood
- Primarily focused on speed
- DRDPM still needs its probability ratios to be of very high quality!
- Thus, the same calibration/refinement arguments as for NLRE apply

DRDPM: How it works



$$d\mathbf{x} = \mathbf{f}(\mathbf{x}, t) dt + g(t) d\mathbf{w} \Rightarrow q(\mathbf{x}_t | \mathbf{x}_0) = \mathcal{N}(\mathbf{x}_t; \sqrt{\bar{\alpha}_t} \mathbf{x}_0, (1 - \bar{\alpha}_t)\mathbf{I})$$

DRDPM classifier NN gives you:

Class conditional probability

$$p(t|x, \theta) = \text{softmax}(z(x, t))$$

A softmax over diffusion noise levels:
“given x , say which time step t it came from”

In inference, only 1 forward pass is needed

What the DRDPM likelihood needs:

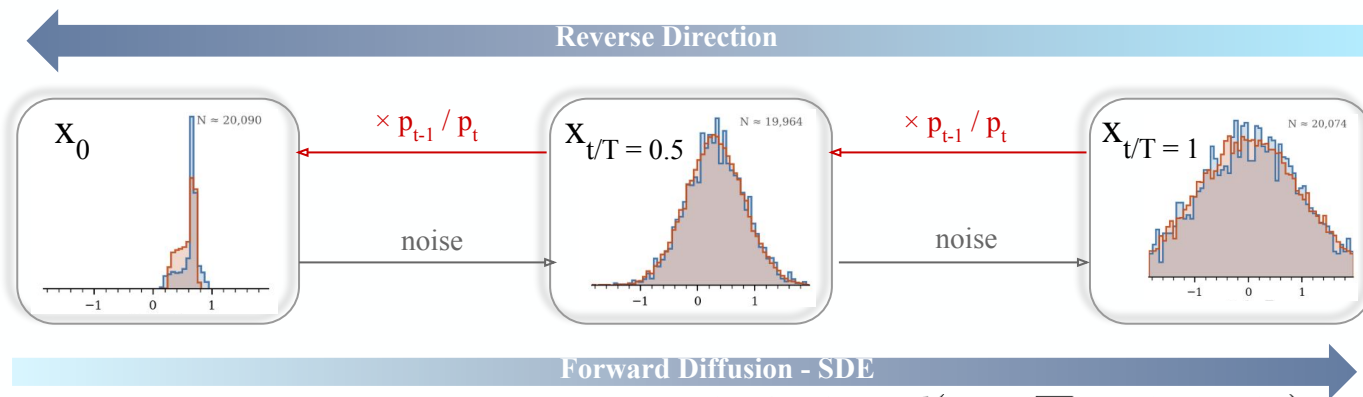
Only the difference of two logits

$$\log(p(x_t|\theta)) \approx \log(p(t|x_t, \theta)) - \log(p(T|x_t, \theta))$$

No Laplacian, no Jacobian trace, no hutchinson estimator calculation is required

0 backward passes in inference

DRDPM: How it works



$$d\mathbf{x} = \mathbf{f}(\mathbf{x}, t) dt + g(t) d\mathbf{w} \Rightarrow q(\mathbf{x}_t | \mathbf{x}_0) = \mathcal{N}(\mathbf{x}_t; \sqrt{\bar{\alpha}_t} \mathbf{x}_0, (1 - \bar{\alpha}_t) \mathbf{I})$$

DRDPM classifier NN gives you:

Class conditional probability

$$p(t|x, \theta) = \text{softmax}(z(x, t))$$

A softmax over diffusion noise levels:
“given x , say which time step t it came from”

In inference, only 1 forward pass is needed

What the DRDPM likelihood needs:

Only the difference of two logits

$$\log(p(x_t|\theta)) \approx \log(p(t|x_t, \theta)) - \log(p(T|x_t, \theta))$$

No Laplacian, no Jacobian trace, no hutchinson estimator calculation is required

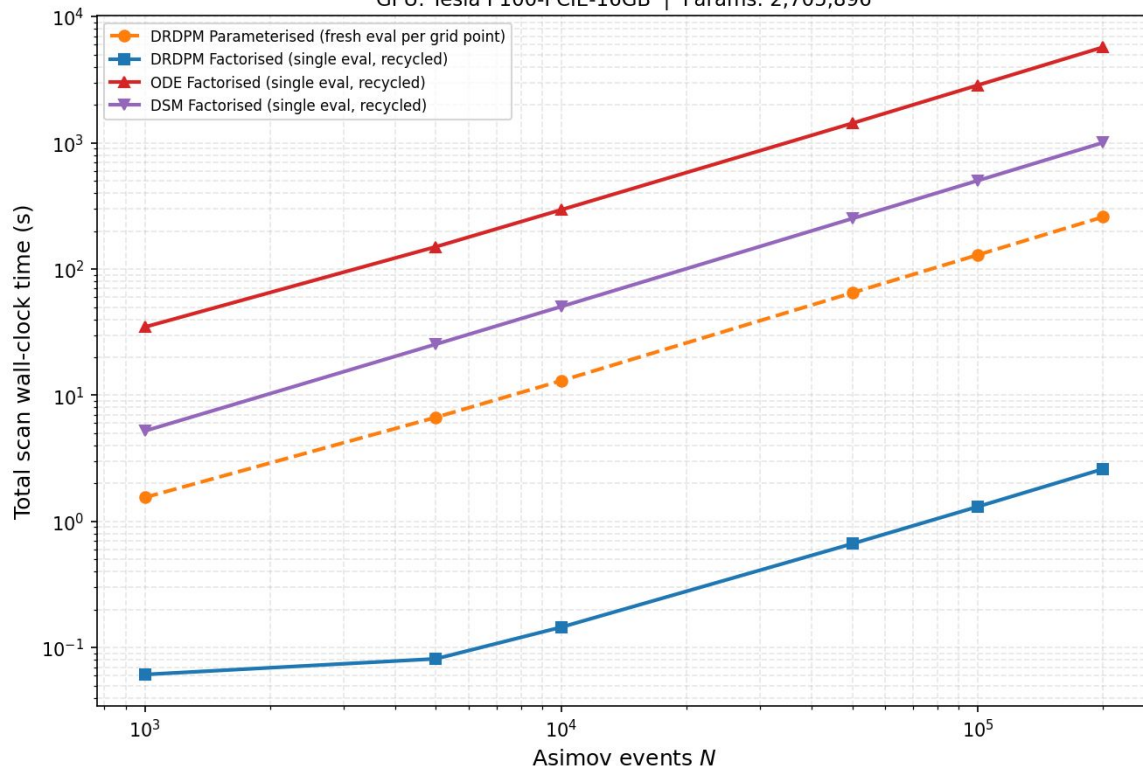
0 backward passes in inference

Cost of a likelihood scan

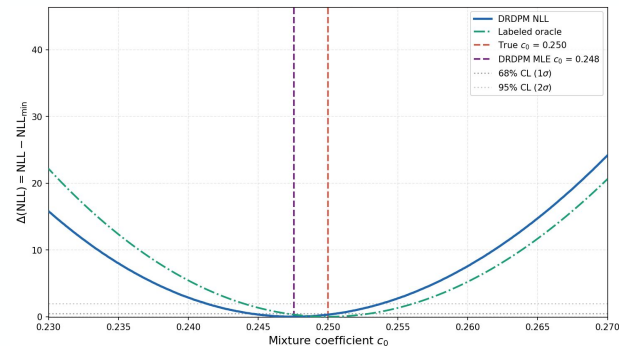
~~T~~ diffusion steps $\times$ ~~1~~ $\times$ ~~K~~ solver evaluations ~~1~~ $\times$ (forward + backward) passes ~~1~~ $\times$ N events $\times$ G scan points in θ

