



UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386



Calomplification: Statistics of Generative Calorimeter Models

Sebastian Bieringer¹, Anja Butter, Sascha Diefenbacher, Engin Enren,
Frank Gaede, Daniel Hundshausen, Gregor Kasieczka,
Benjamin Nachman, Tilman Plehn, Mathias Trabs

¹Institut für Experimentalphysik, Universität Hamburg, Germany

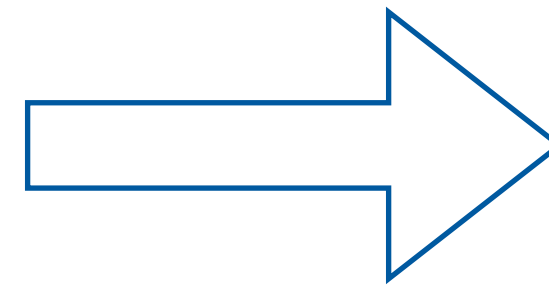
sebastian.guido.bieringer@uni-hamburg.de

CDCS Opening Symposium 2022

Introduction

Need to speed up MC

- Event generation
- Calorimeter simulation



Use generative machine learning models like

- Generative Adversarial Networks (GANs)
- or Variational Autoencoders (VAEs)

$$\text{simulation speed} = \frac{\# \text{ samples}}{\text{time}}$$

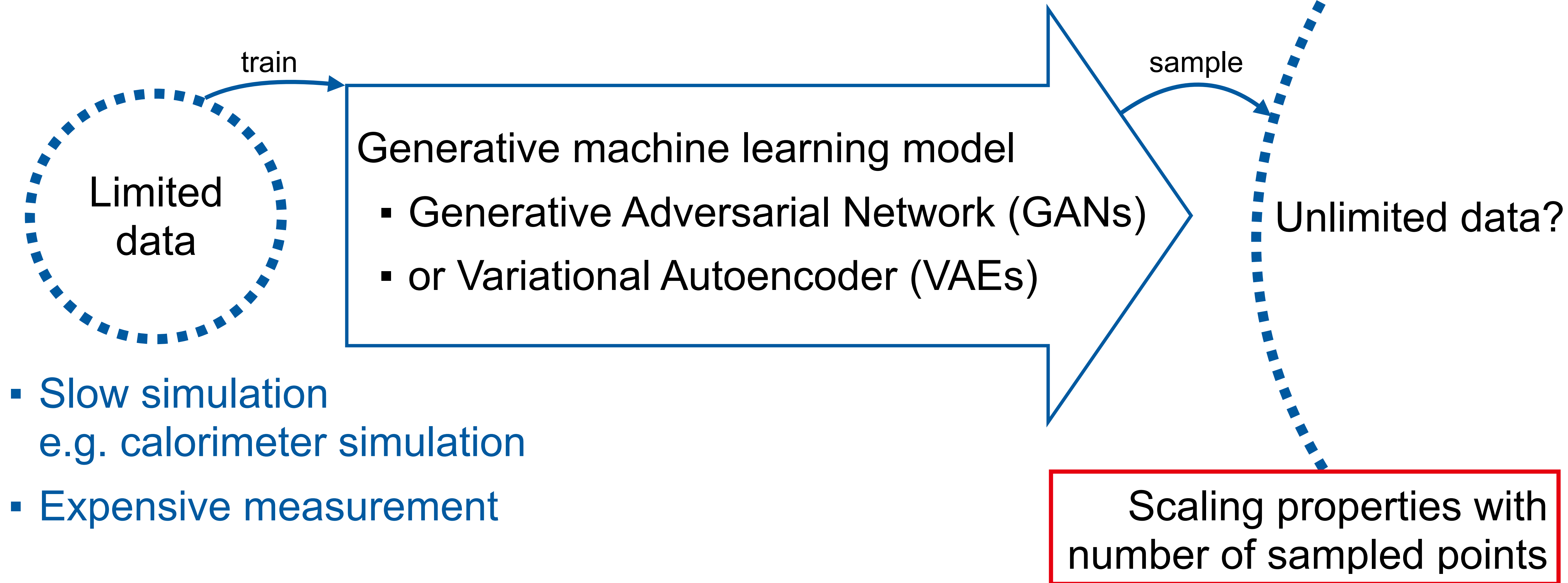
What about # samples?

A. Butter et al. *GANplifying Event Samples*. 2021. arXiv: [2008.06545 \[hep-ph\]](https://arxiv.org/abs/2008.06545)

S. Bieringer et al. *Calomplification -- The Power of Generative Calorimeter Models*. 2022. arXiv: [2202.07352 \[hep-ph\]](https://arxiv.org/abs/2202.07352)

Introduction

DASHH

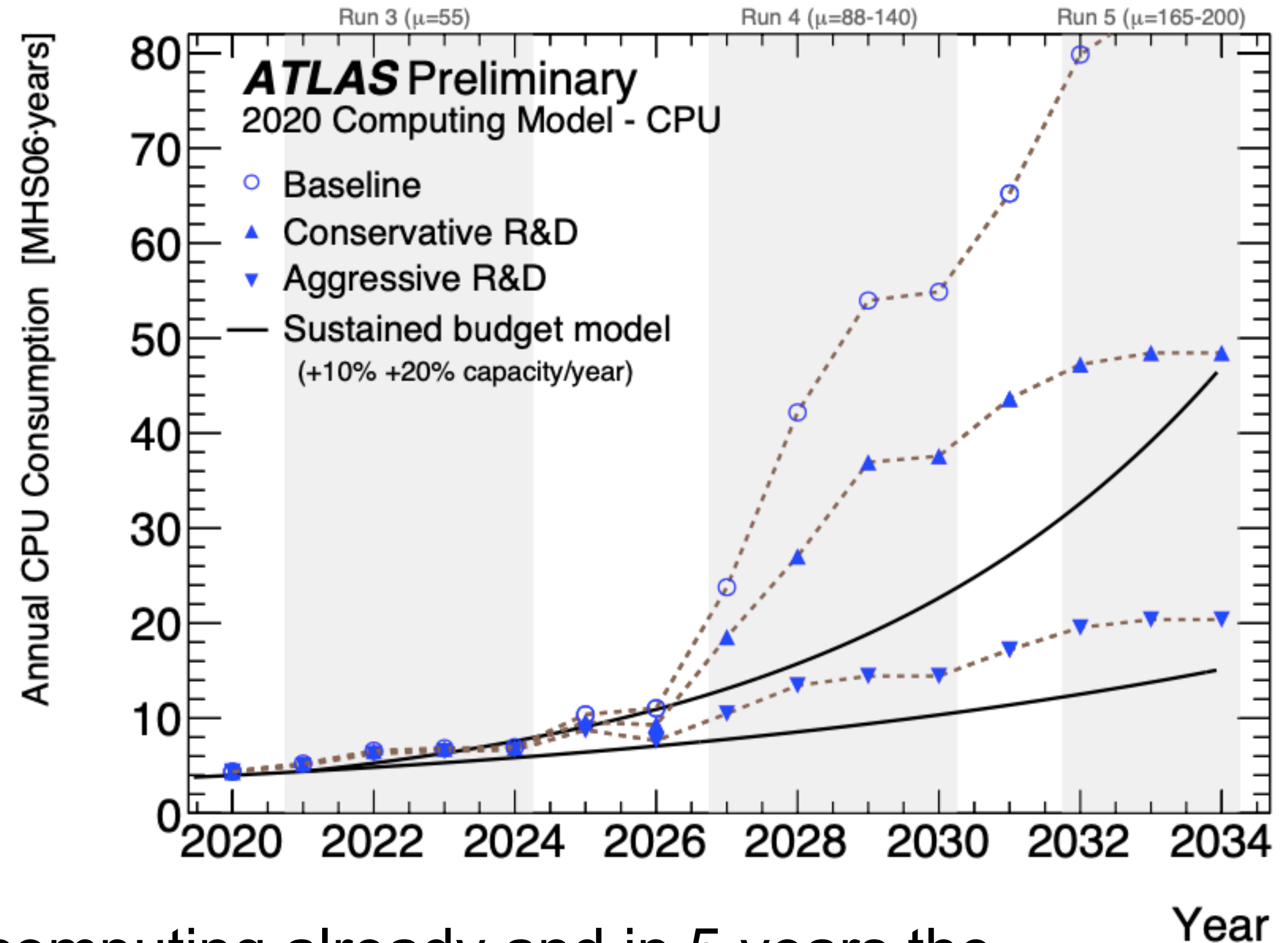
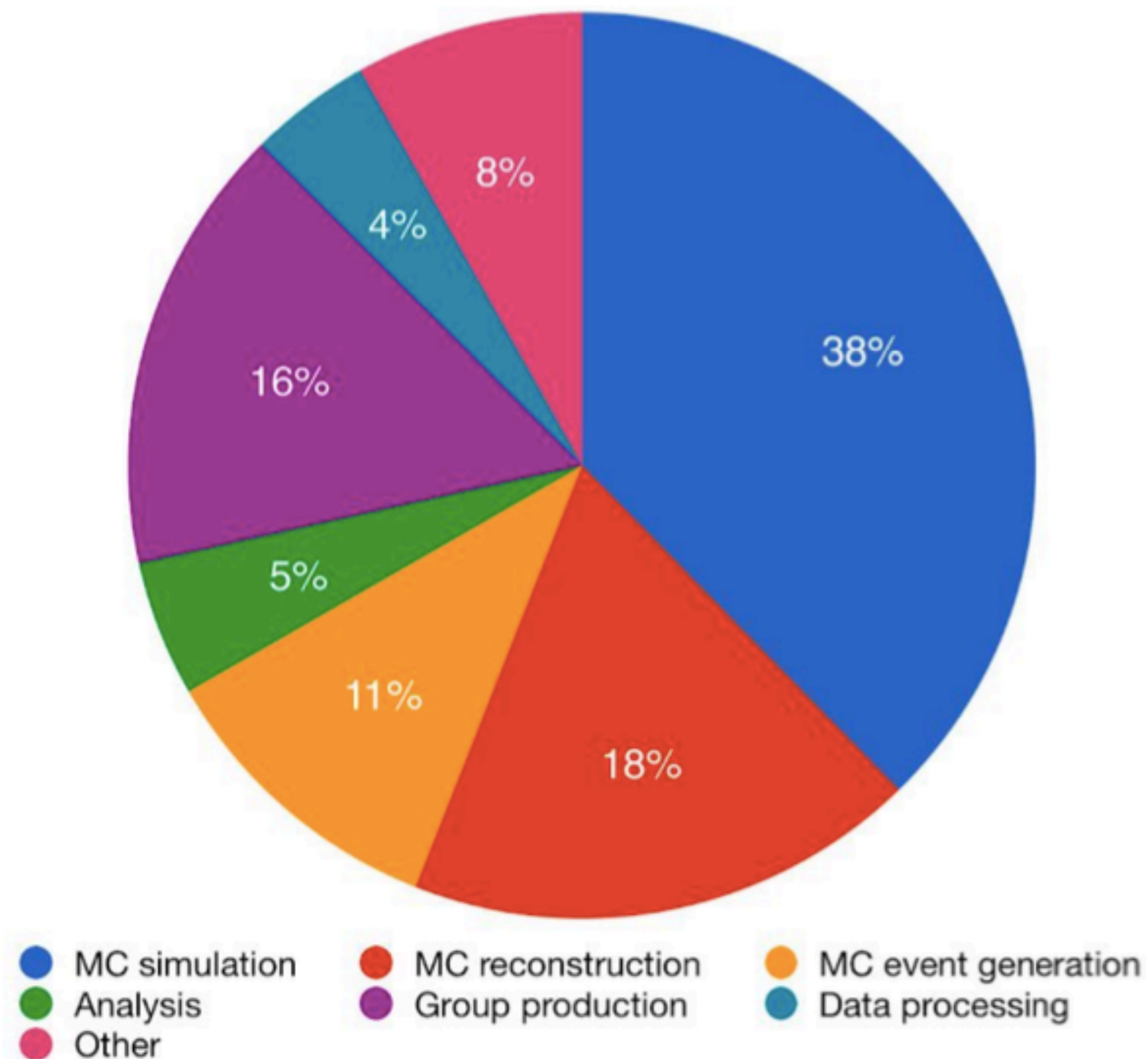


A. Butter et al. *GANplifying Event Samples*. 2021. arXiv: [2008.06545 \[hep-ph\]](https://arxiv.org/abs/2008.06545)

S. Bieringer et al. *Calomplification -- The Power of Generative Calorimeter Models*. 2022. arXiv: [2202.07352 \[hep-ph\]](https://arxiv.org/abs/2202.07352)

Introduction

Wall clock consumption per workflow

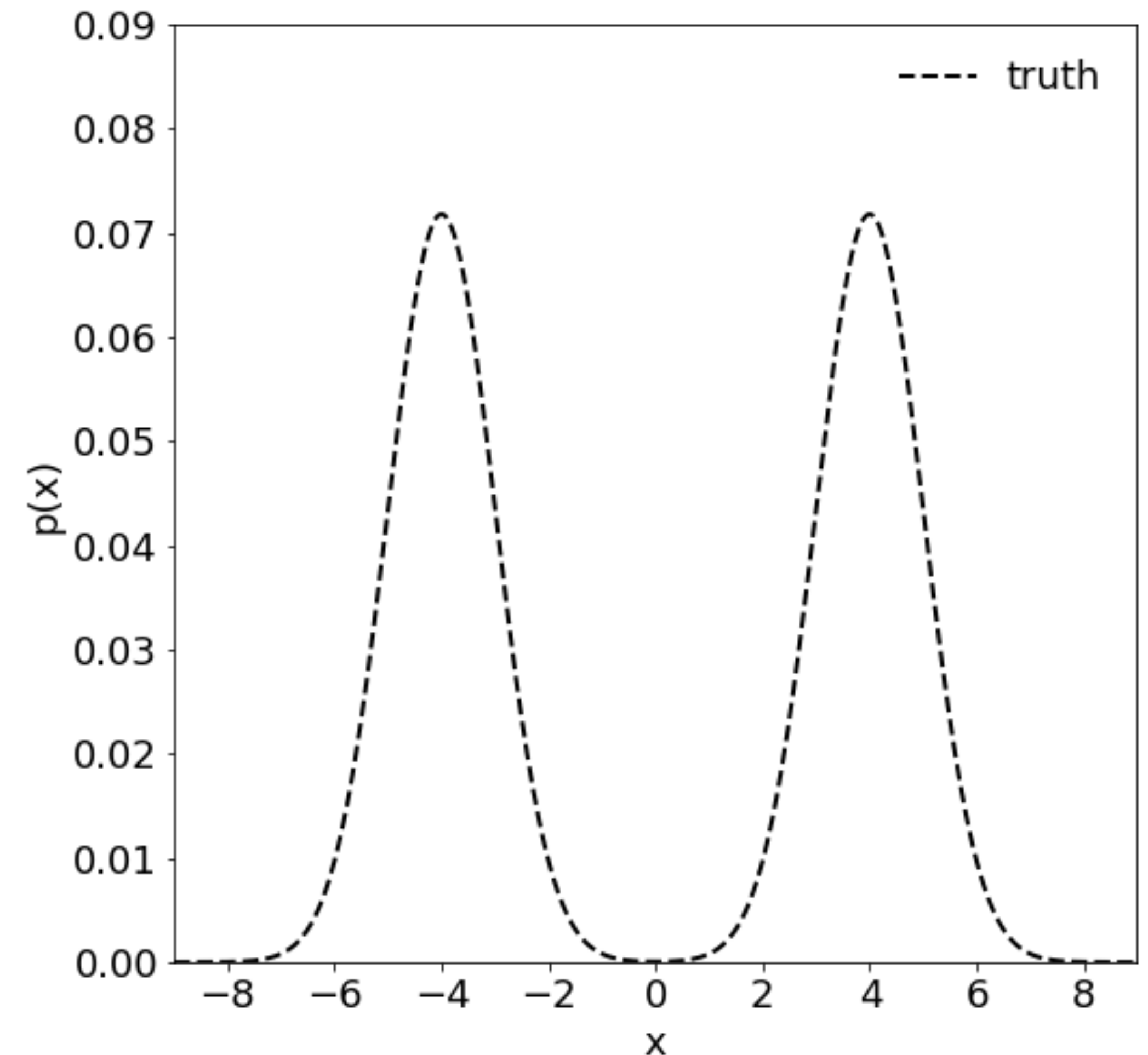


MC is a big bottleneck in computing already and in 5 years the demand for CPU time will outgrow the supply [0]

Toy Model: Setup

- Underlying function:

$$P(x) = \frac{1}{2} \left(\mathcal{N}_{-4,1}(x) + \mathcal{N}_{4,1}(x) \right)$$

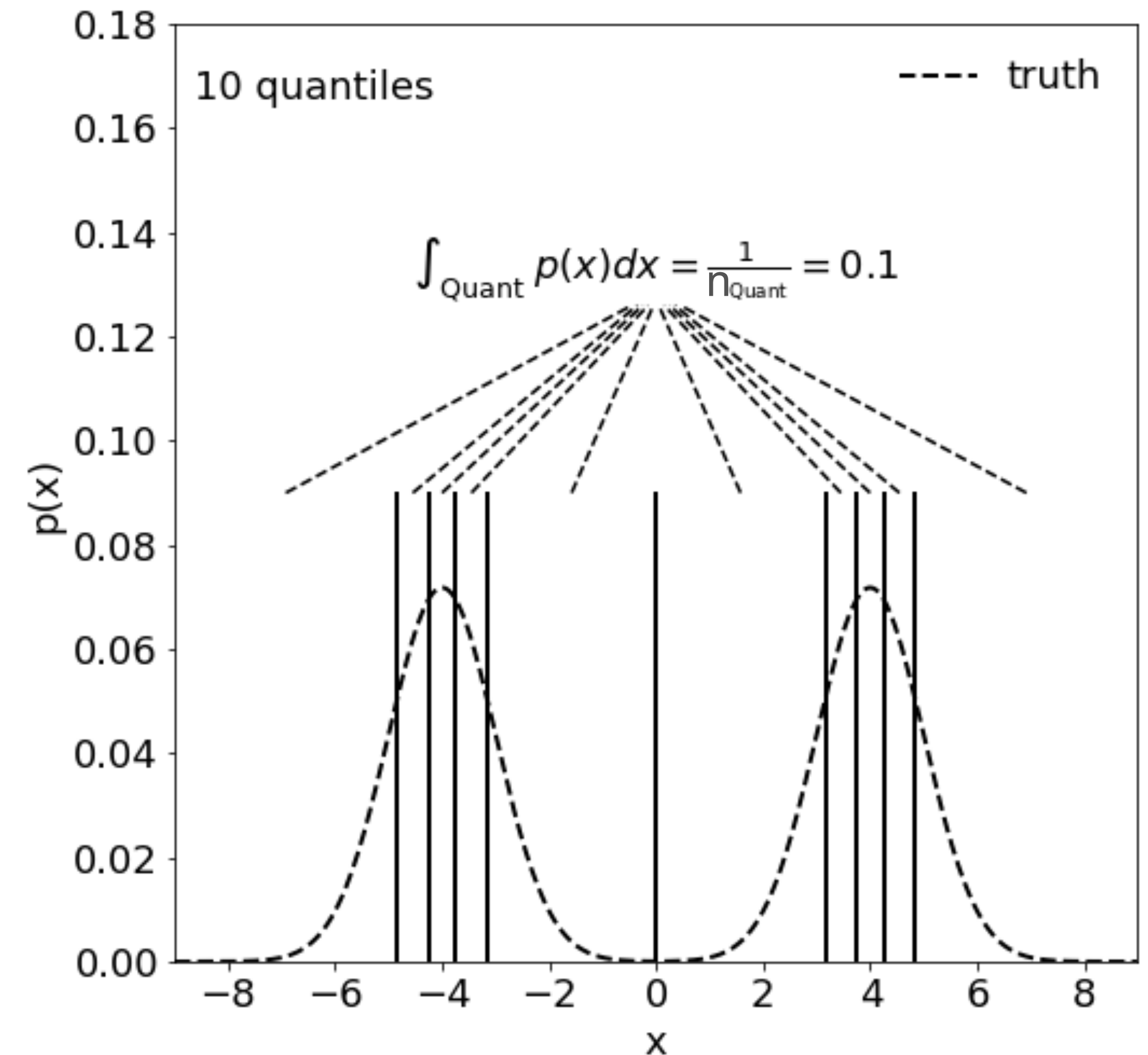


Toy Model: Setup

- Underlying function:

$$P(x) = \frac{1}{2} \left(\mathcal{N}_{-4,1}(x) + \mathcal{N}_{4,1}(x) \right)$$

- "Pearson χ^2 -test":
 - Introduce equal probability quantiles



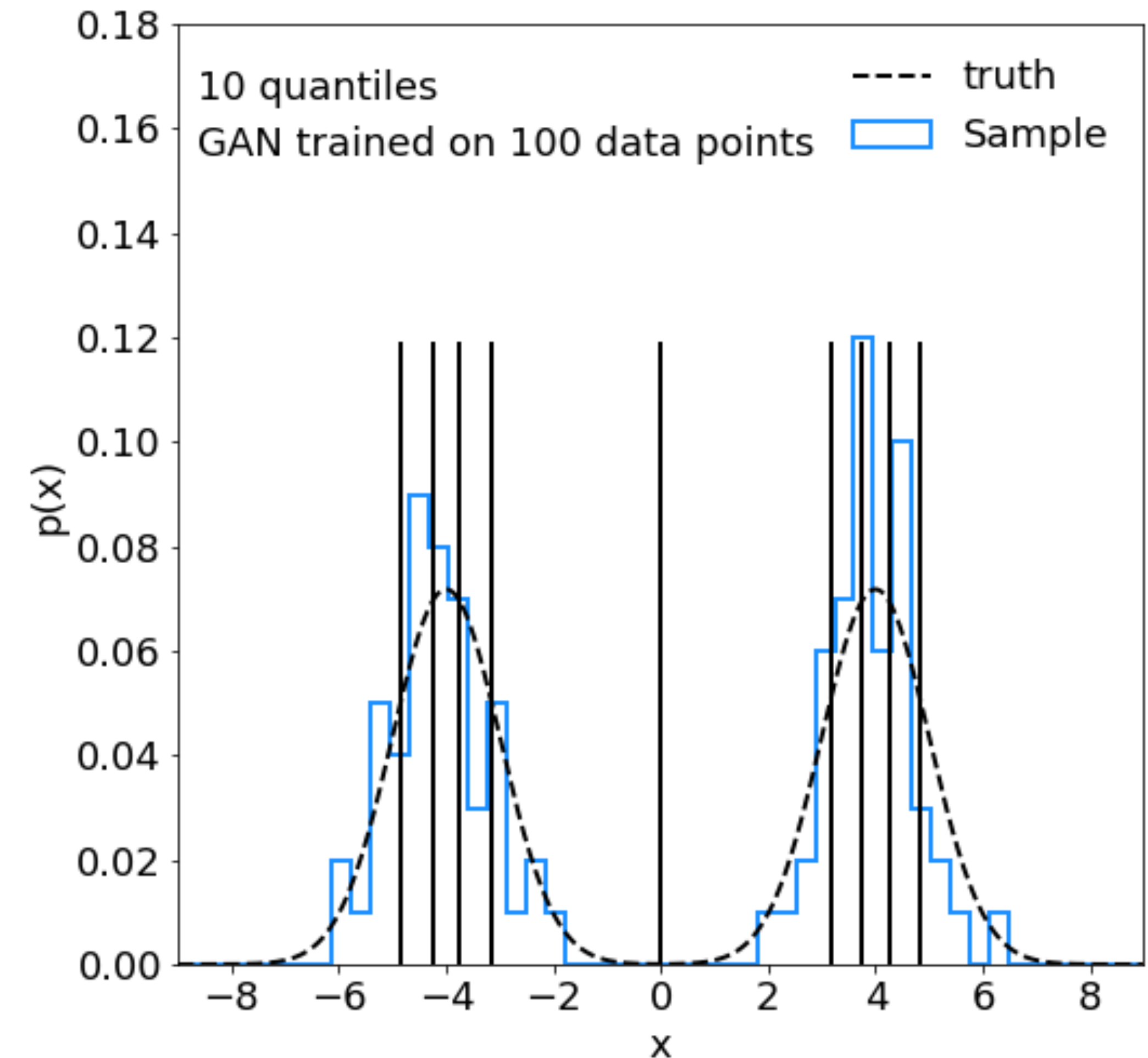
Toy Model: Setup

- Underlying function:

$$P(x) = \frac{1}{2} \left(\mathcal{N}_{-4,1}(x) + \mathcal{N}_{4,1}(x) \right)$$

- "Pearson χ^2 -test":
 - Introduce equal probability quantiles
 - Generate data
 - Calculate deviation metric

$$\hat{\chi}_{n_{\text{quant}}}^2 = n_{\text{quant}} \sum_{j=0}^{n_{\text{quant}}} \left(x_j - \frac{1}{n_{\text{quant}}} \right)^2$$



Toy Model: Setup

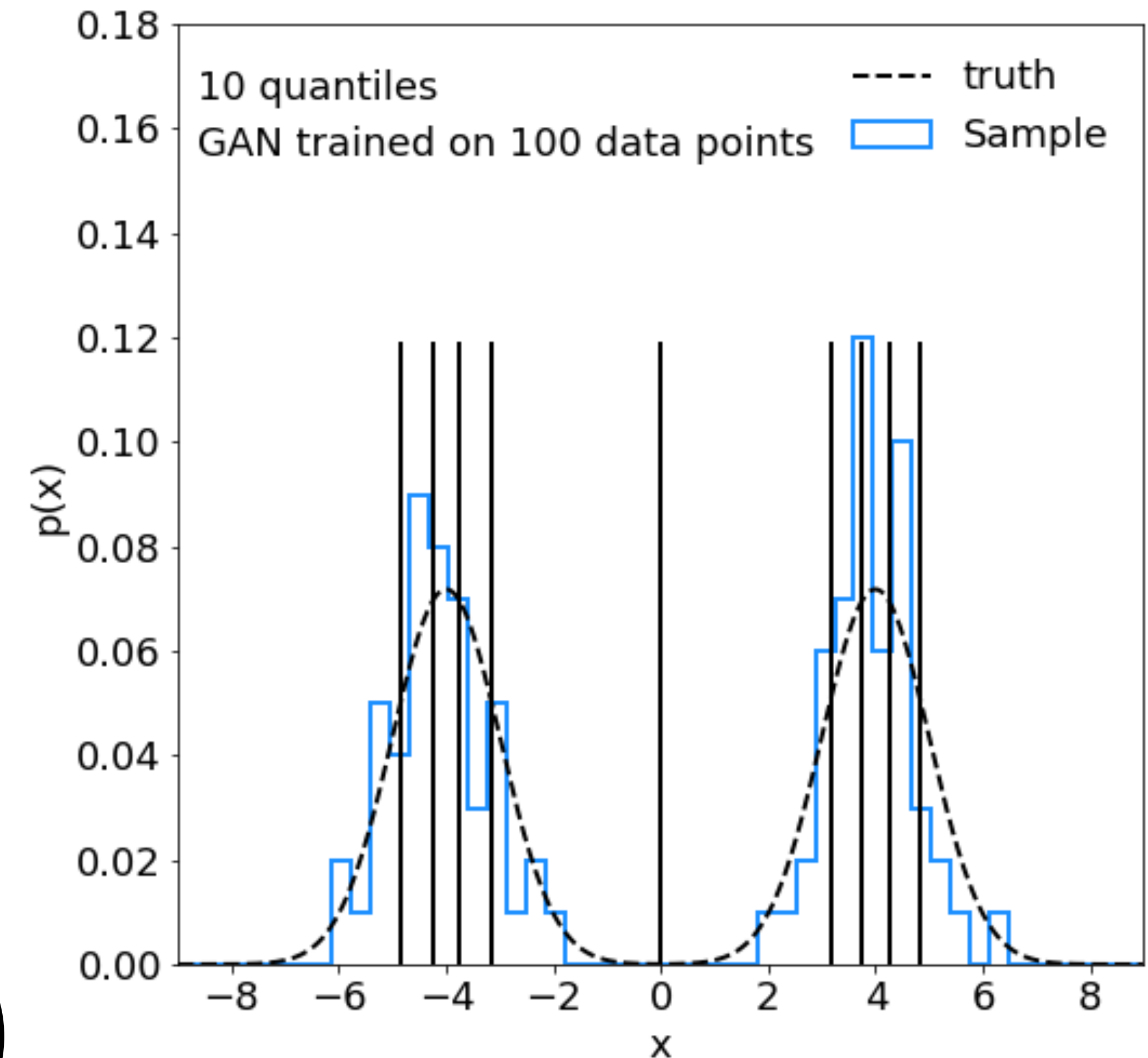
- Underlying function:

$$P(x) = \frac{1}{2} \left(\mathcal{N}_{-4,1}(x) + \mathcal{N}_{4,1}(x) \right)$$

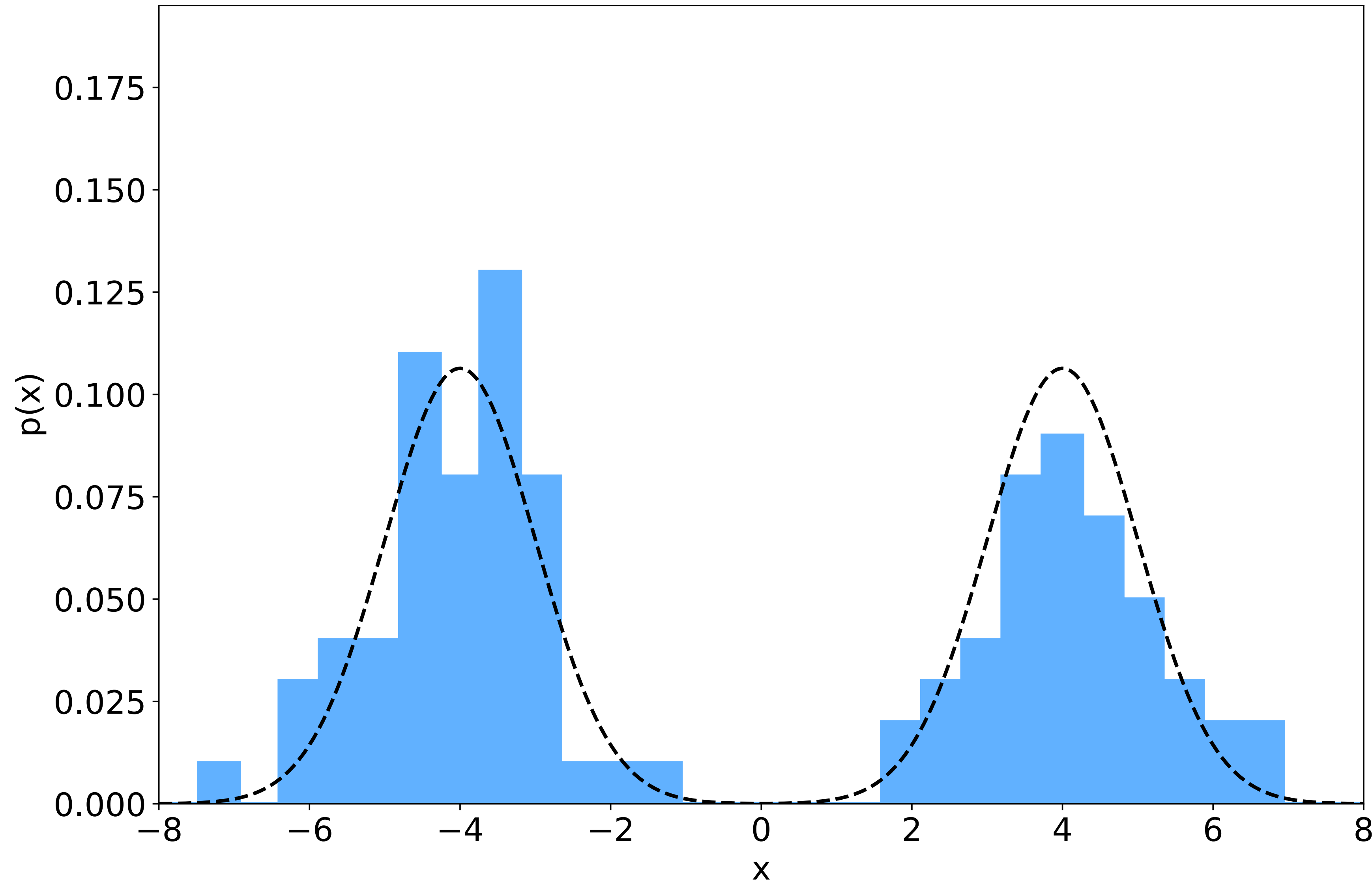
- "Pearson χ^2 -test":
 - Introduce equal probability quantiles
 - Generate data and calculate

$$\hat{\chi}_{n_{\text{quant}}}^2 = n_{\text{quant}} \sum_{j=0}^{n_{\text{quant}}} \left(x_j - \frac{1}{n_{\text{quant}}} \right)^2$$

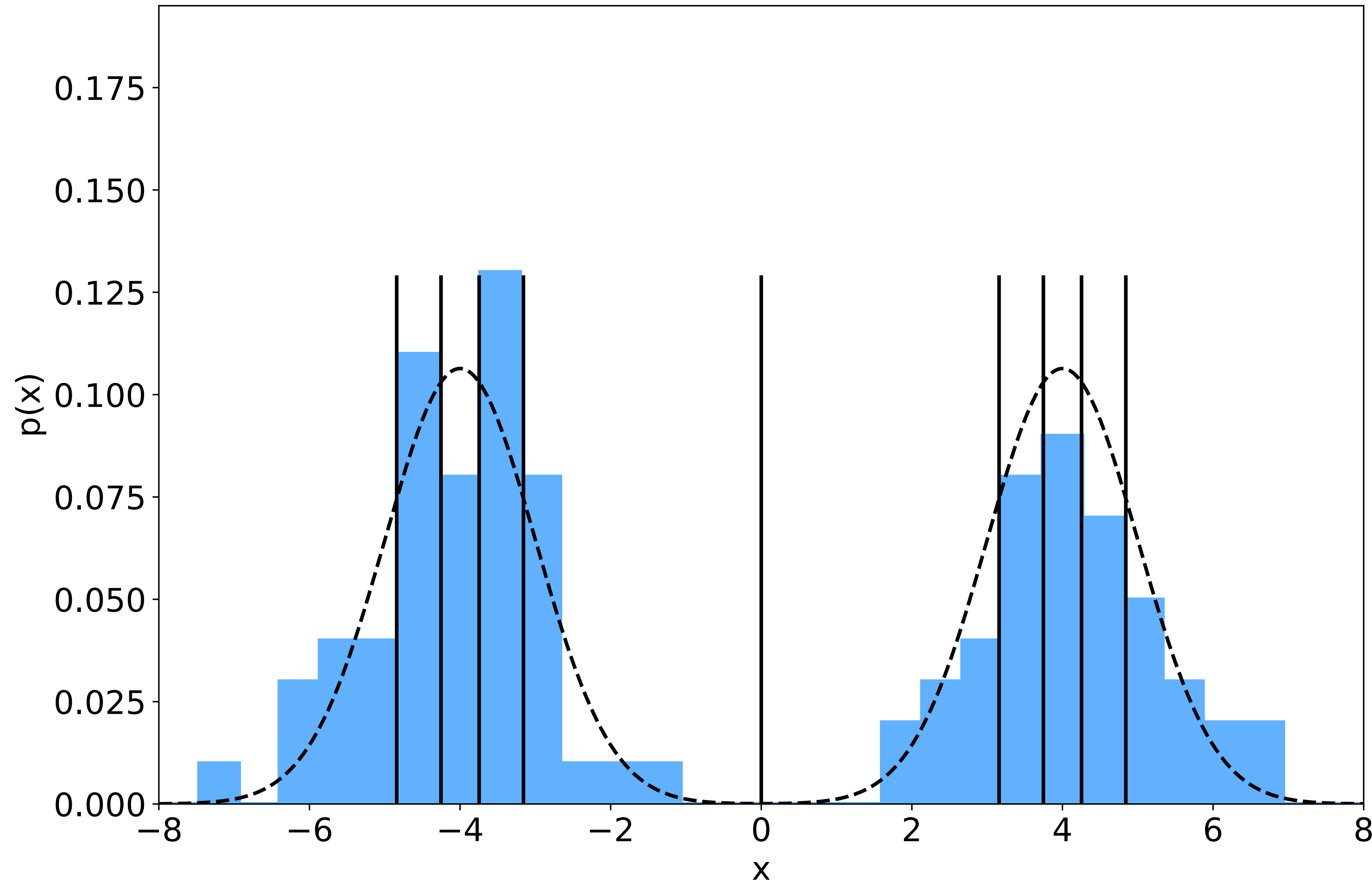
with $\hat{\chi}_{n_{\text{quant}}}^2 \xrightarrow{n_{\text{quant}} \rightarrow \infty} \chi^2 \left(P(x), P_{\text{samples}} \right)$



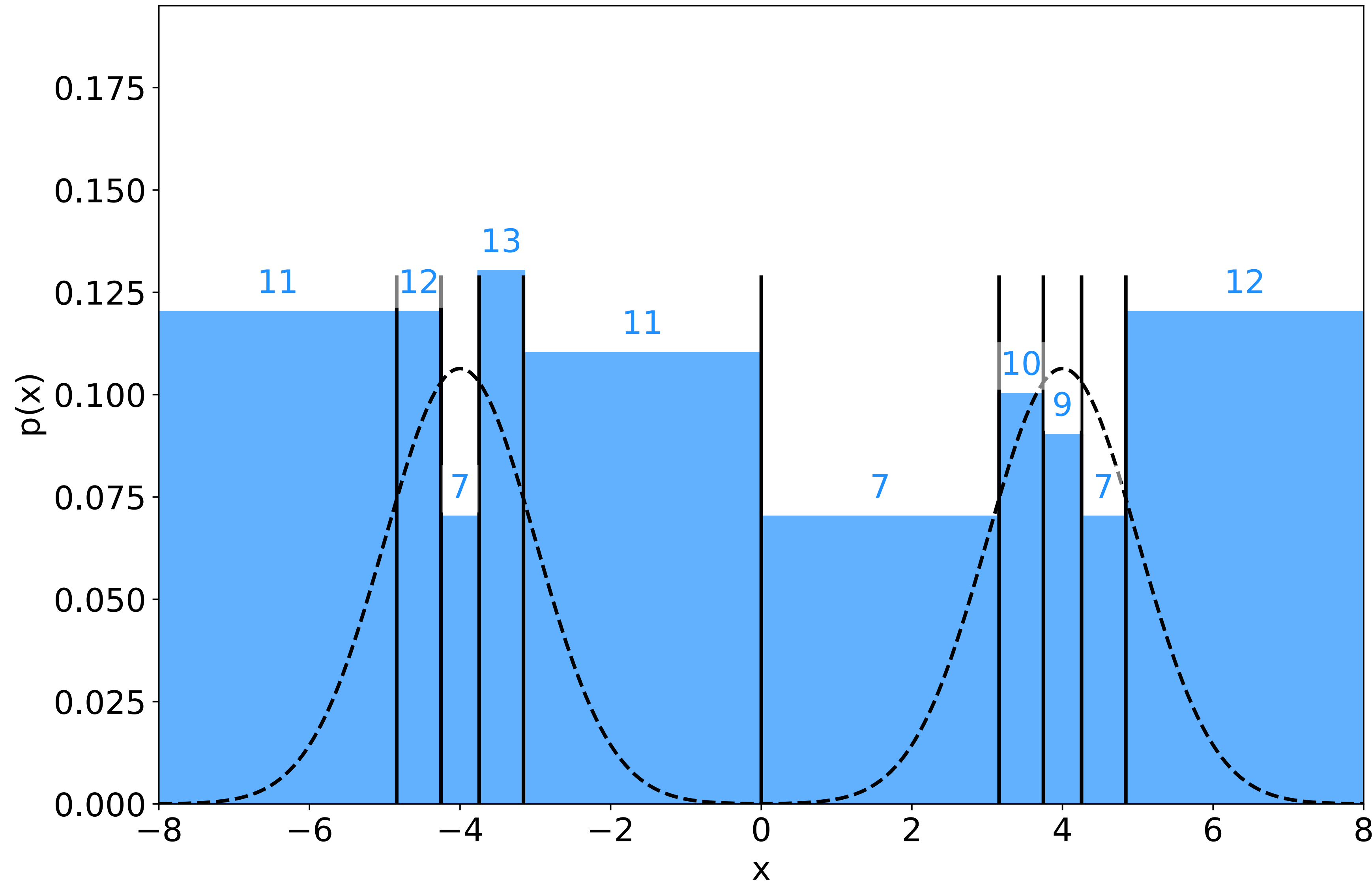
Bonus: How does the χ^2 measure work?