DRDPM: Inference speed

Total Scan Time vs N ($G=100$ grid points)
GPU: Tesla P100-PCIE-16GB | Params: 2,705,896



And it works on asimov signal injection fits! (see backup)



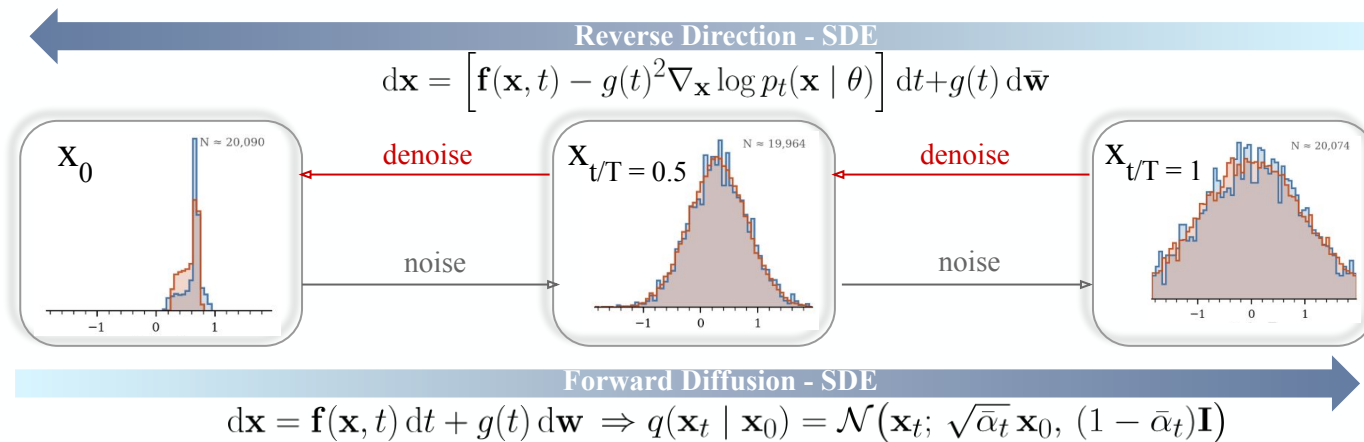
Cost of a likelihood scan

~~T~~ diffusion steps $\times$ ~~1~~ $\times$ ~~K~~ solver evaluations $\times$ ~~1~~ $\times$ (forward + backward) passes $\times$ ~~1~~ $\times$ N events $\times$ G scan points in θ



Backup: DRDPM

Diffusion probabilistic models, inference cost



What the network gives you

A Gradient

$$\nabla_{\mathbf{x}} \log p_t(\mathbf{x} | \theta) = -\frac{\epsilon_{\phi}(\mathbf{x}_t, t)}{\sqrt{1 - \bar{\alpha}_t}}$$

Denoising score matching learns the **first** derivative of the log-density, a vector with d components.

1 forward pass

What the likelihood needs

Divergence $\rightarrow$ differentiate the network

$$\nabla_{\mathbf{x}}^2 \log p_t(\mathbf{x} | \theta) \approx -\frac{1}{\sqrt{1 - \bar{\alpha}_t}} \mathbf{u}^{\top} \left(\frac{\partial \epsilon_{\phi}(\mathbf{x}, t)}{\partial \mathbf{x}} \right) \mathbf{u}$$

∇^2 is the **Laplacian** obtained by back-propagating through the score.