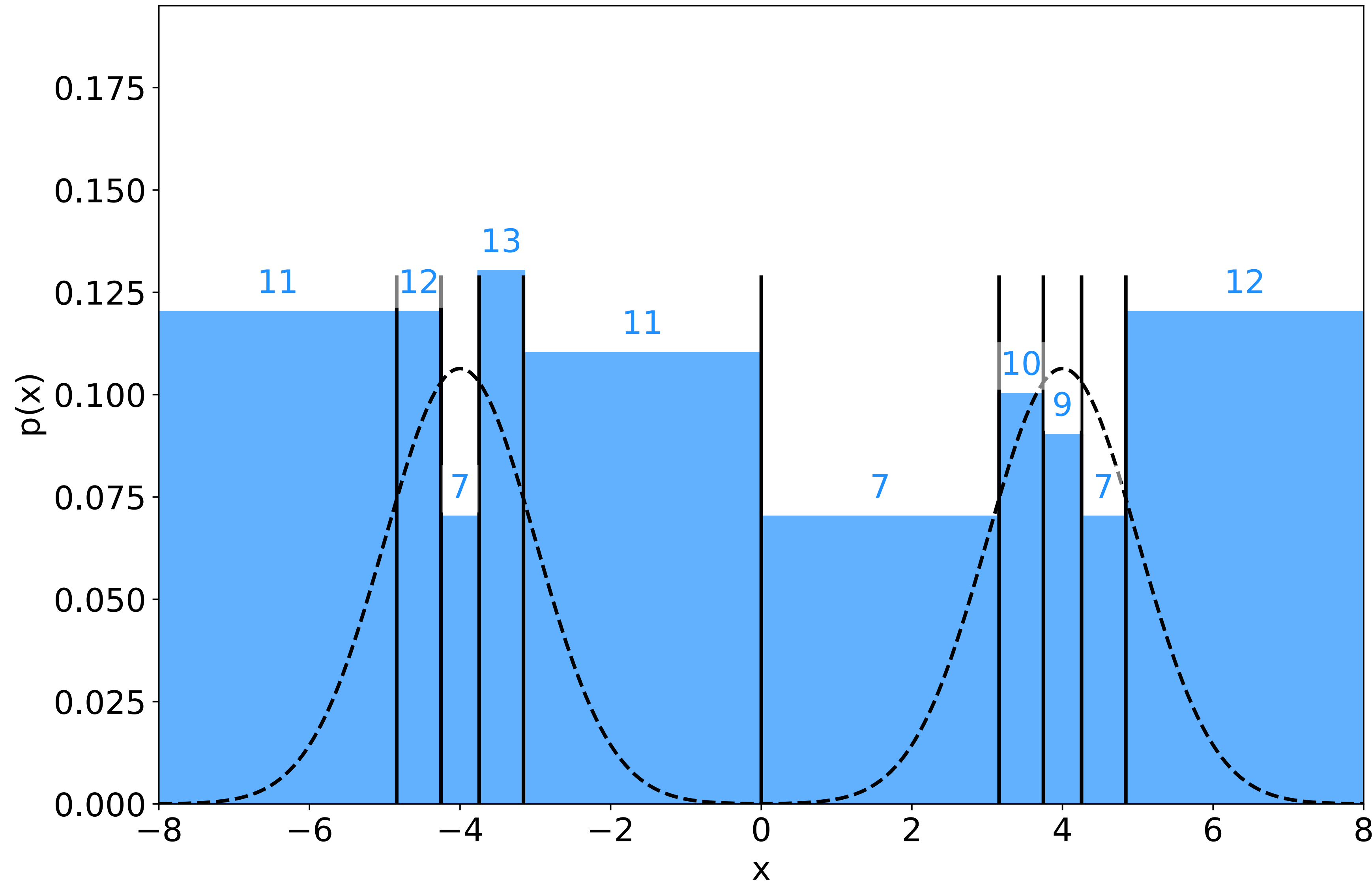
Bonus: How does the χ^2 measure work?



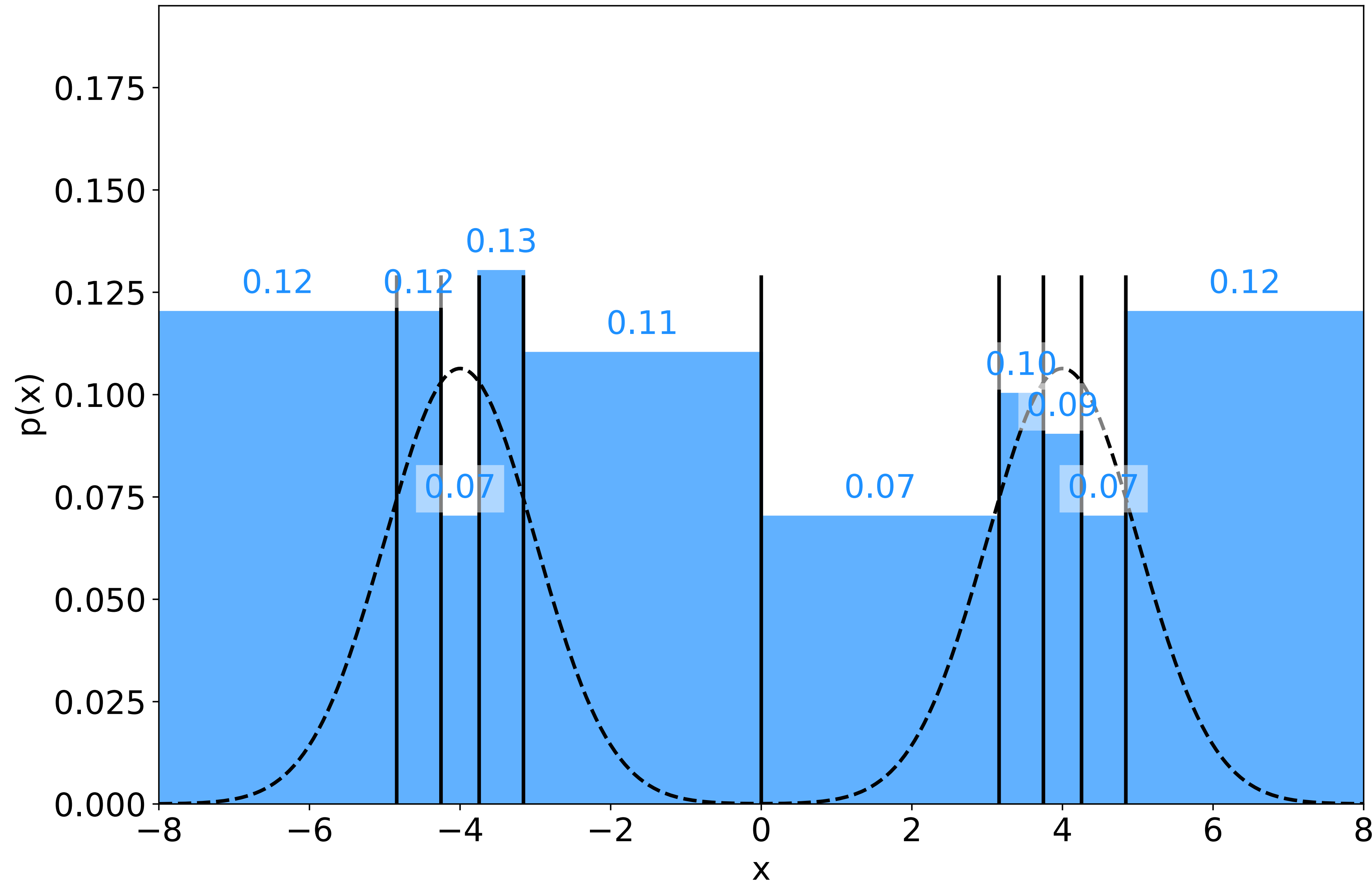
Bonus: How does the χ^2 measure work?



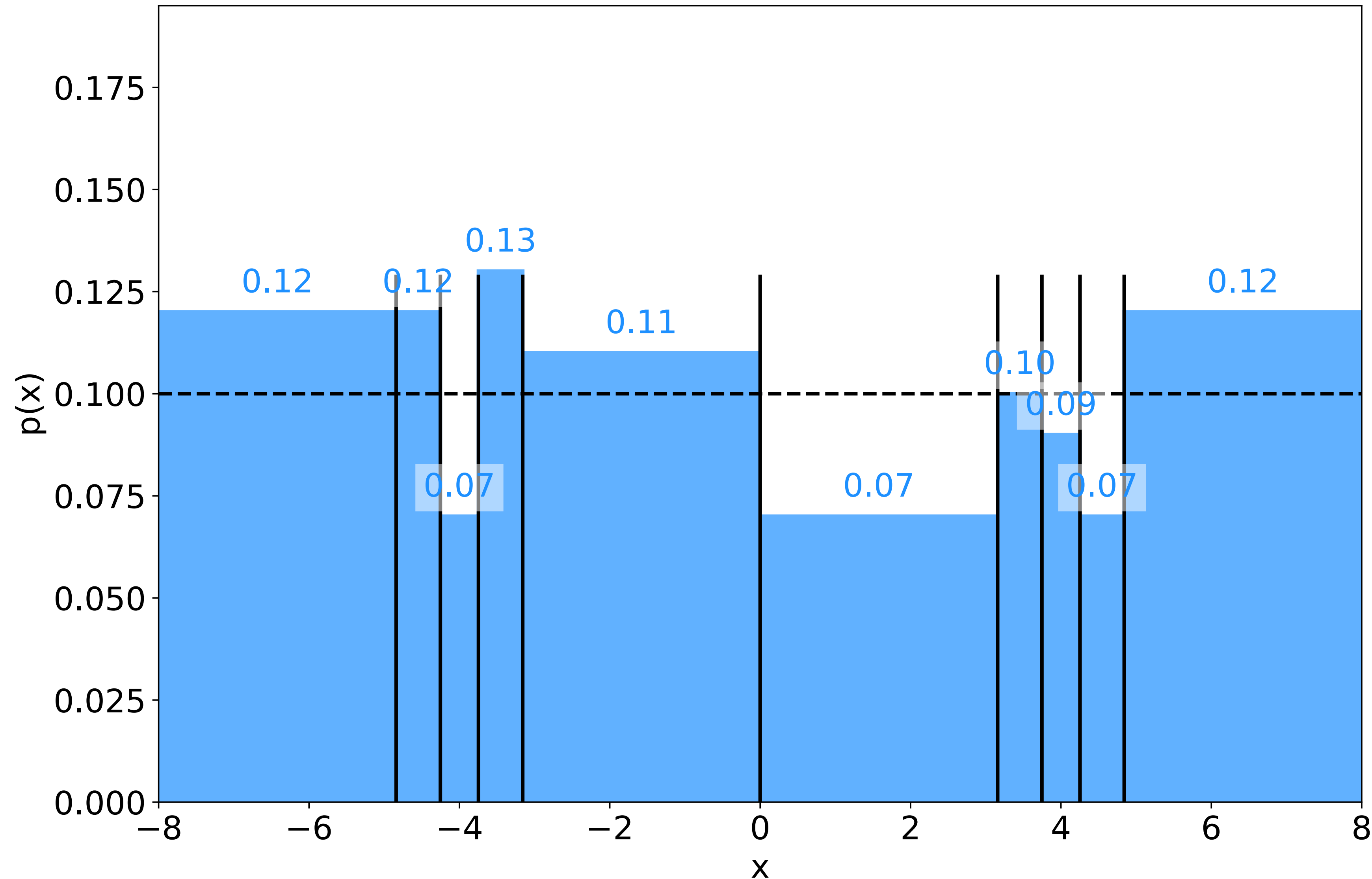
Bonus: How does the χ^2 measure work?



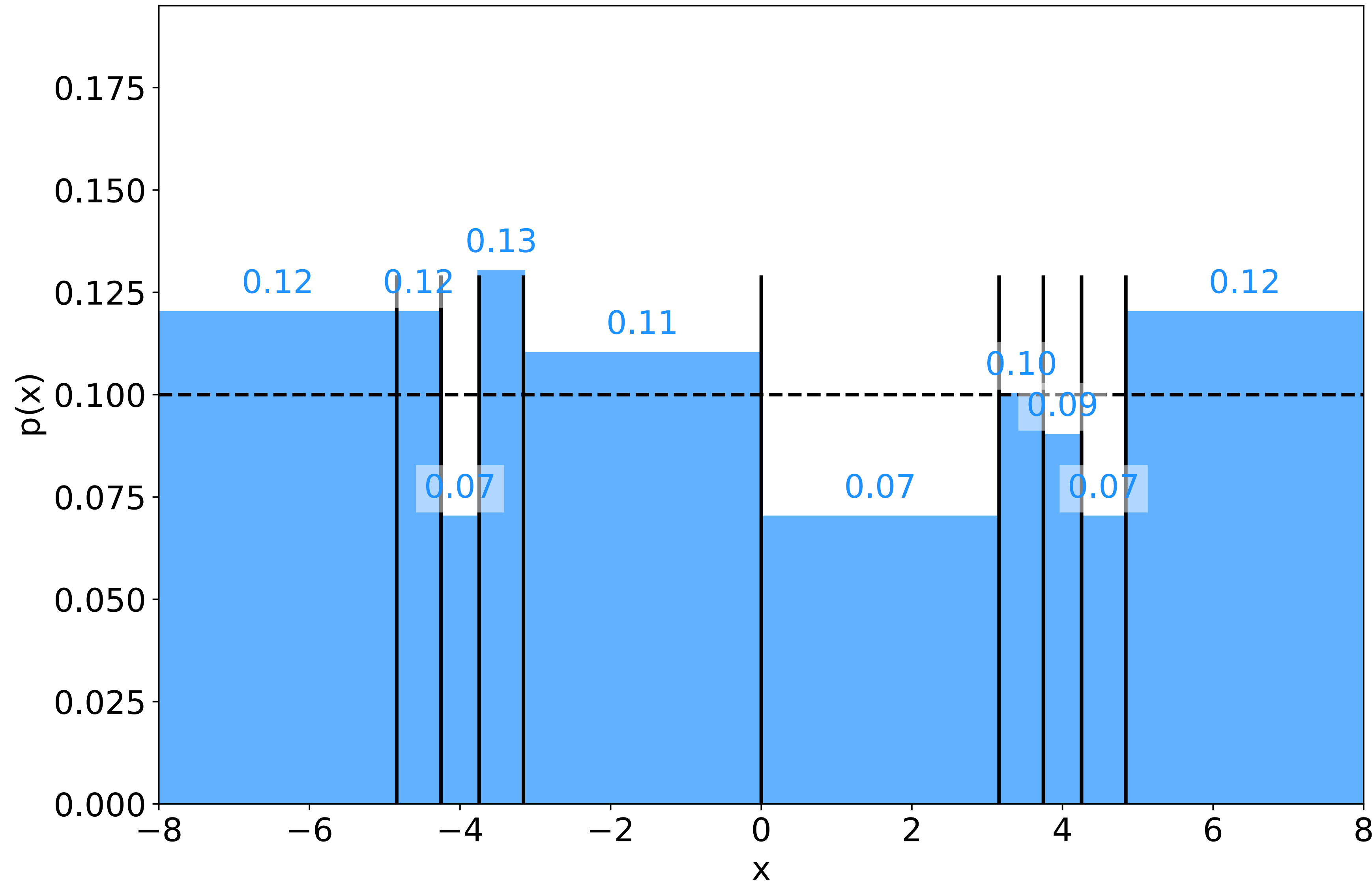
Bonus: How does the χ^2 measure work?



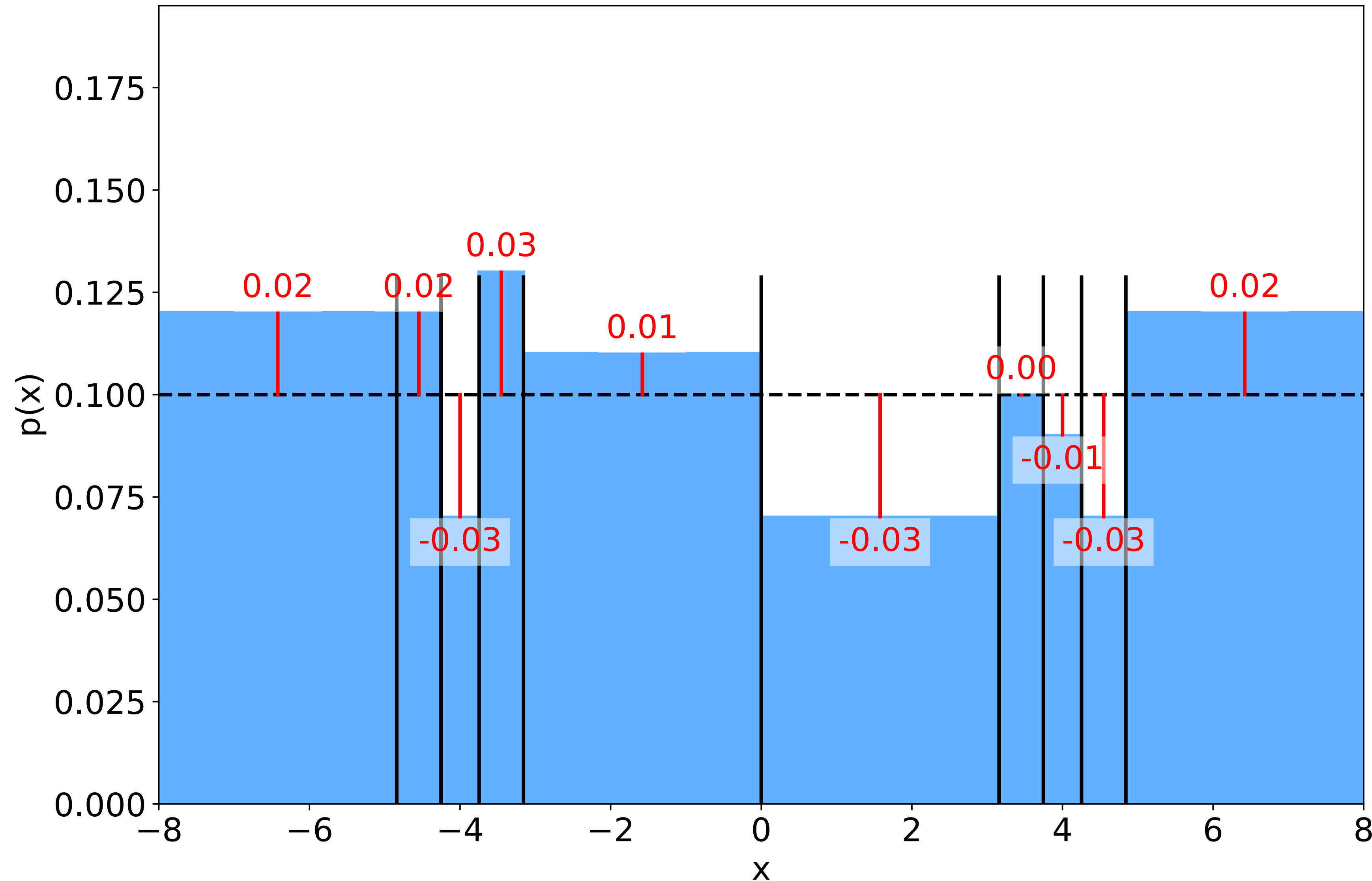
Bonus: How does the χ^2 measure work?



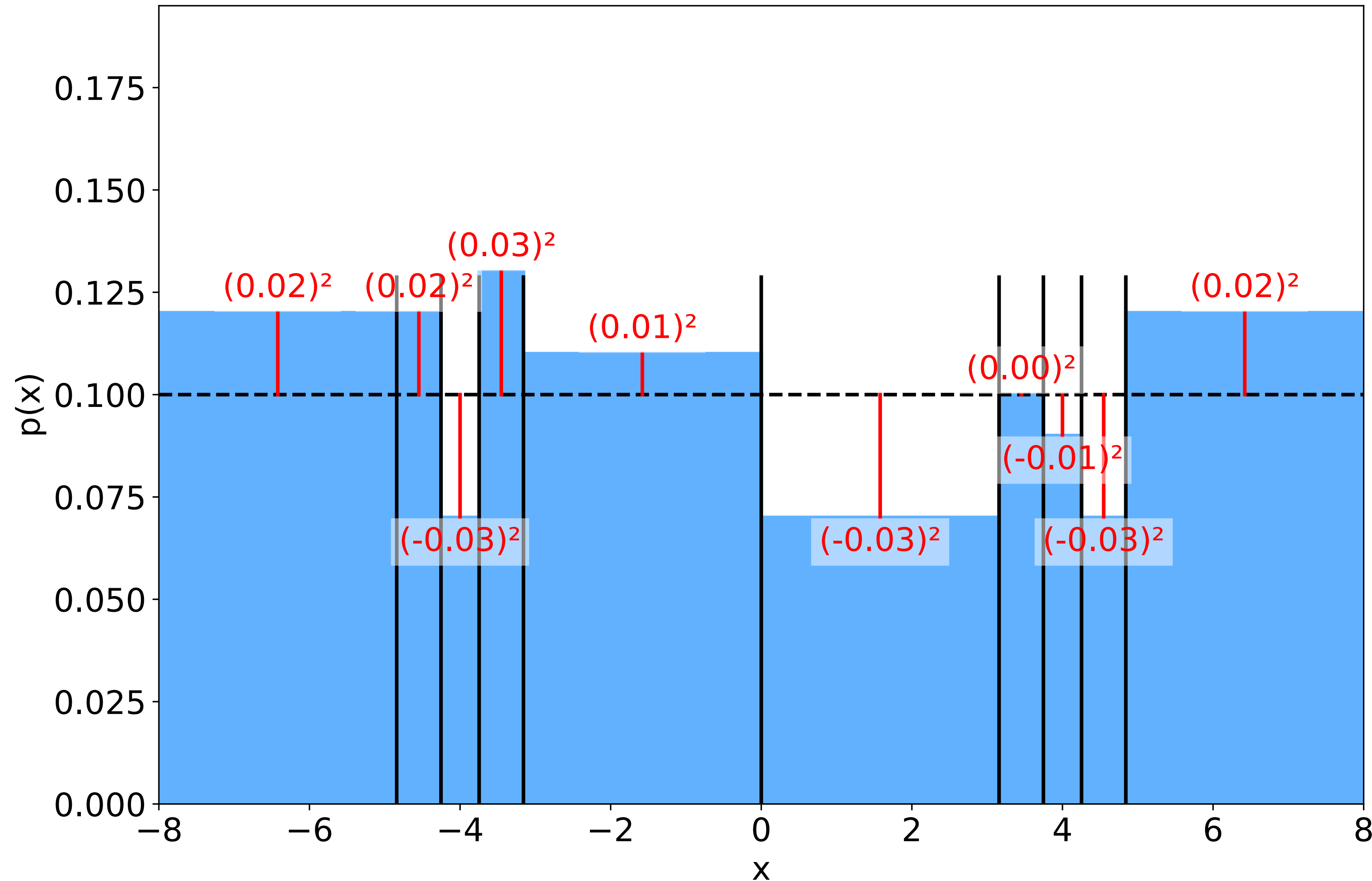
Bonus: How does the χ^2 measure work?



Bonus: How does the χ^2 measure work?

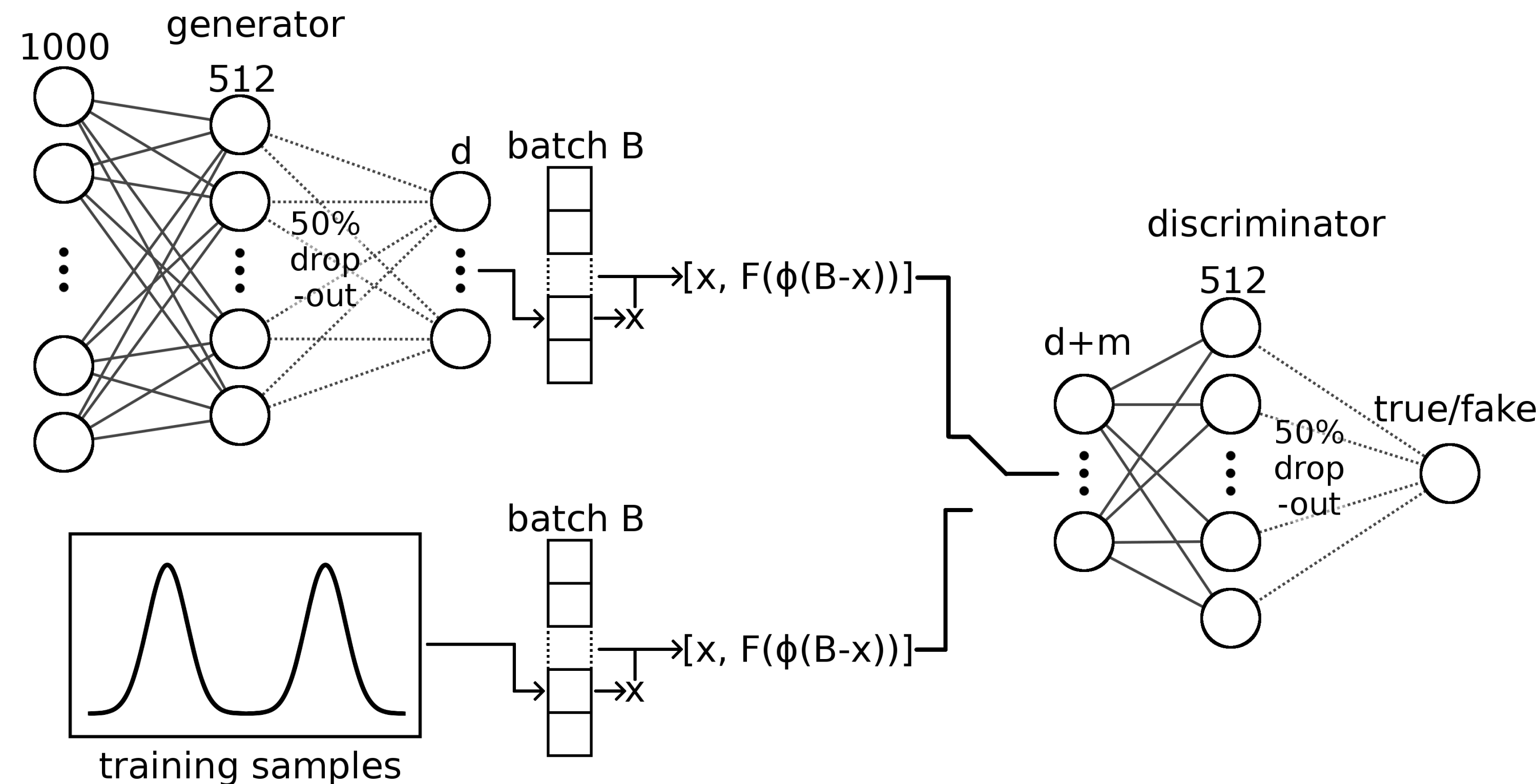


Bonus: How does the χ^2 measure work?



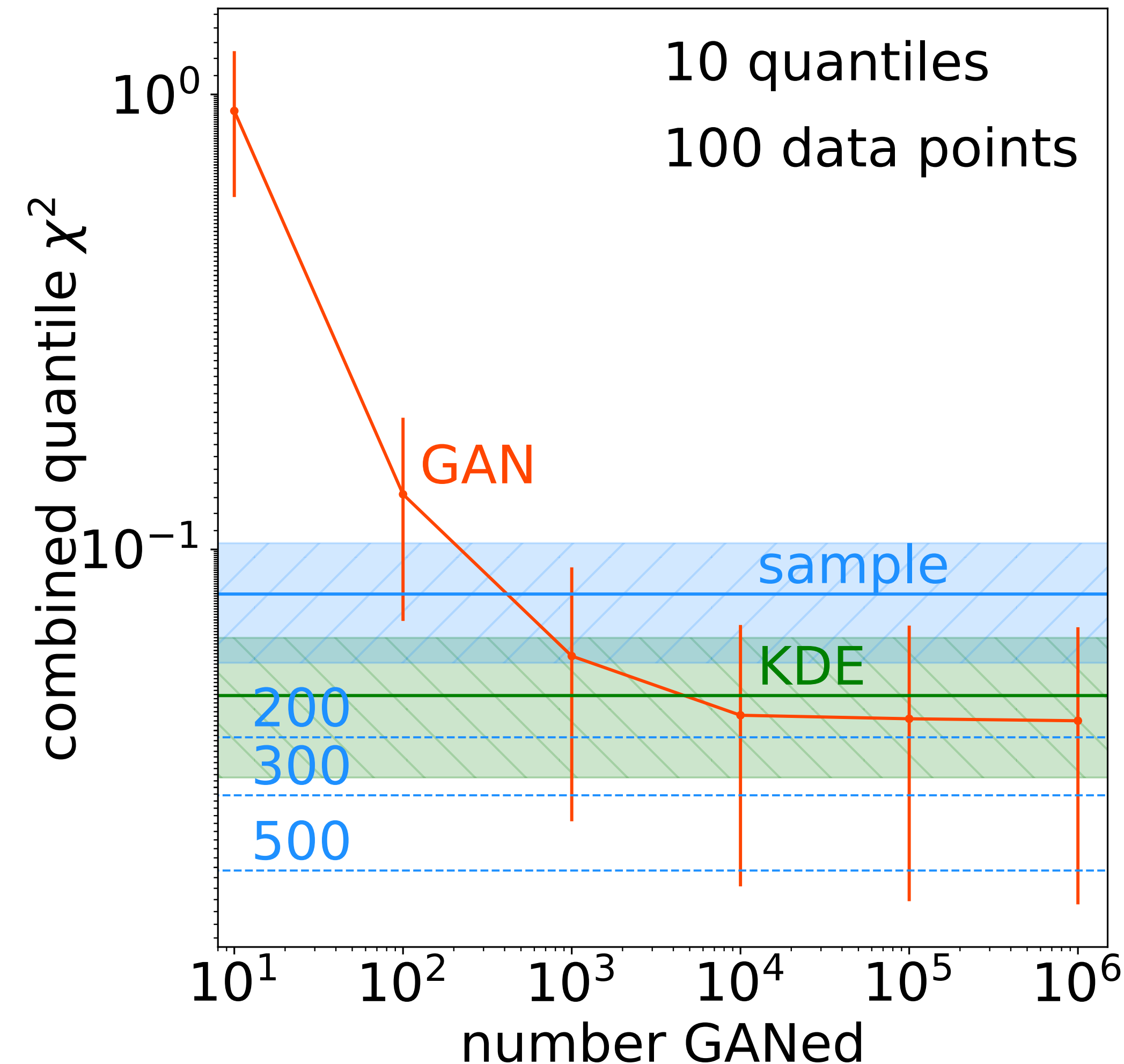
Toy Model: Generative Network

- Train on $n_{\text{data}} = 100$ data points generated from $P(x)$
- Prone to mode-collapse and overfitting:
 - Dropout
 - Noise augmentation
 - Batch-statistics
- Generate high amounts of data from Network



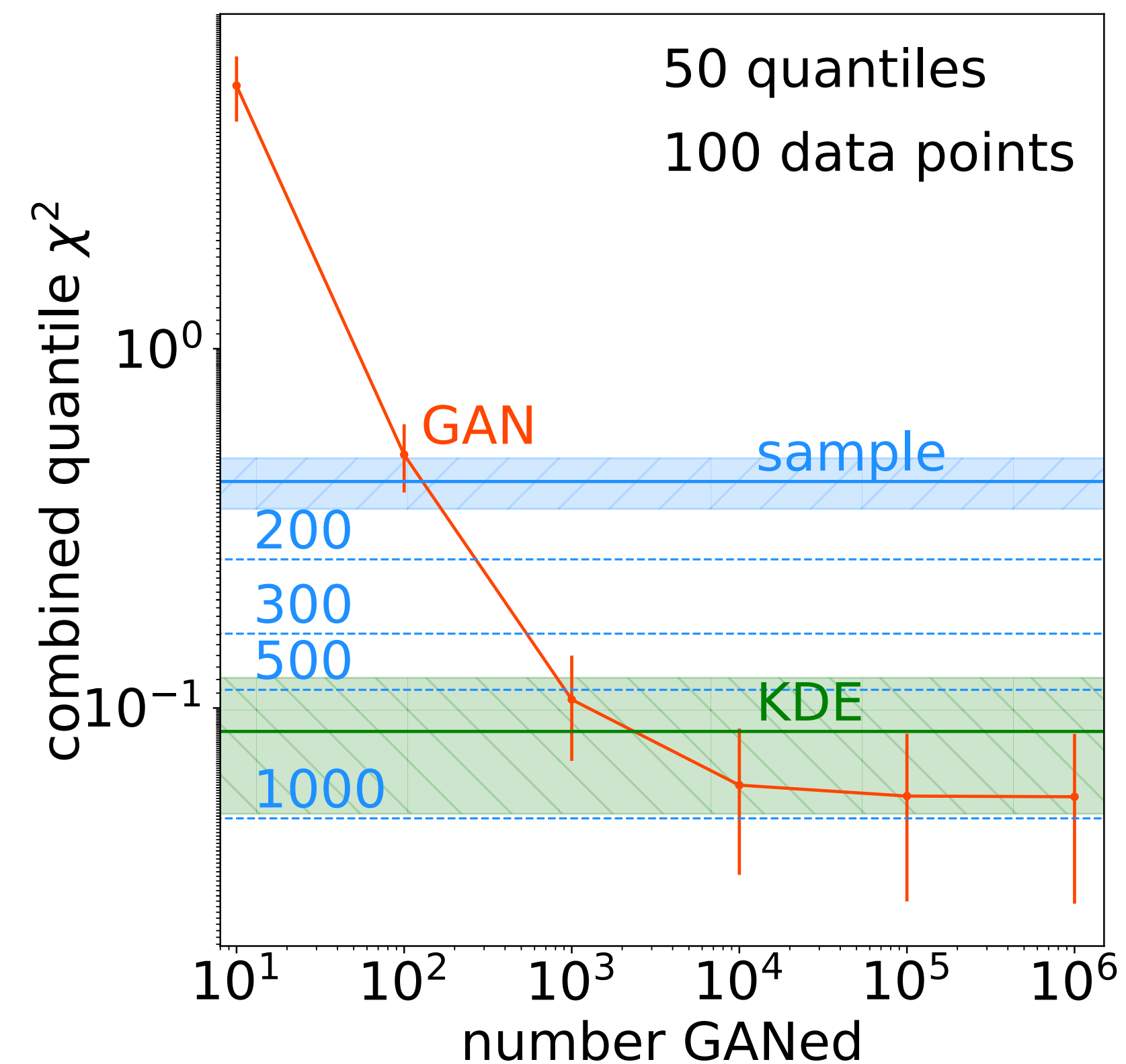
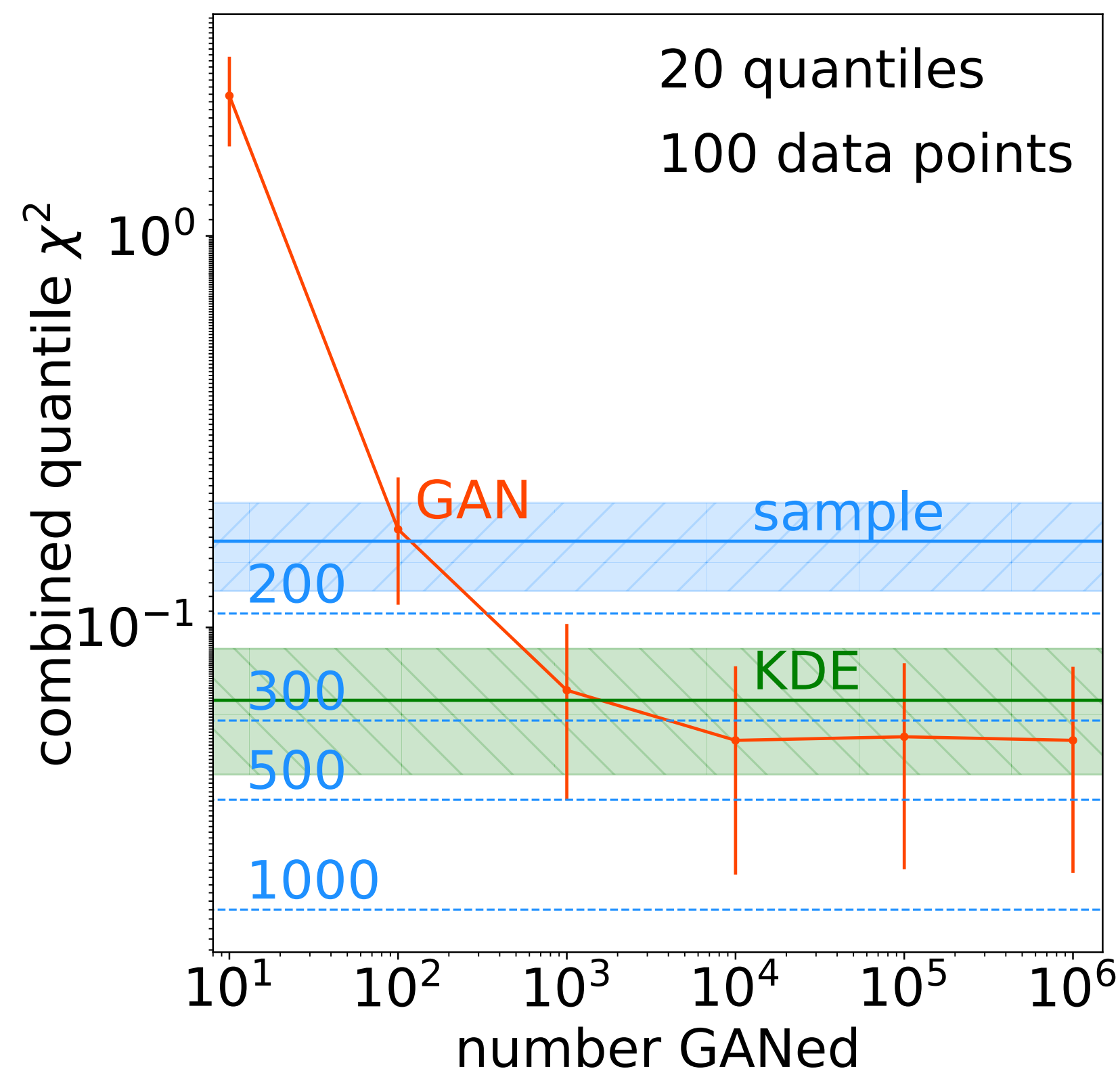
Toy Model: 1D

- GAN (red) and KDE (green) reach higher value than training data
 - sample: only data points
 - KDE: data + smooth, continuous function
 - GAN: data + smooth, continuous function
- 10.000 GANed points match 180 true ones
- Statistical uncertainty of training data becomes systematic uncertainty of the model



Toy Model: 1D

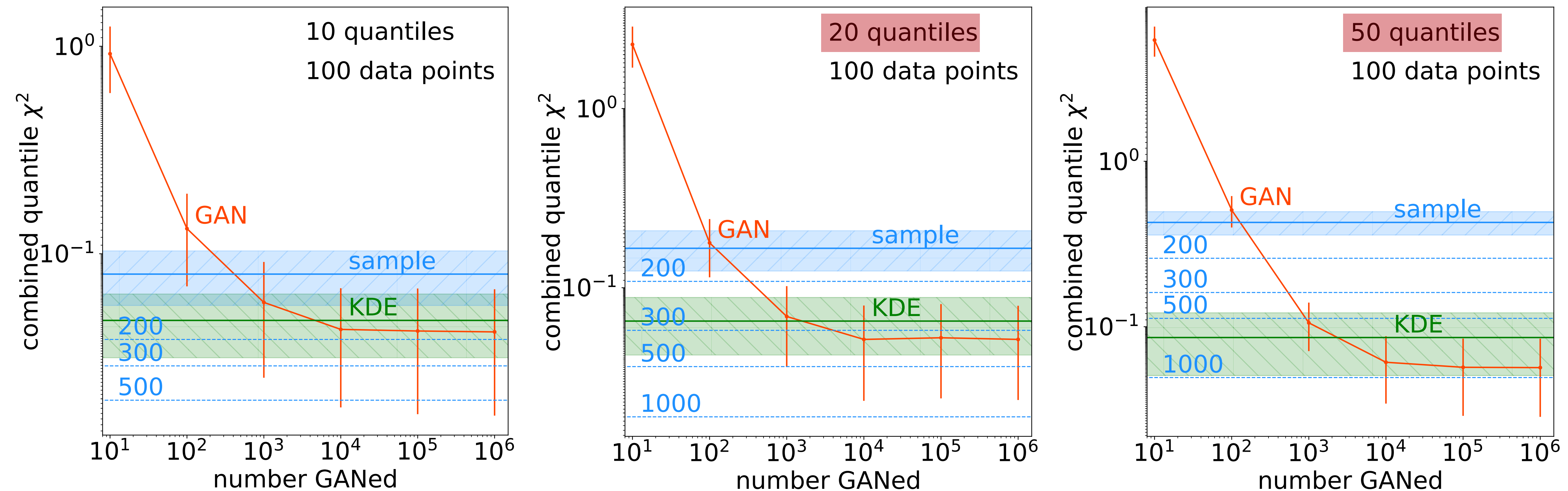
- Interpolation more noticeable for sparse data



- However**, quantile measure breaks down for sparse data

Toy Model: 1D

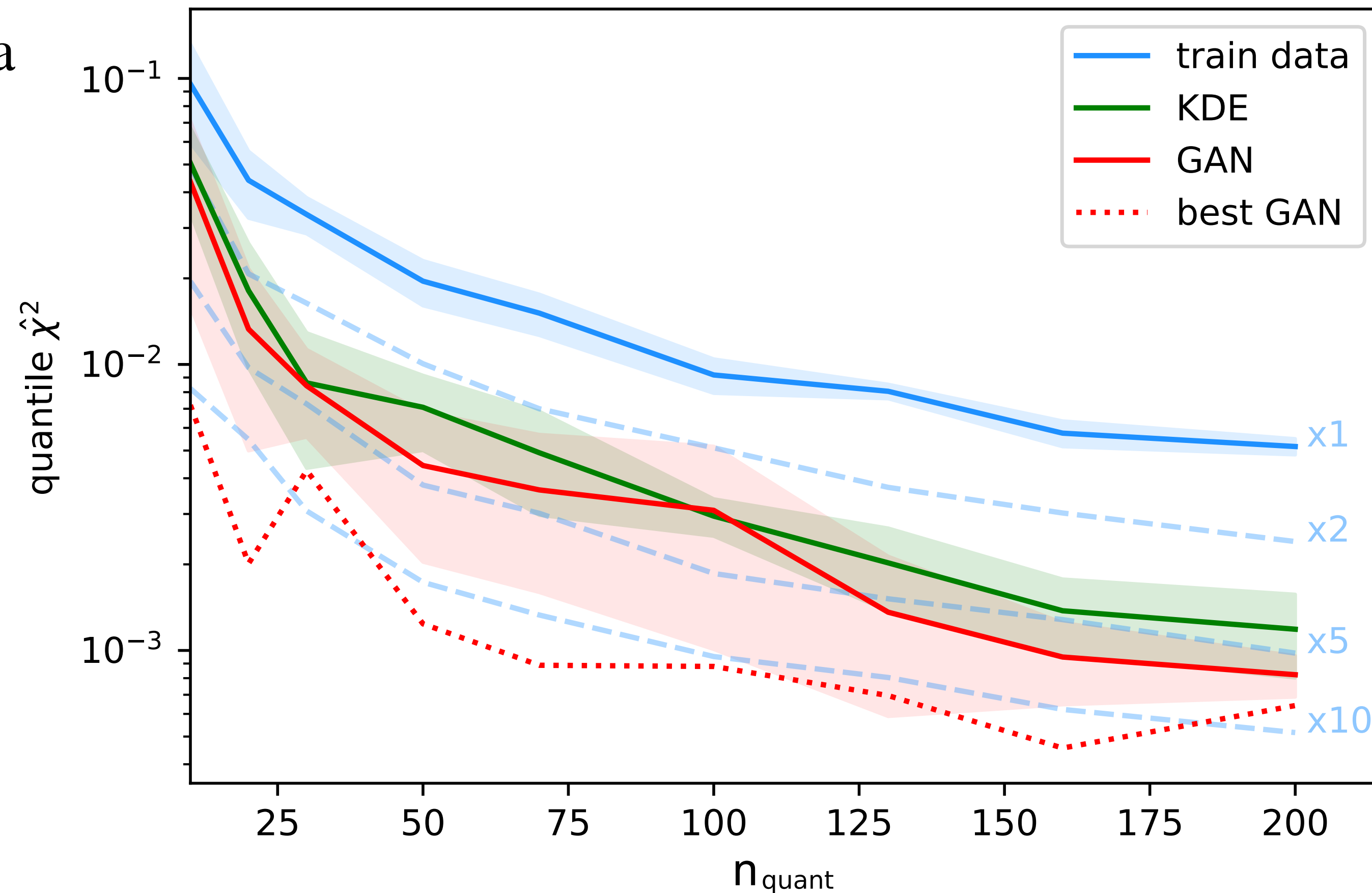
- large n_{quant} $\longrightarrow$ global properties of the fit