+ 1 backward pass

What needs to be solved

An ODE trajectory

$$\log p_0(\mathbf{x}_0) = \log p_T(\mathbf{x}_T) - \frac{1}{2} \int_0^T \beta(t) \left[d + \nabla_{\mathbf{x}}^2 \log p_t(\mathbf{x}_t | \theta) \right] dt$$

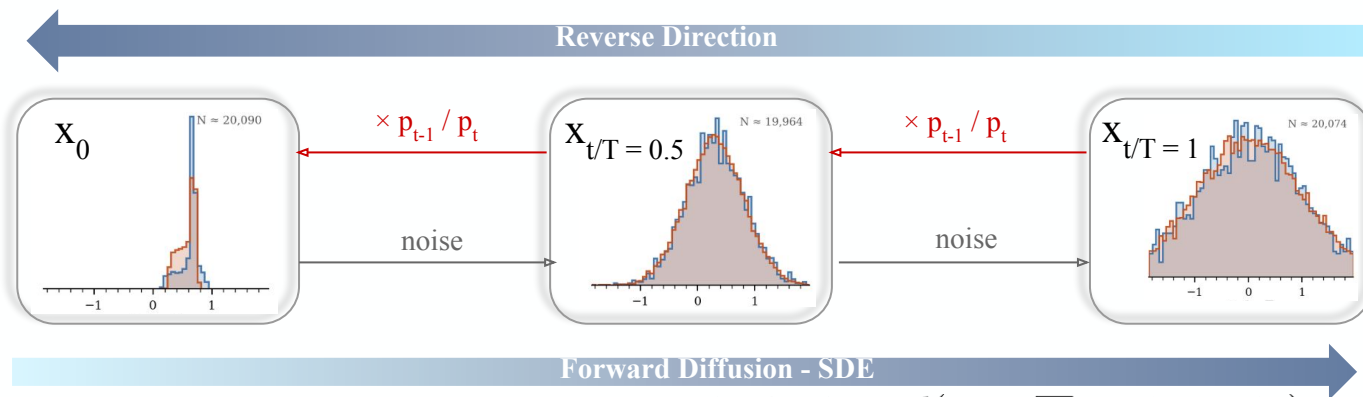
An adaptive solver takes **O(T)** evaluations, and runs **K steps** in solver.

× every solver step

Cost of one likelihood scan

T diffusion steps **×** **K** solver evaluations **×** (forward+backward) pass **×** **N** events **×** **G** scan points in θ

DRDPM, inference cost



$$d\mathbf{x} = \mathbf{f}(\mathbf{x}, t) dt + g(t) d\mathbf{w} \Rightarrow q(\mathbf{x}_t | \mathbf{x}_0) = \mathcal{N}(\mathbf{x}_t; \sqrt{\bar{\alpha}_t} \mathbf{x}_0, (1 - \bar{\alpha}_t)\mathbf{I})$$

DRDPM classifier NN gives you:

Class conditional probability

$$p(t|x, \theta) = \text{softmax}(z(x, t))$$

A softmax over diffusion noise levels:
“given x , say which time step t it came from”

In inference, only 1 forward pass is needed

What the DRDPM likelihood needs:

Only the difference of two logits

$$\log(p(x_t|\theta)) \approx \log(p(t|x_t, \theta)) - \log(p(T|x_t, \theta))$$

No Laplacian, no Jacobian trace, no hutchinson estimator calculation is required

0 backward passes in inference

Cost of a likelihood scan

~~T~~ diffusion steps $\times$ ~~1~~ $\times$ ~~K~~ solver evaluations ~~1~~ $\times$ (forward + backward) passes ~~1~~ $\times$ N events $\times$ G scan points in θ

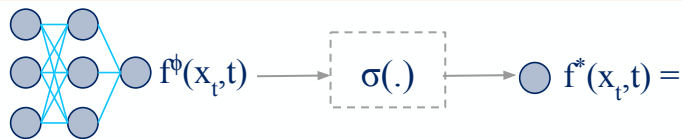
DRDPM: more details on the model used and method basics

- **Forward Diffusion Process:**

- Acts as an interpolation strategy between the target distribution $p(x_0 | \theta)$ and a **Gaussian anchor** $p(x_T | \theta)$