- However**, quantile measure breaks down for sparse data

Toy Model: 1D

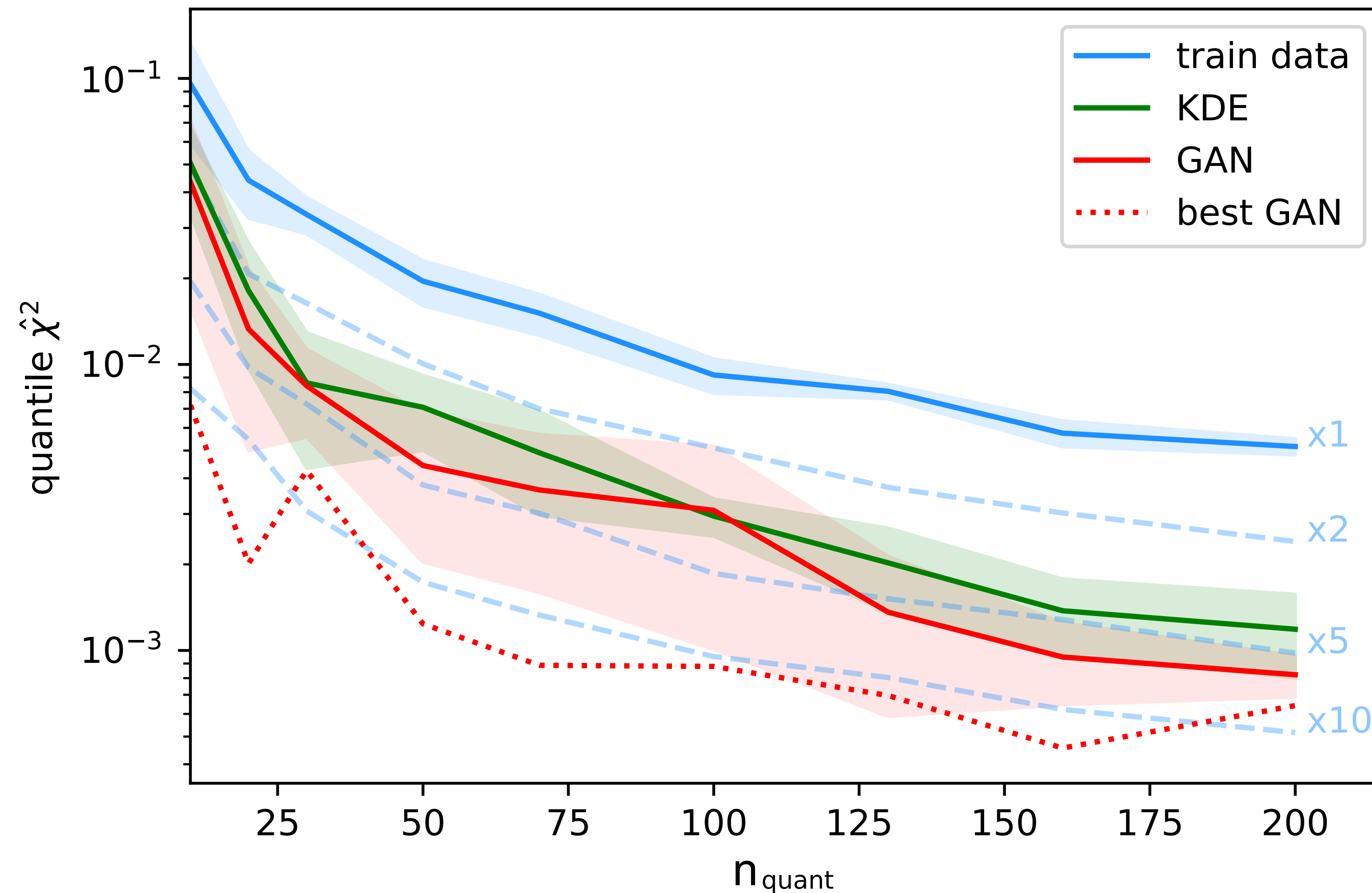
- Examine high n_{quant} and high n_{data}
 - Train on $n_{\text{data}} = n_{\text{quant}}^2$
 - Generate $100 \cdot n_{\text{data}}$
- Examine which data converges to 0 fastest
- GAN amplifies data by a factor ~ 5



Toy Model: 1D

$$\hat{\chi}_{N_{\text{quant}}}^2 \xrightarrow{n_{\text{quant}} \rightarrow \infty} \chi^2(P(x), P_{\text{data}}(x))$$
$$\xrightarrow{n_{\text{data}} \rightarrow \infty} \chi^2(P(x), P(x)) = 0$$

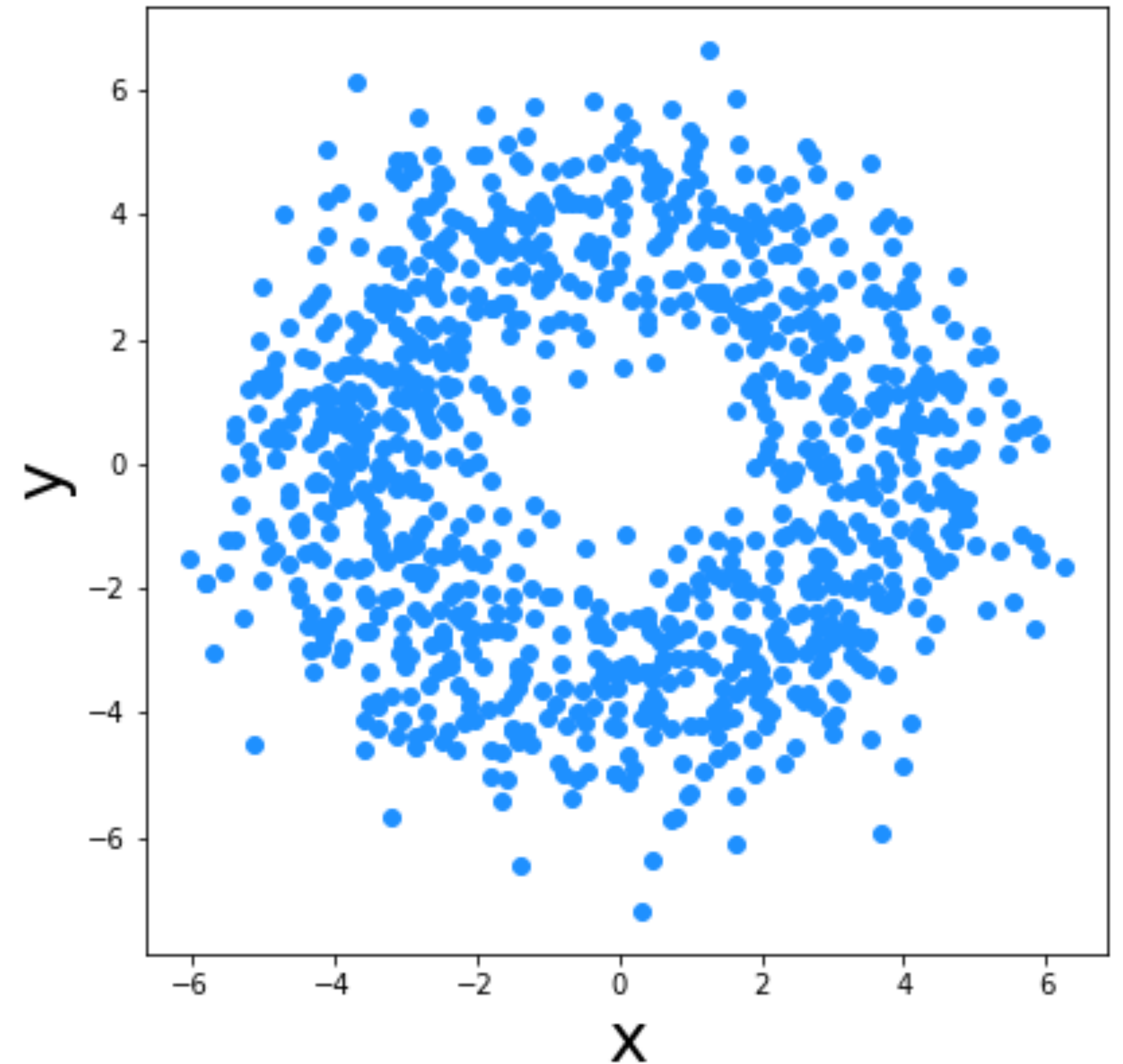
- Examine which data converges to 0 (fastest)
- GAN amplifies data by a factor ~5



Toy Model: 2D

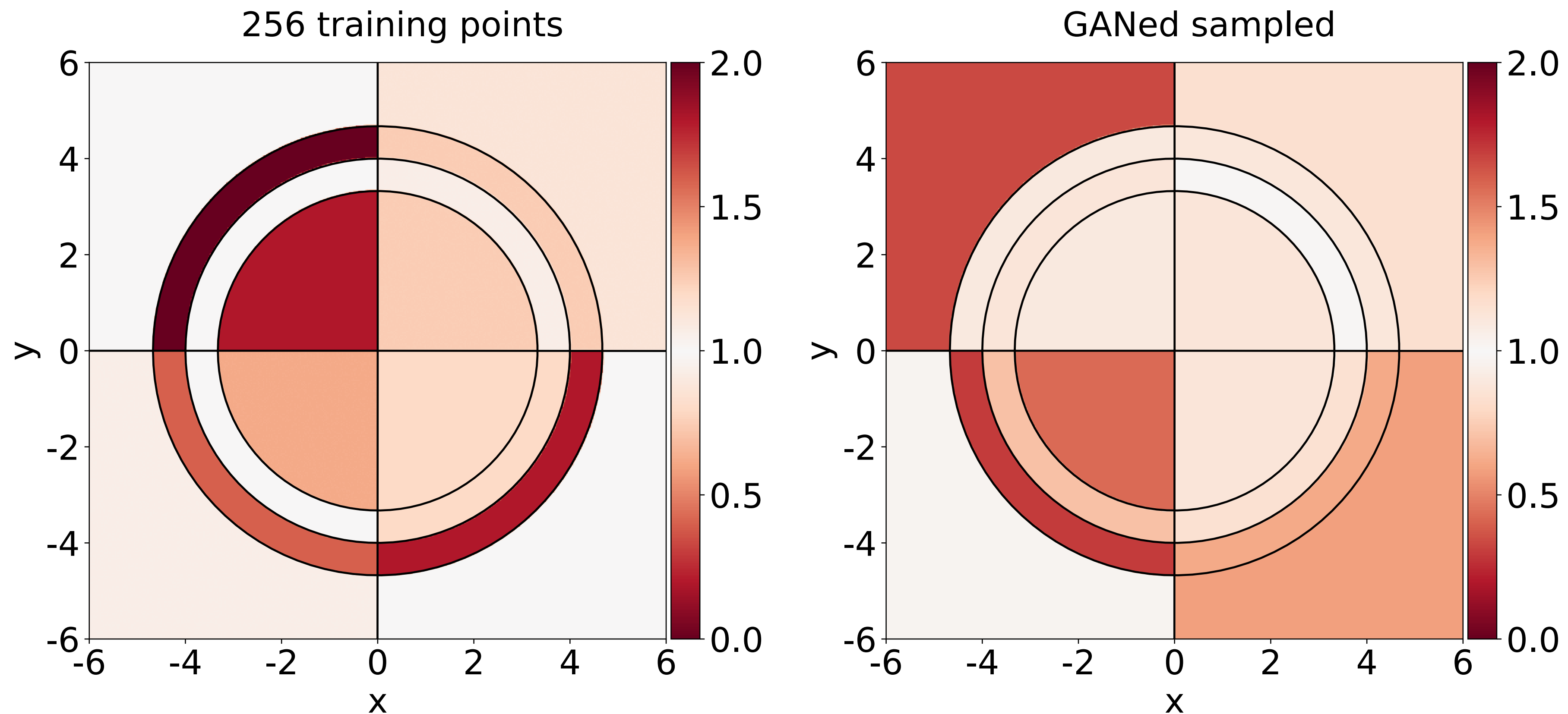
- Ring with gaussian radius
- GAN is trained on cartesian coordinates
- Quantiles are calculated on polar coordinates

GAN has to learn correlations



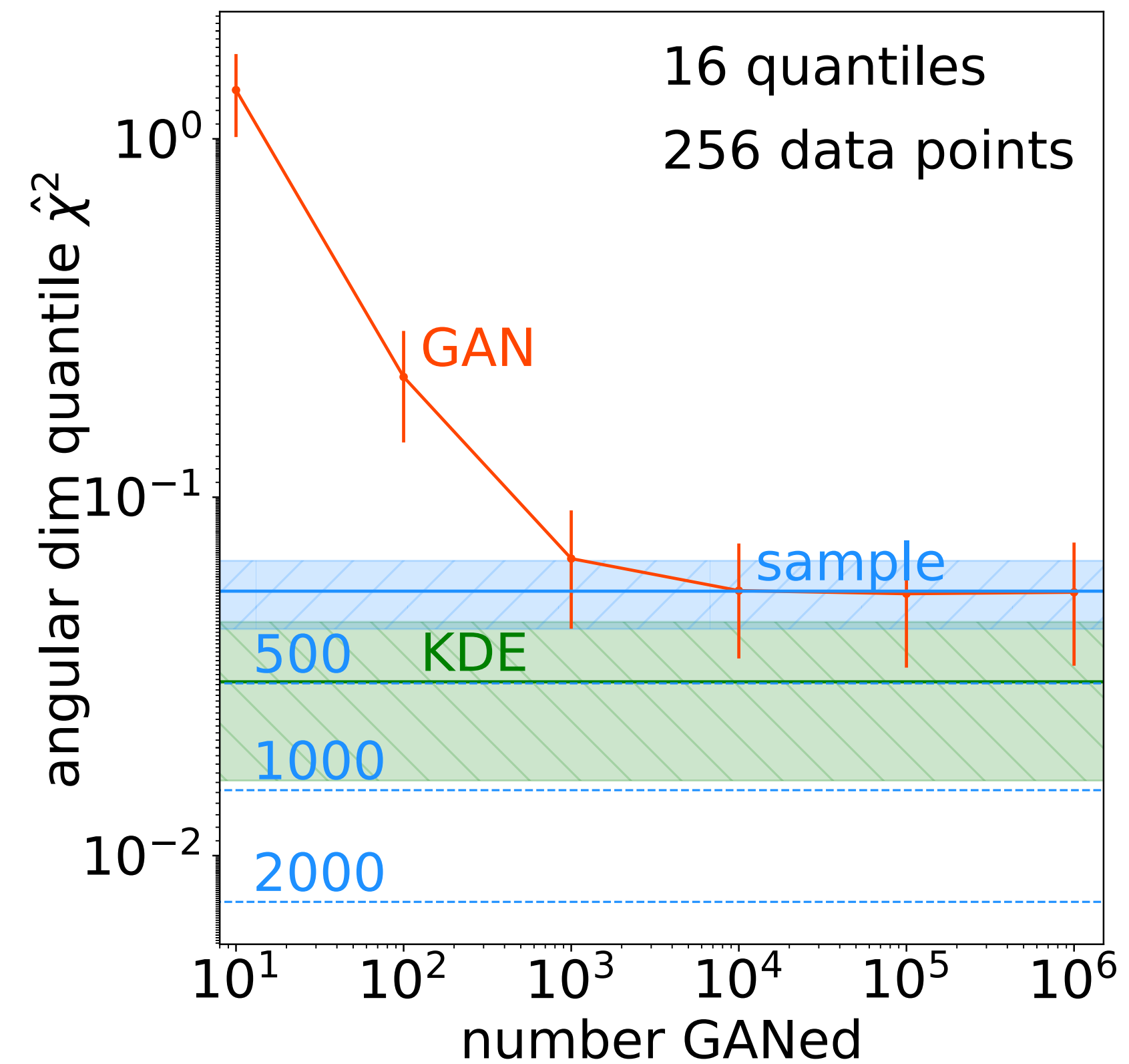
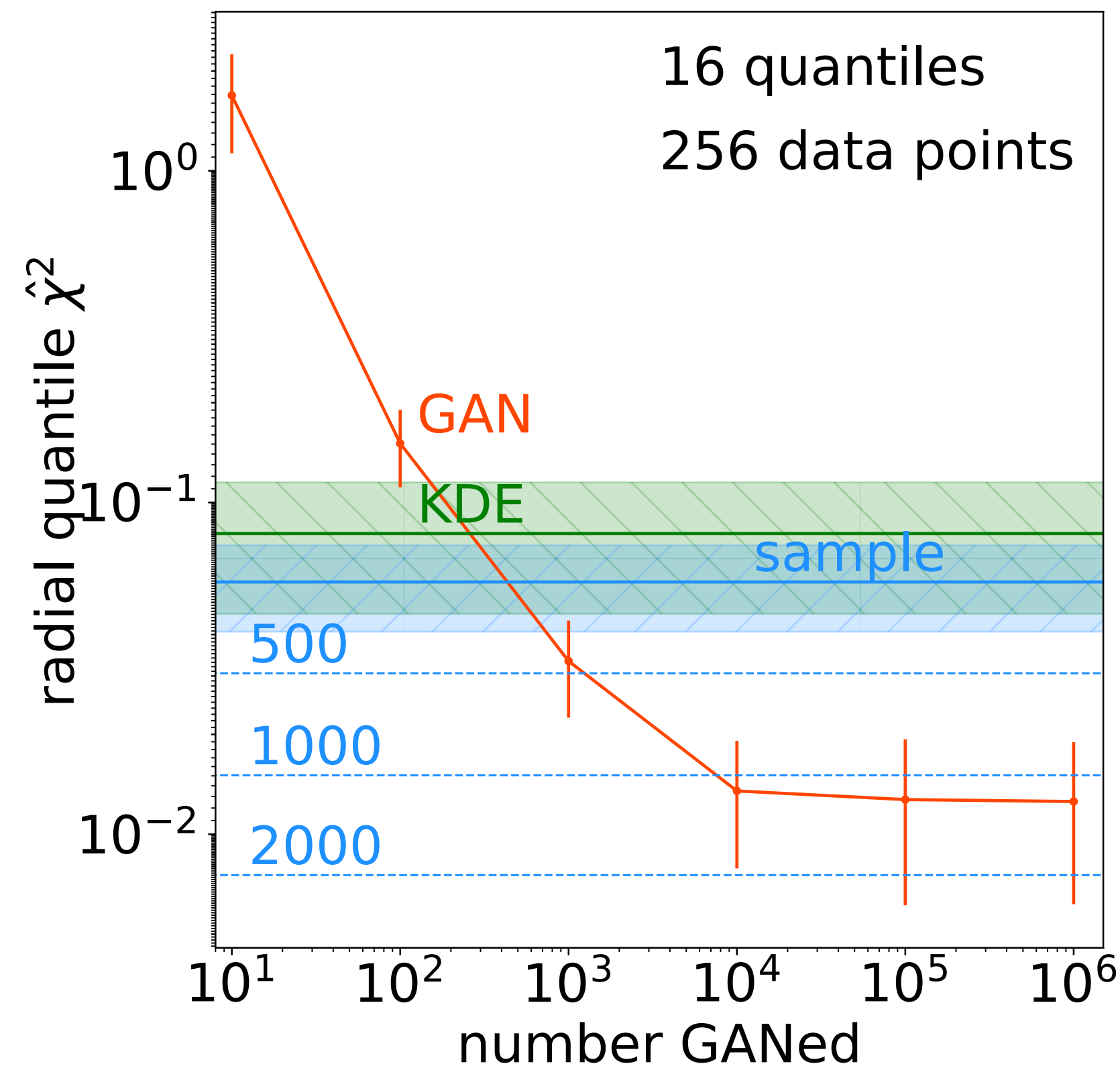
Toy Model: 2D