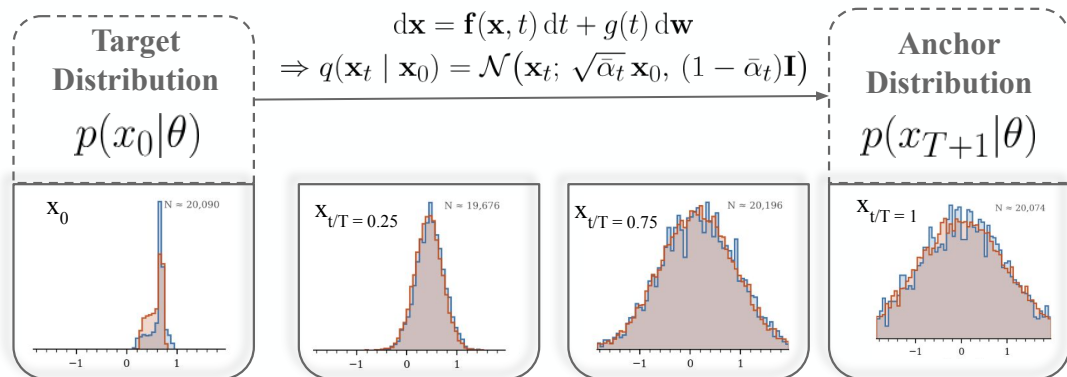
- **Neural Surrogate:**

- The neural surrogate learns a **time step** multi-classifier:



- **Likelihood:**

- The likelihood can be extracted analytically from the Gaussian probability path by the **density ratio** of the end logits **t=0 & t=T**



$$p_{\mathbf{x}_t|\theta}(\mathbf{x} | \theta) = \frac{p_t(T+1)}{p_t(t)} \frac{p_{t|\bar{\mathbf{x}},\theta}(t | \mathbf{x}, \theta)}{p_{t|\bar{\mathbf{x}},\theta}(T+1 | \mathbf{x}, \theta)} \mathcal{N}(\mathbf{x}; \mathbf{0}, \mathbf{I}),$$

DRDPM: more details on the model used and method basics

- **Forward Diffusion Process:**

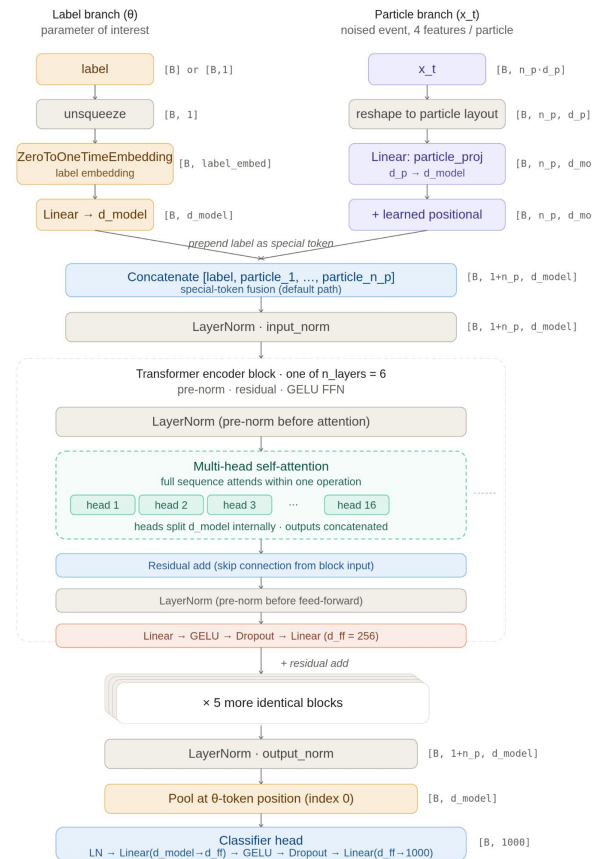
- Acts as an interpolation strategy between the target distribution $p(x_0 | \theta)$ and a **Gaussian anchor** $p(x_T | \theta)$

- **Neural Surrogate:**

- The neural surrogate learns a **time step** multi-classifier:

- **Likelihood:**

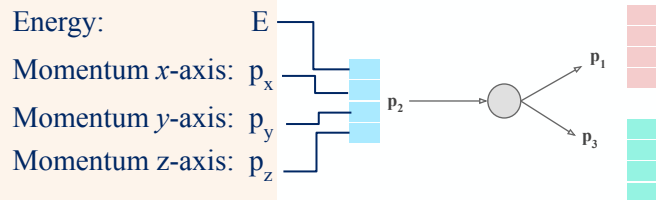
- The likelihood can be extracted analytically from the Gaussian probability path by the **density ratio** of the end logits **t=0 & t=T**



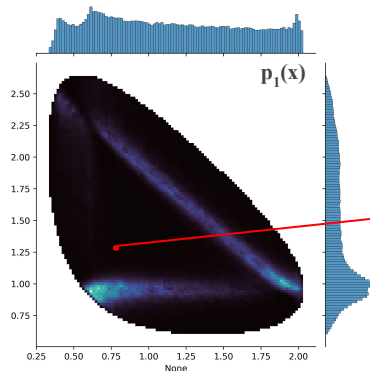
DRDPM: 3-body decay dataset details

• Monte Carlo - 12-Dim. Generator Space

- Generate 3-particles with 4-vector information:

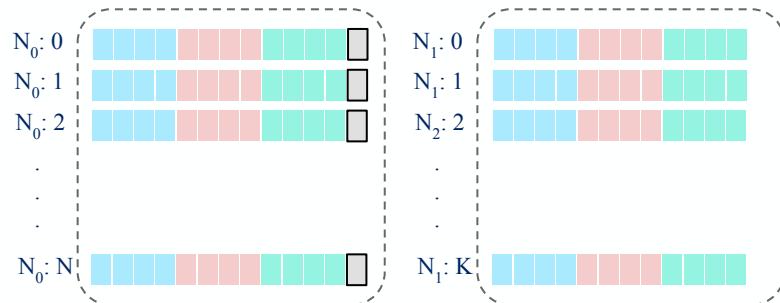
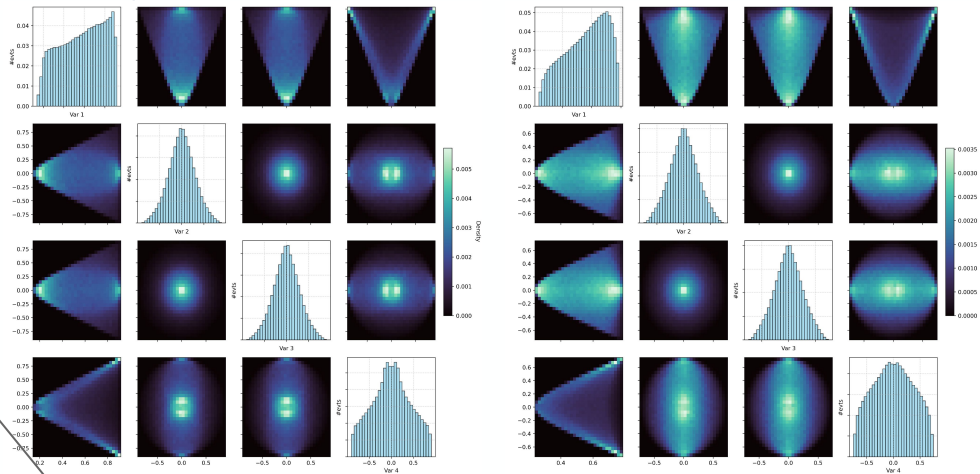


- Record matrix elements of each sample point in the (s_{12}, s_{23}) space:



- Evaluate the matrix element
- Calculate the phase-space element
- Normalise by the solid angle of emission in 3D space
- Normalise by the width of the decay

$$p_{12d}(x | \theta) = \frac{|\mathcal{M}(x; \theta)|^2 d\Phi_3(\theta)}{\Gamma(\theta) V_{\text{angles}}}$$