- Quantiles in radial and angular direction



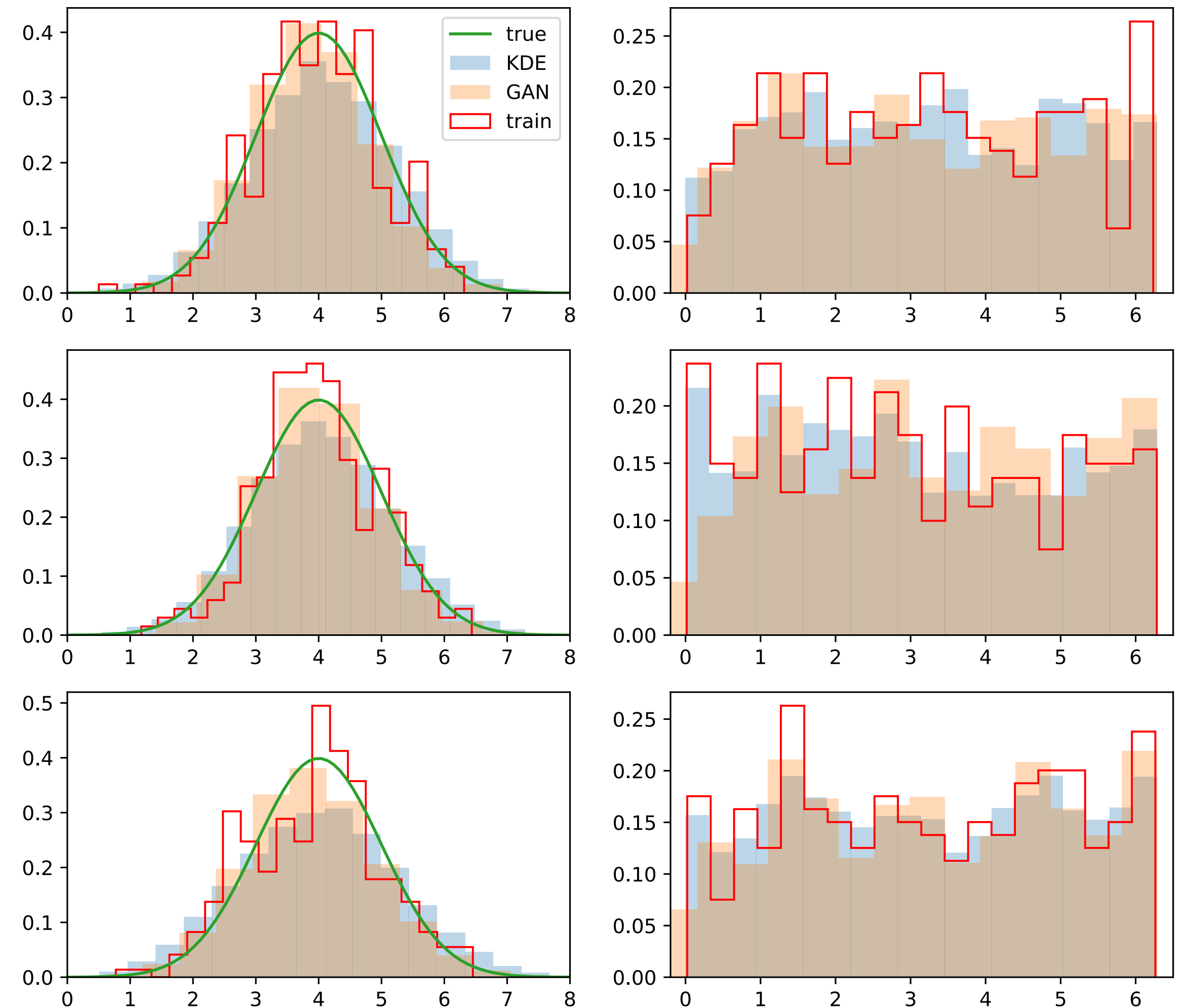
Toy Model: 2D

- Quantiles in radial and angular direction

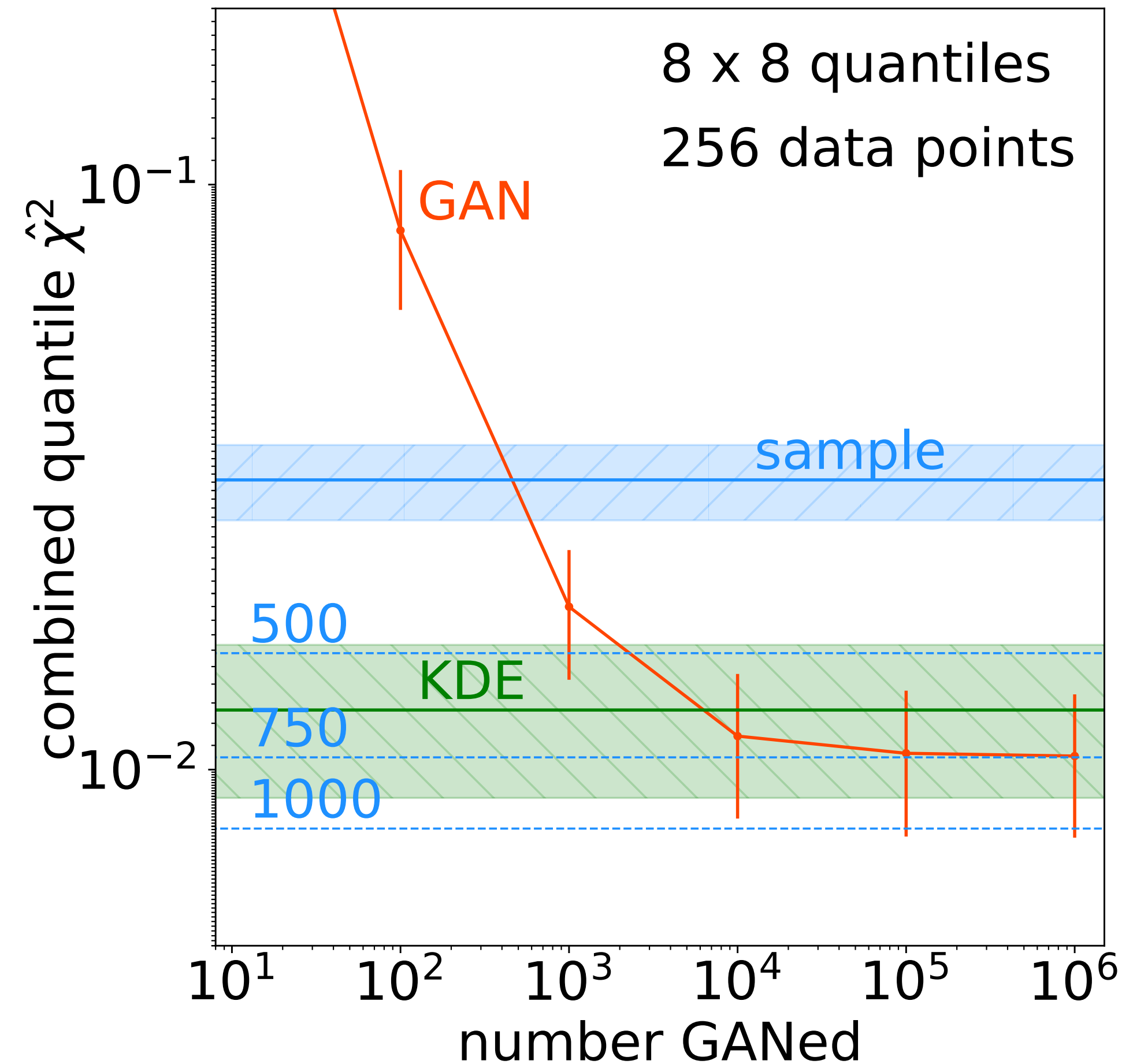
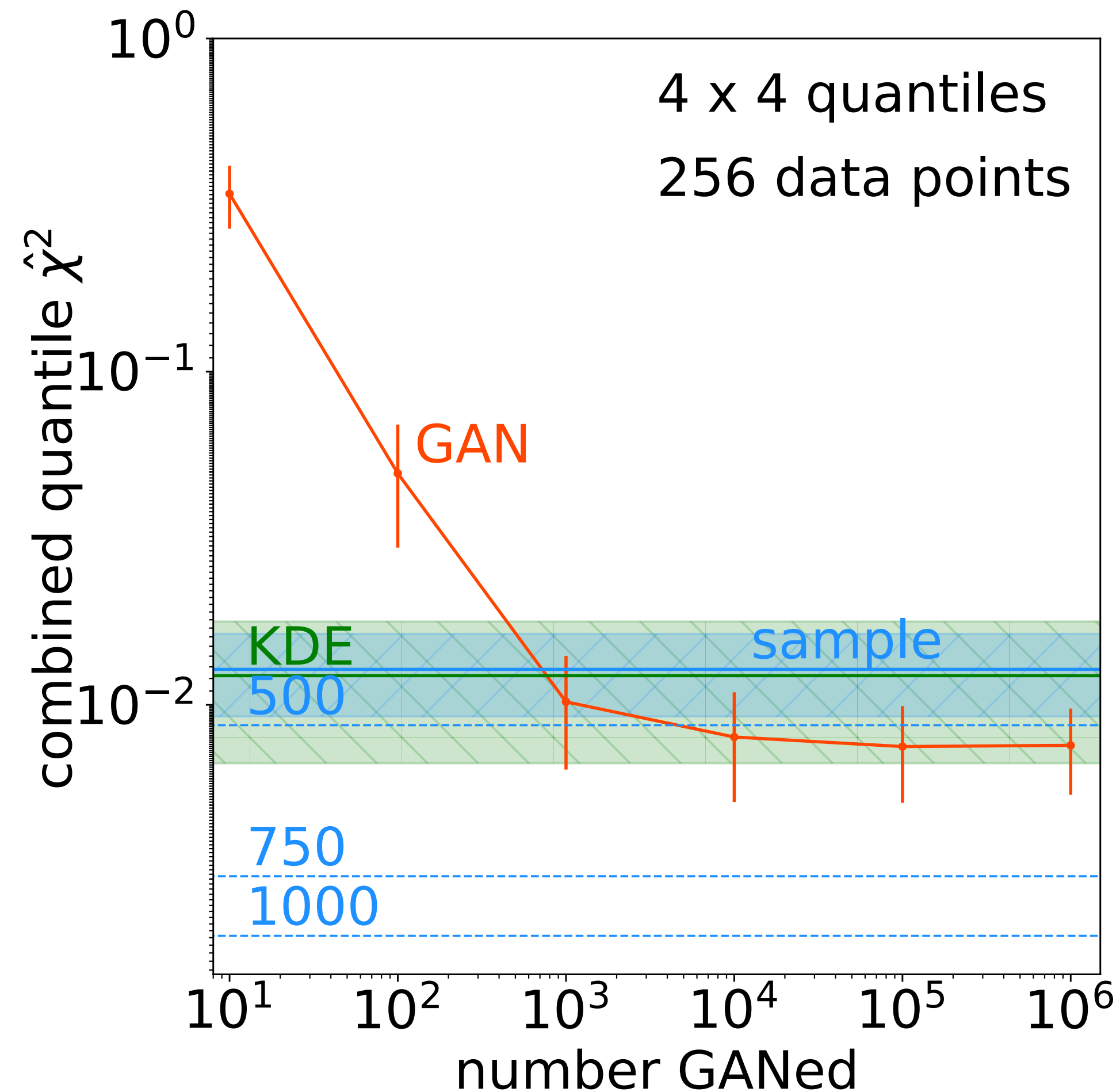


Bonus: Why is the KDE χ^2 so bad?

- KDE „underfits“ the gaussian radial distribution
- GAN „overfits“ the uniform angular data



Toy Model: 2D

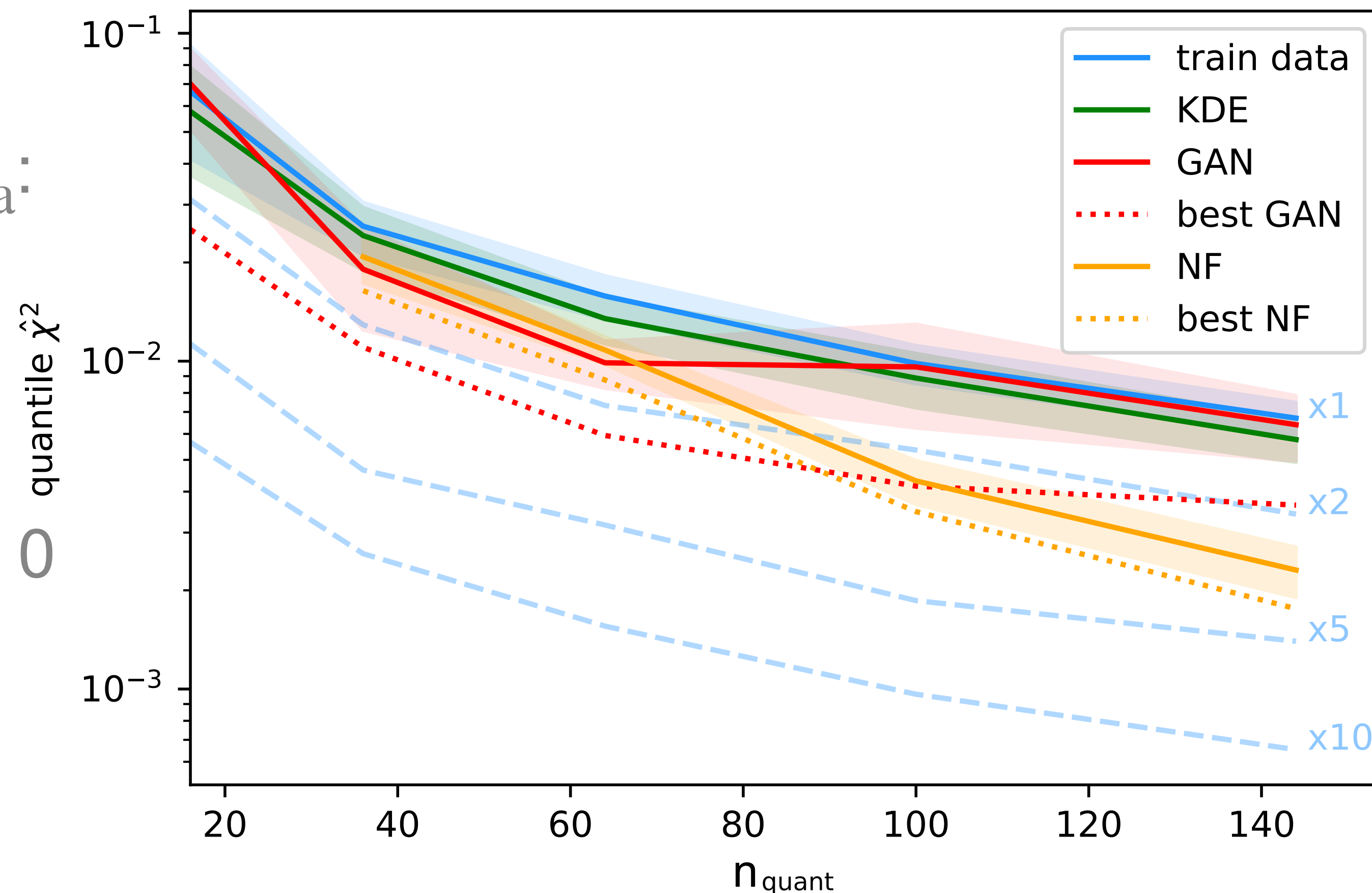


- Unsurprisingly, the dependence on the number of quantiles persists

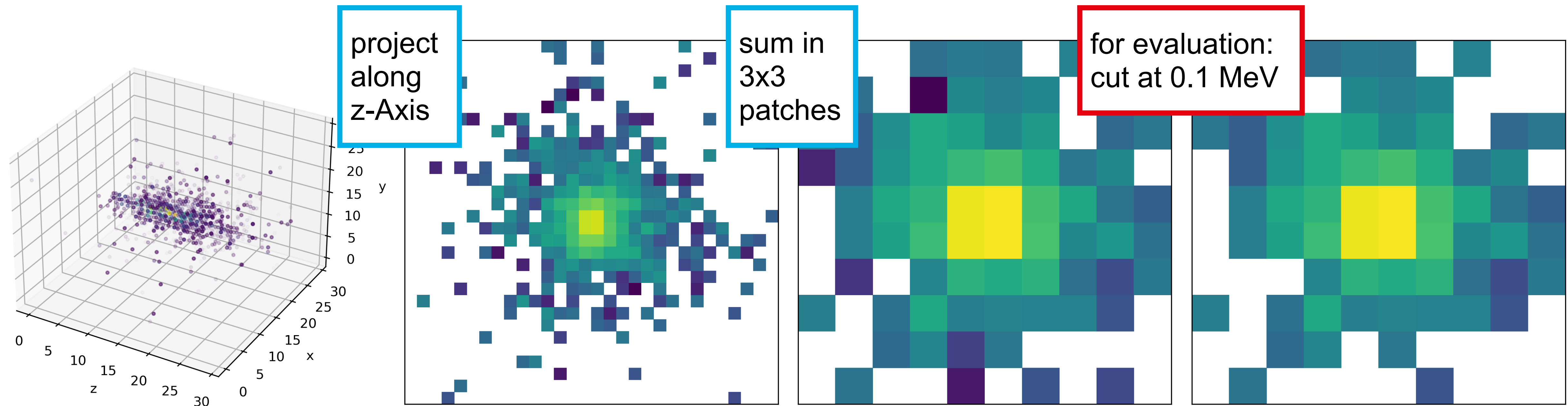
Toy Model: 2D

Do the same thing again:

- Examine high n_{quant} and high n_{data} :
 - Train on n_{quant}^2 data points
 - Generate $100 \cdot n_{\text{data}}$
- Examine which data converges to 0 (fastest)



Calorimeter Simulations: Data



- 269k photon showers at 50 GeV in International Large Detector [1]

Image-shaped data

Unknown true distribution, limited data

Harder learning task → training on multiple training set sizes unfeasible

Calorimeter Simulations: Architecture

- Change to location-aware VAE-GAN architecture → [2202.07352 \[hep-ph\]](#) [2]

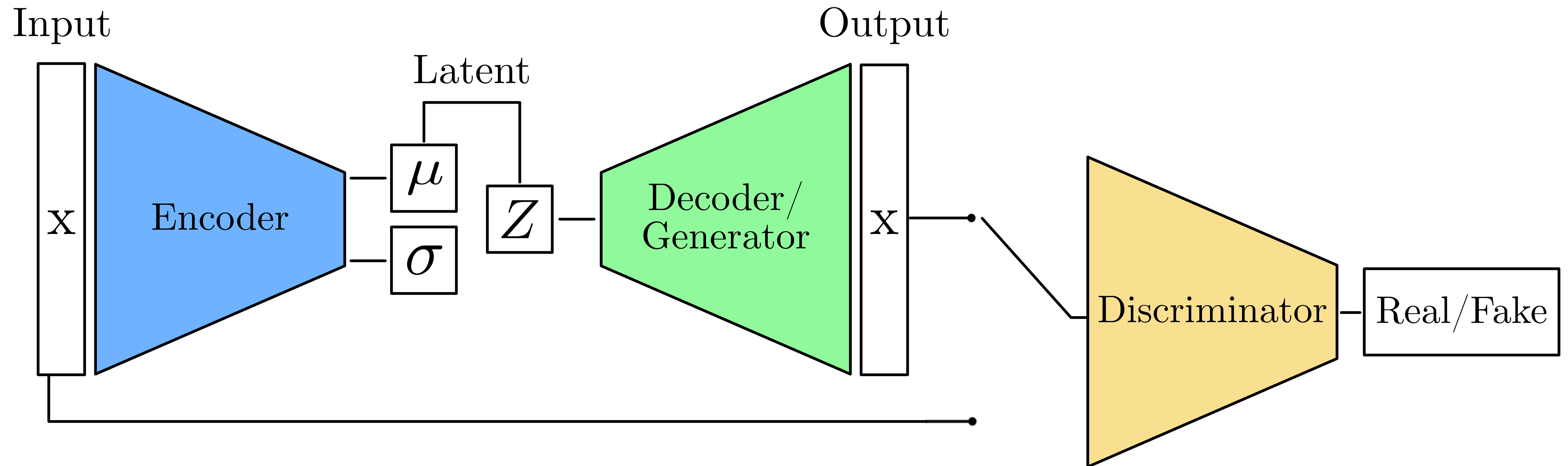


Image-
shaped data

Unknown true
distribution, limited data

Harder learning task → training on
multiple training set sizes unfeasible

Calorimeter Simulations: Architecture



- Change to *location aware* VAE-GAN architecture [2,3,4]