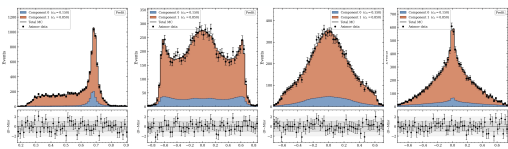
DRDPM: asimov fit results and setup

- **Monte Carlo Asimov**

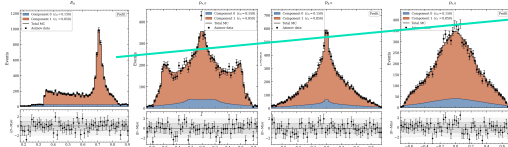
- Construct an entirely MC based ‘observed’ dataset with the task of extracting the mixture coefficient ‘ c ’:

- **Black points:** MC based asimov ‘observed’ data
- **Histograms:** MC Expectation

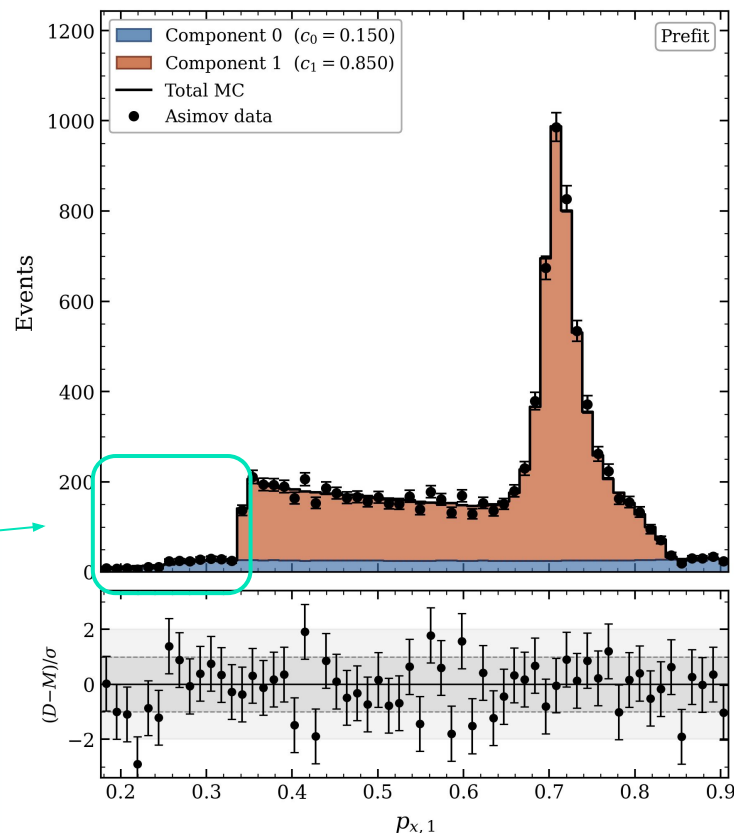
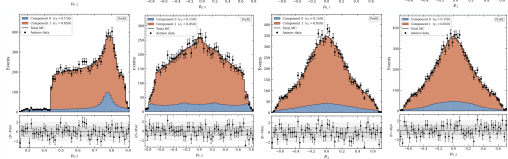
Particle 1



Particle 2



Particle 3



DRDPM: asimov fit results and setup

- **Monte Carlo Asimov**

- Construct an entirely MC based ‘observed’ dataset with the task of extracting the mixture coefficient ‘ c ’:

- **Black points:** MC based asimov ‘observed’ data
- **Histograms:** MC Expectation

- **POI Scan:**

- Run a scan over the mixture coefficient range $c \in [0.05, 0.45]$

$$p(x|c) = c \cdot p(x|\theta_0) + (1 - c) \cdot p(x|\theta_1)$$