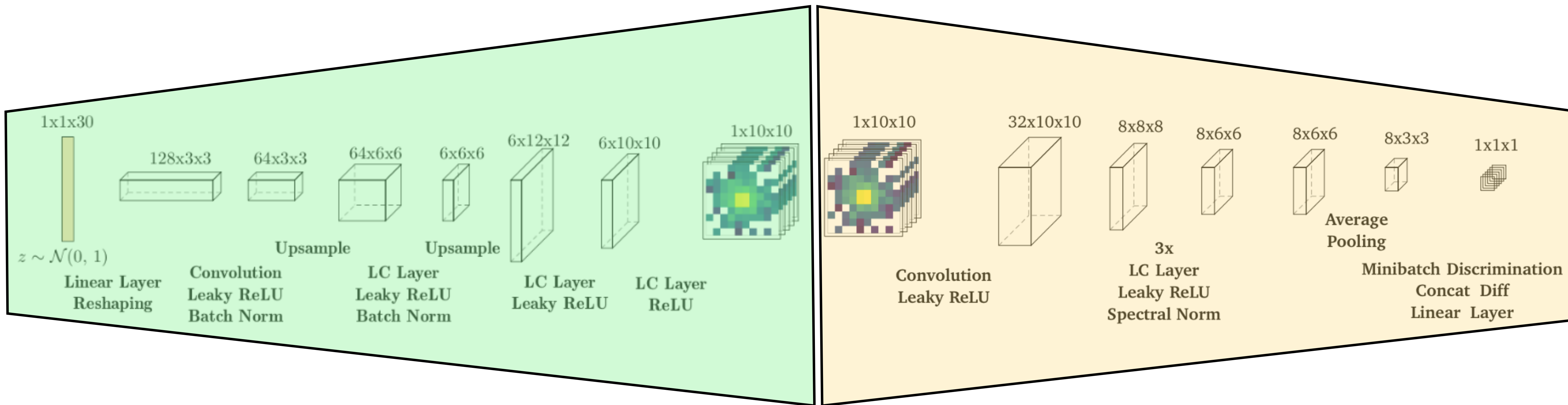
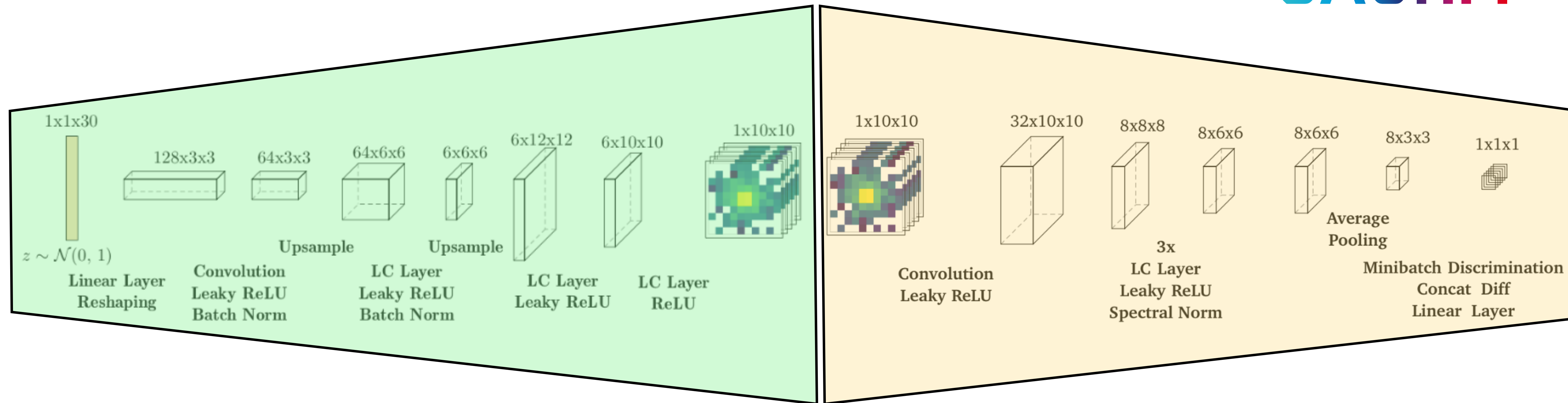


Image-shaped data

Unknown true distribution, limited data

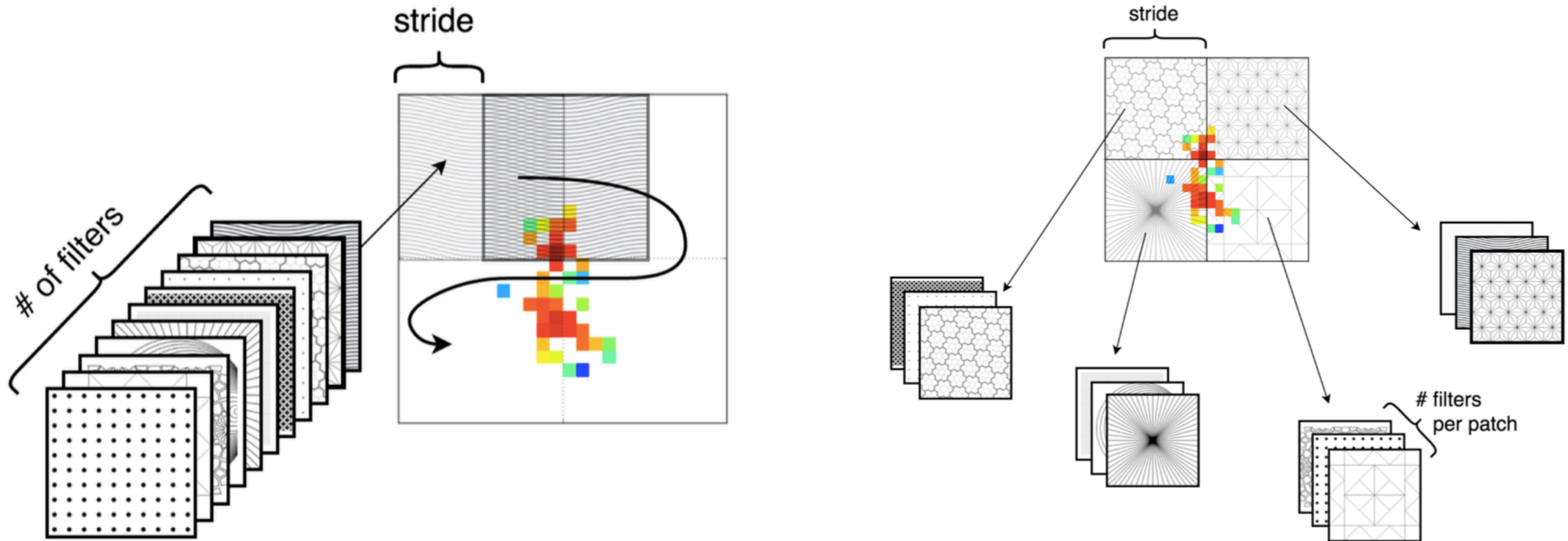
Harder learning task → training on multiple training set sizes unfeasible

Bonus: VAE-LAGAN Architecture



- Change to *location aware* VAE-GAN architecture [2,3,4]:
 - discriminator is trained twice per generator batch
 - learning rate 8×10^{-6}
 - 24h training ~ 50000 epochs of the 1k dataset
 - model selection using JSD to training data

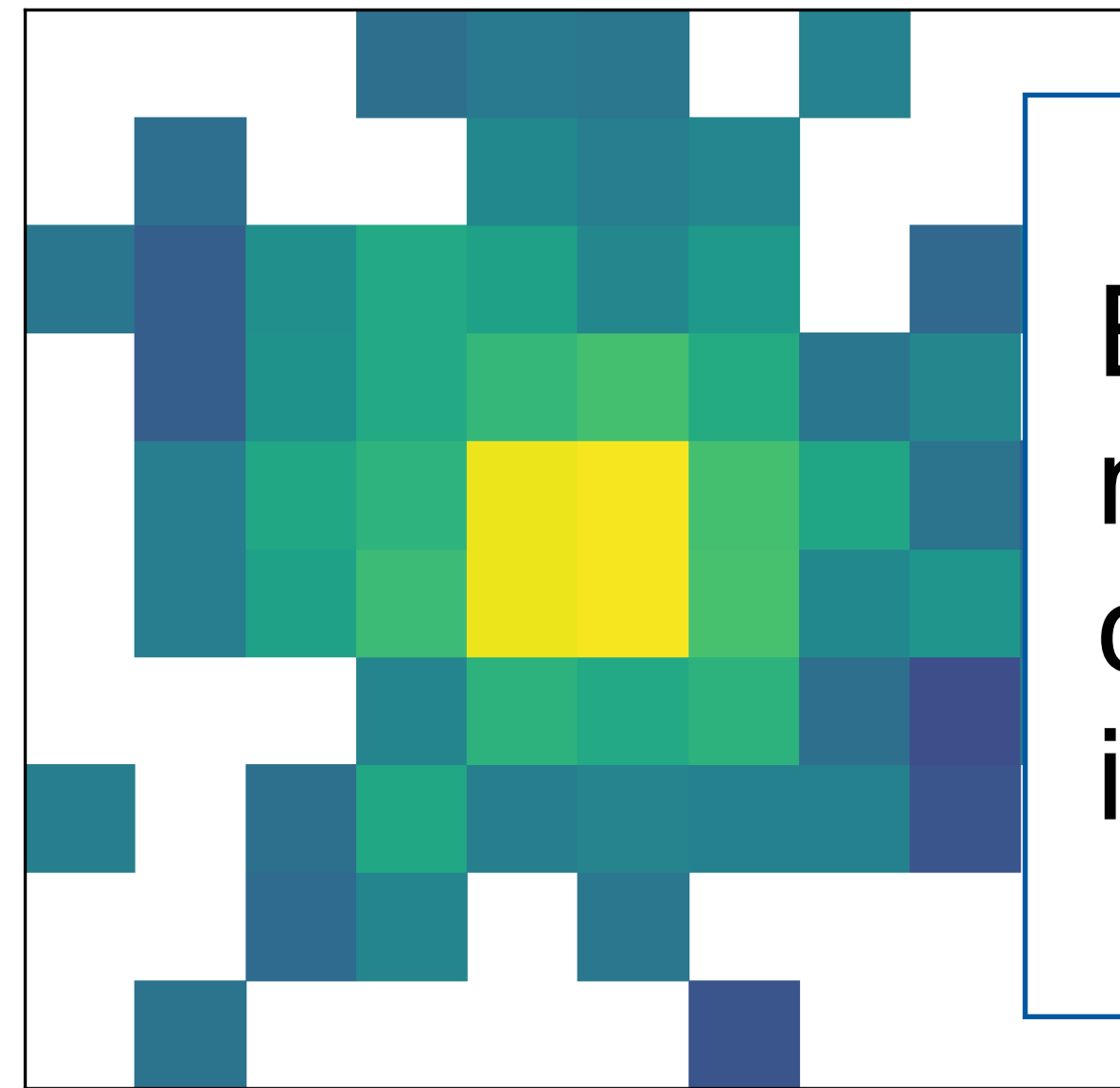
Bonus: VAE-LAGAN Architecture



Convolutional Layers (left) use the same filters for every position of the image, Locally Connected (LC) Layers (right) use different filters [2].

Calorimeter Simulations: Setup

DASHH



Evaluate 1D
metrics,
calculated on
images

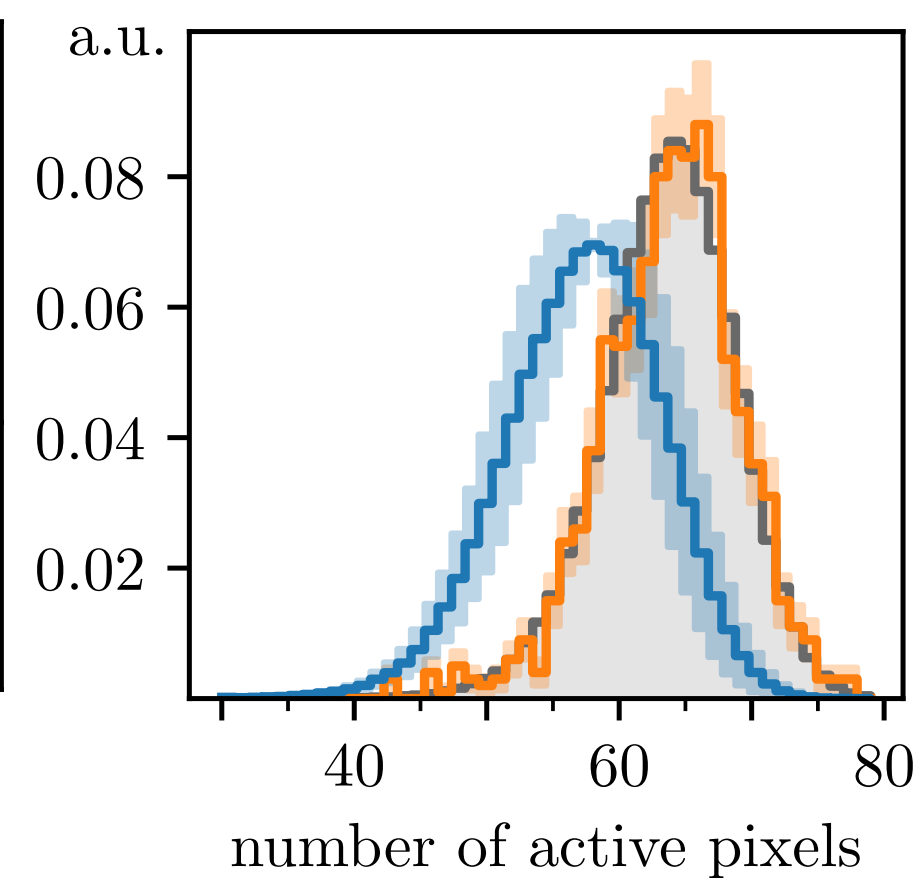
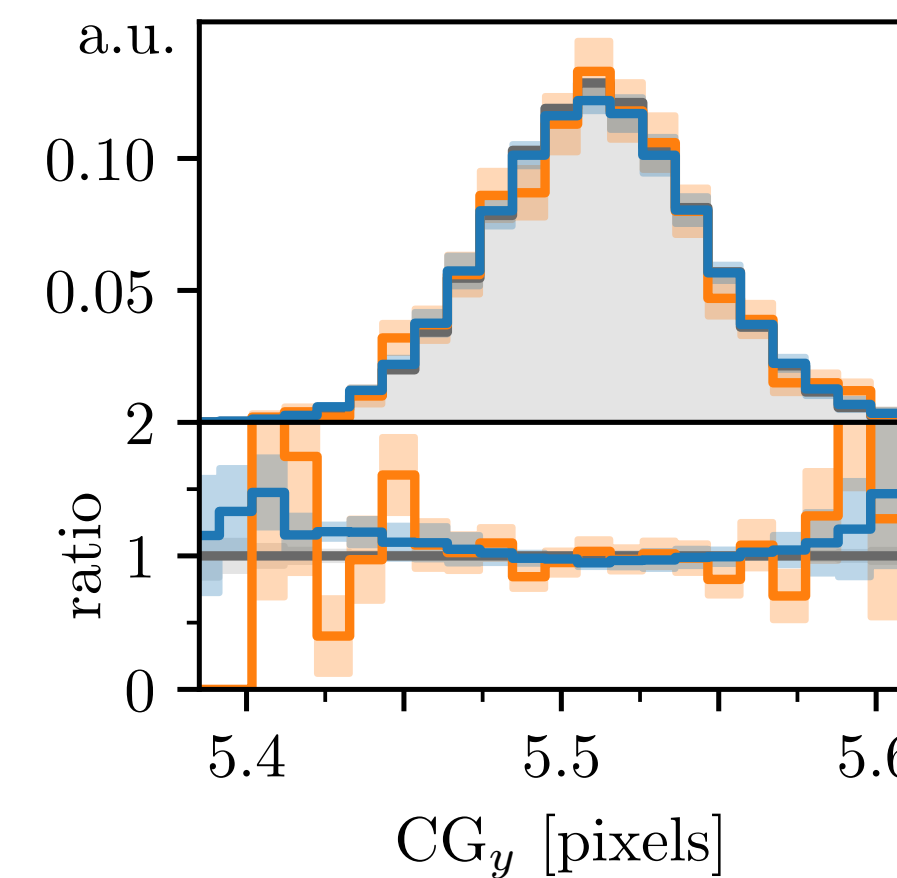
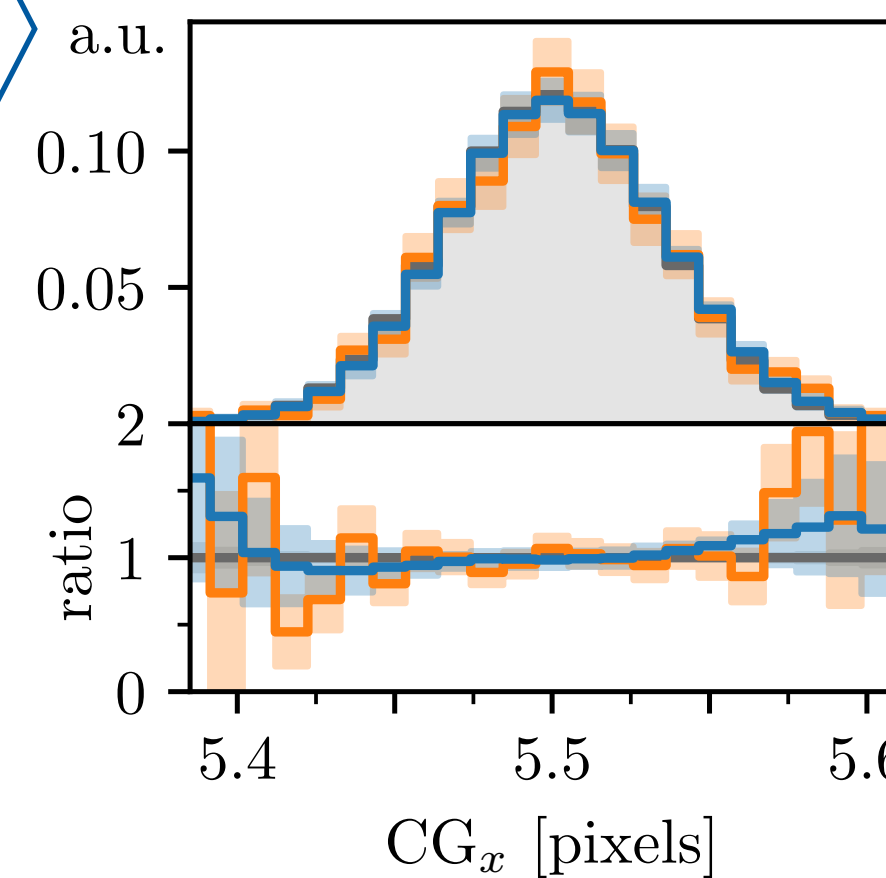
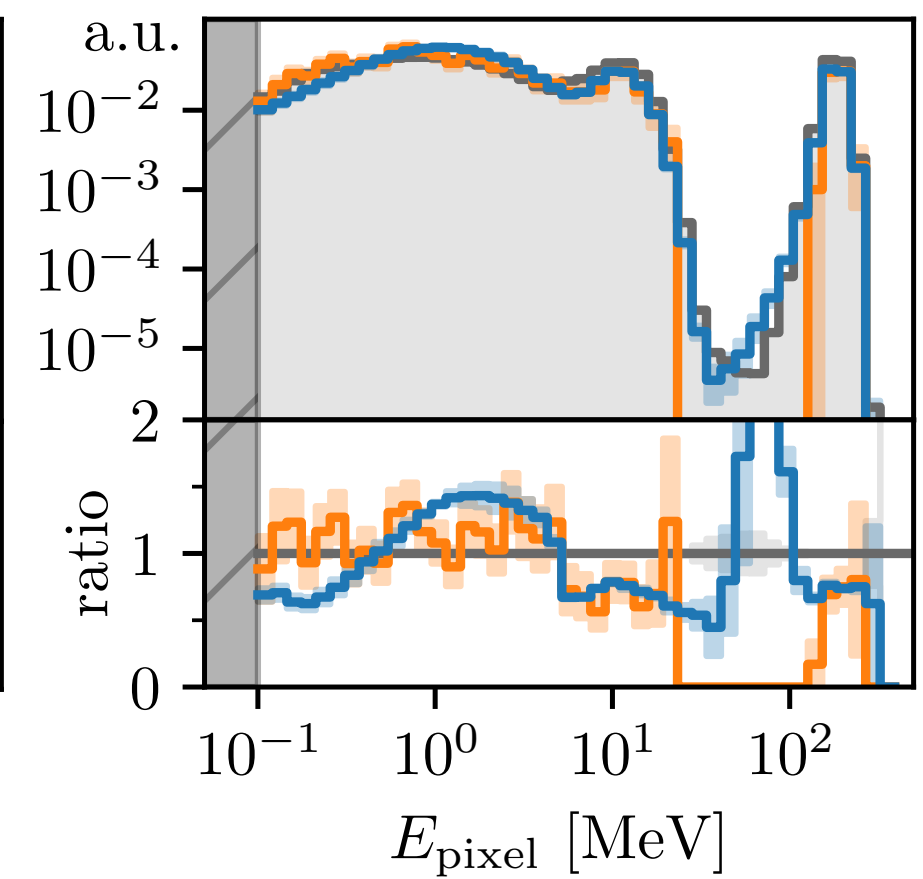
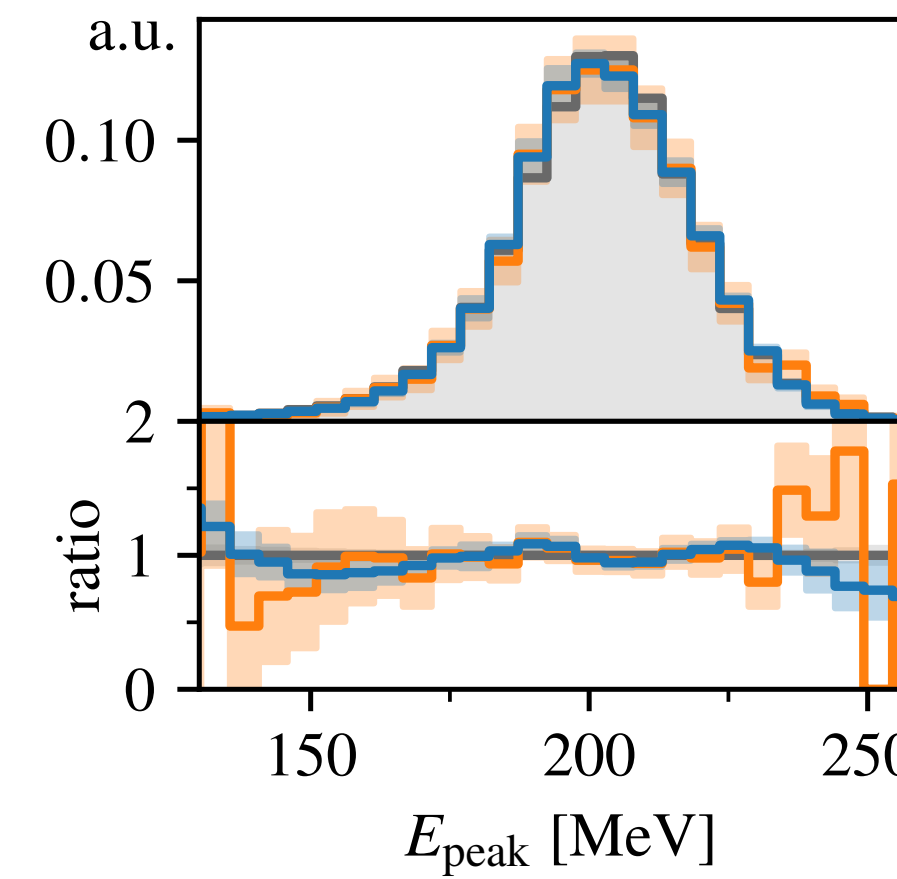
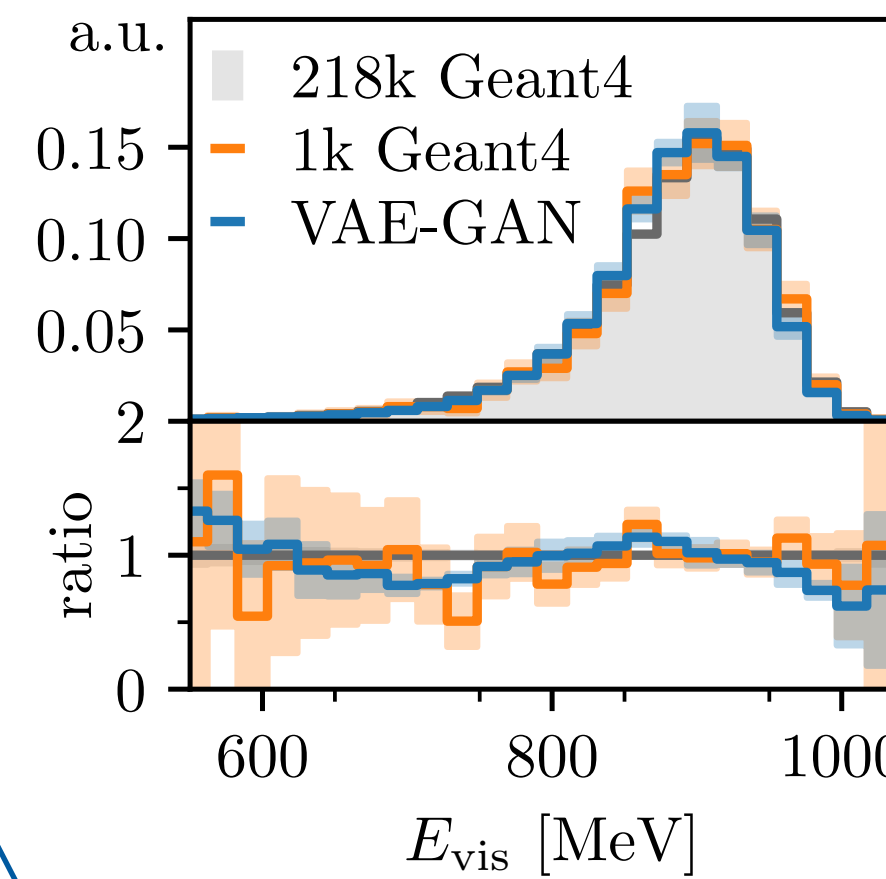


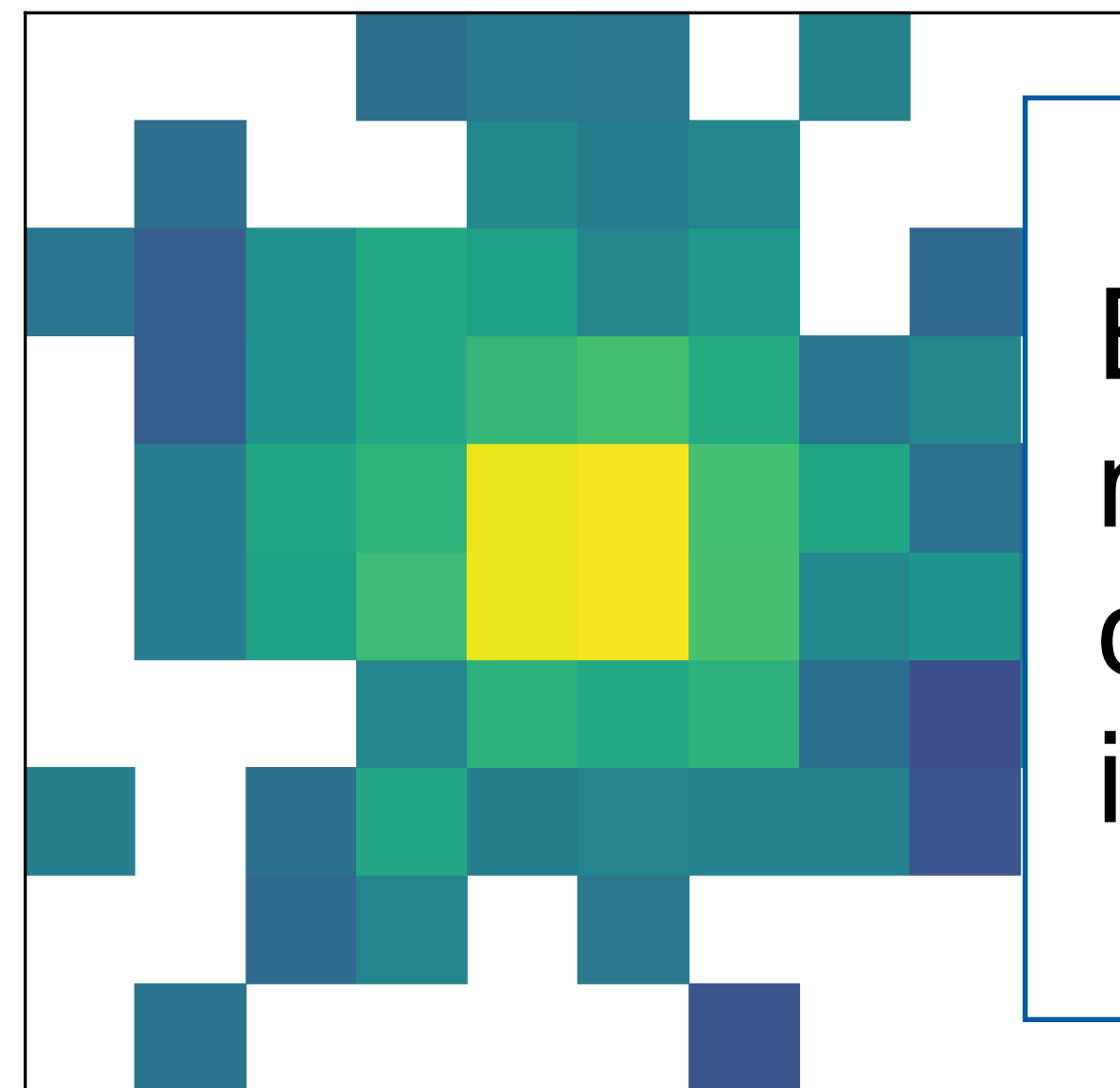
Image-
shaped data

Unknown true
distribution, limited data

Harder learning task → training on
multiple training set sizes unfeasible

Calorimeter Simulations: Setup

DASHH



Evaluate 1D
metrics,
calculated on
images

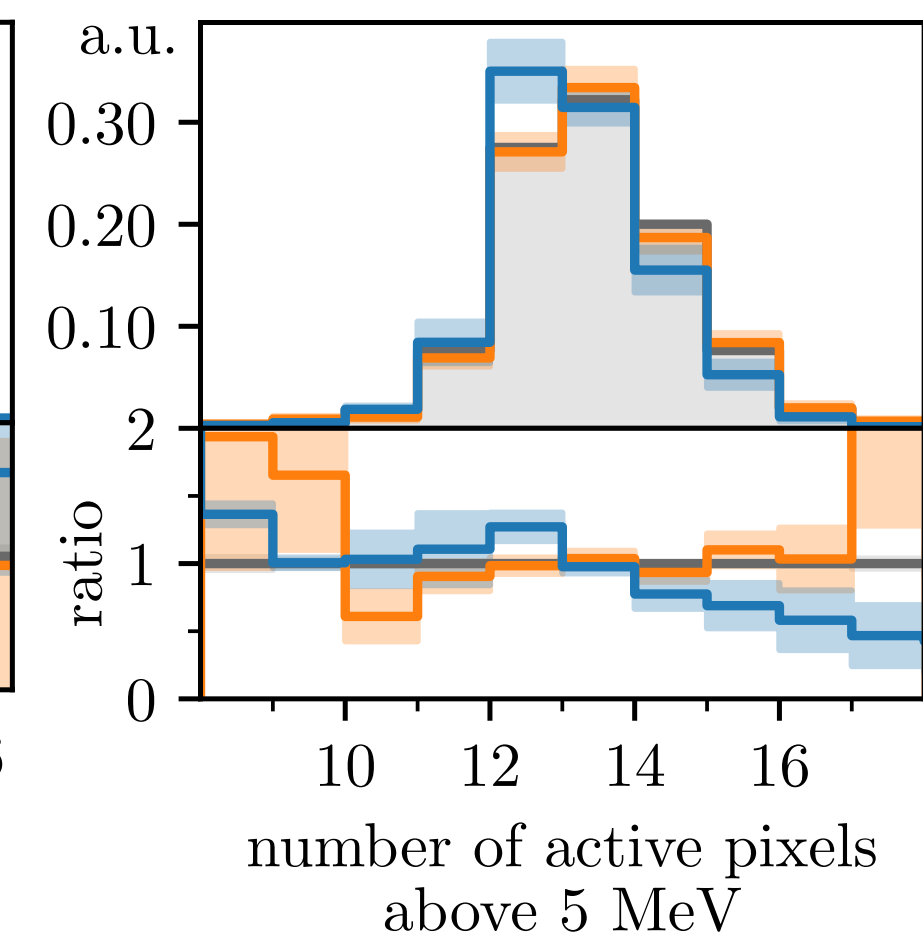
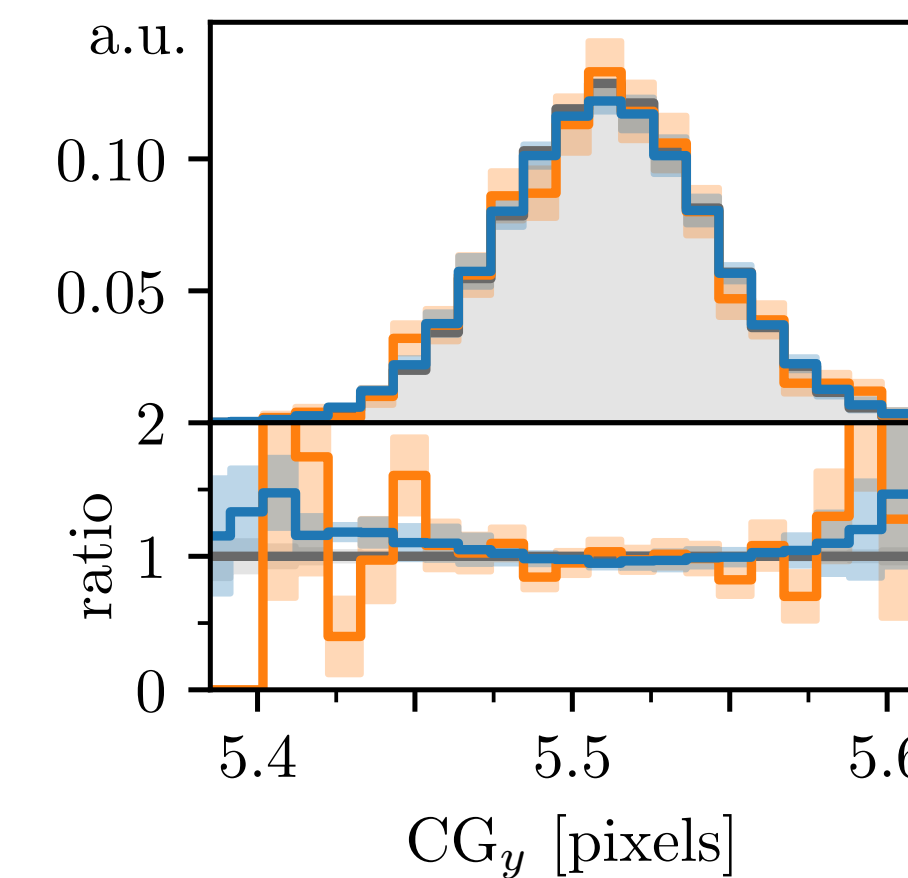
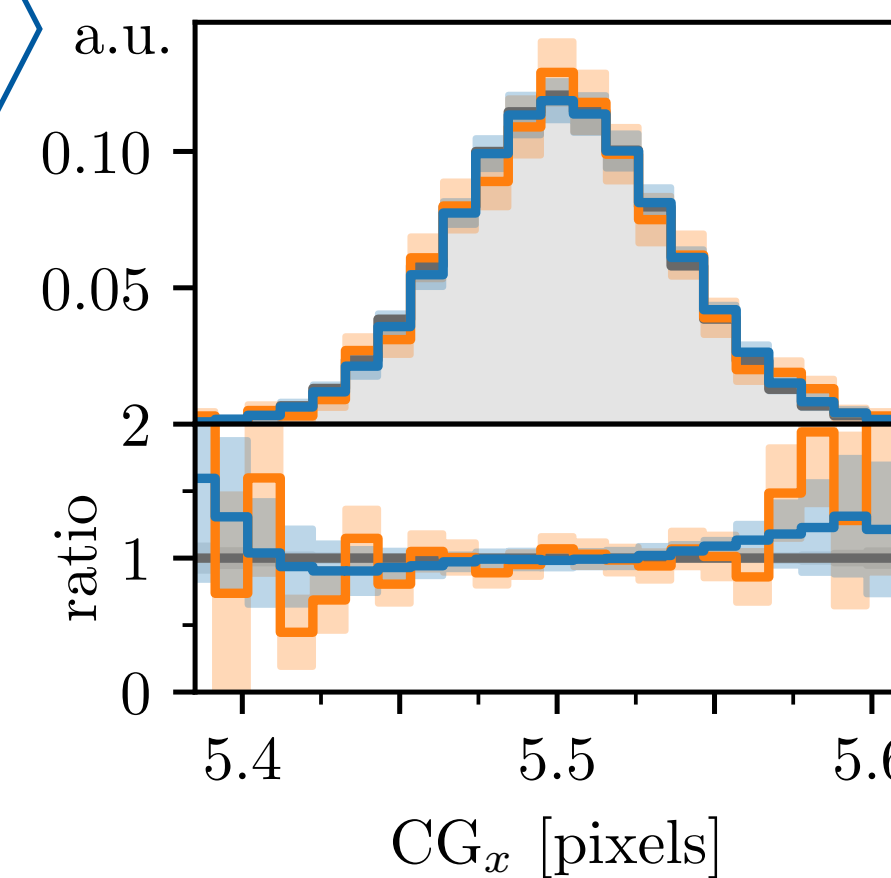
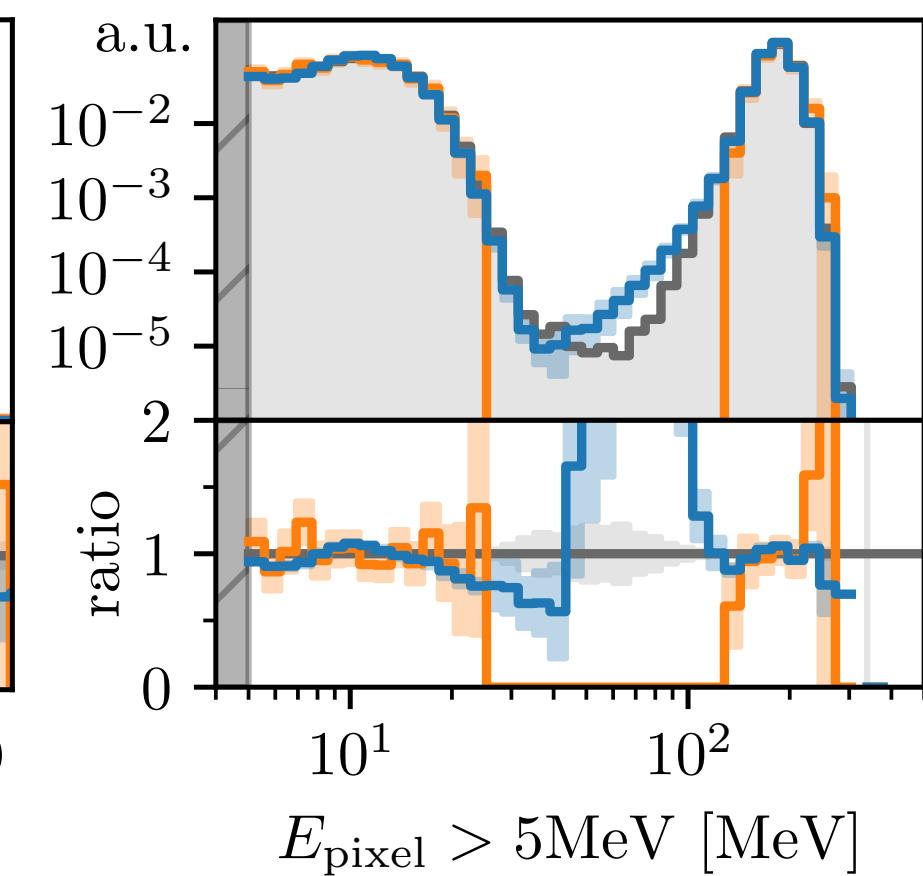
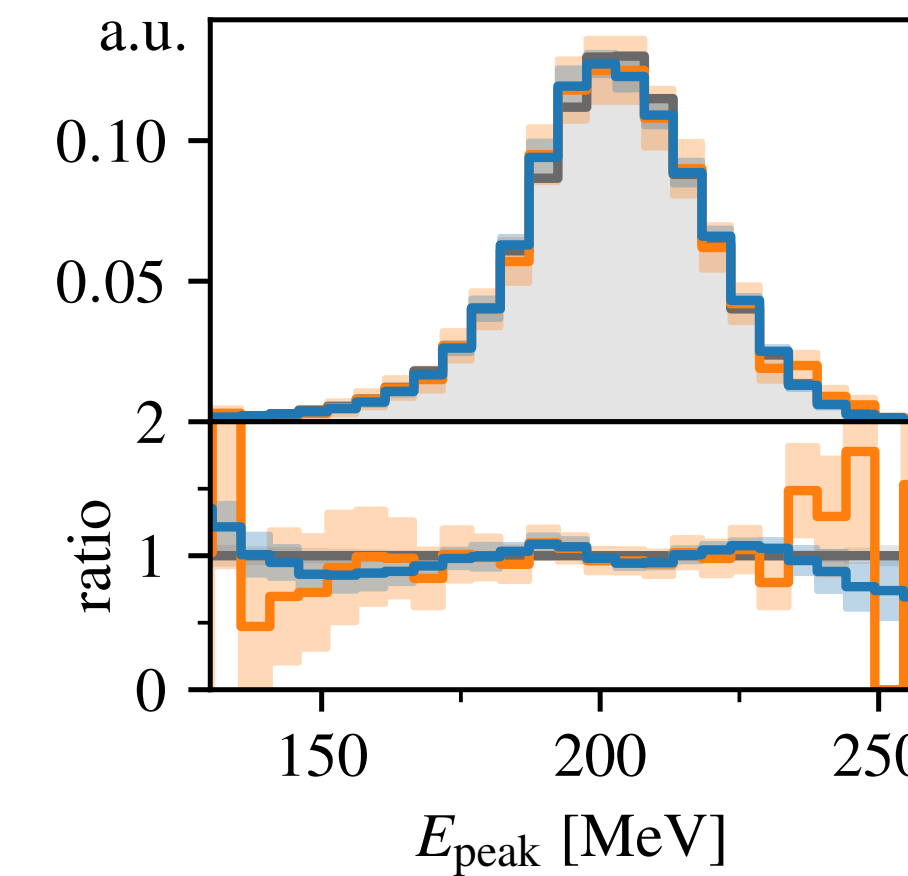
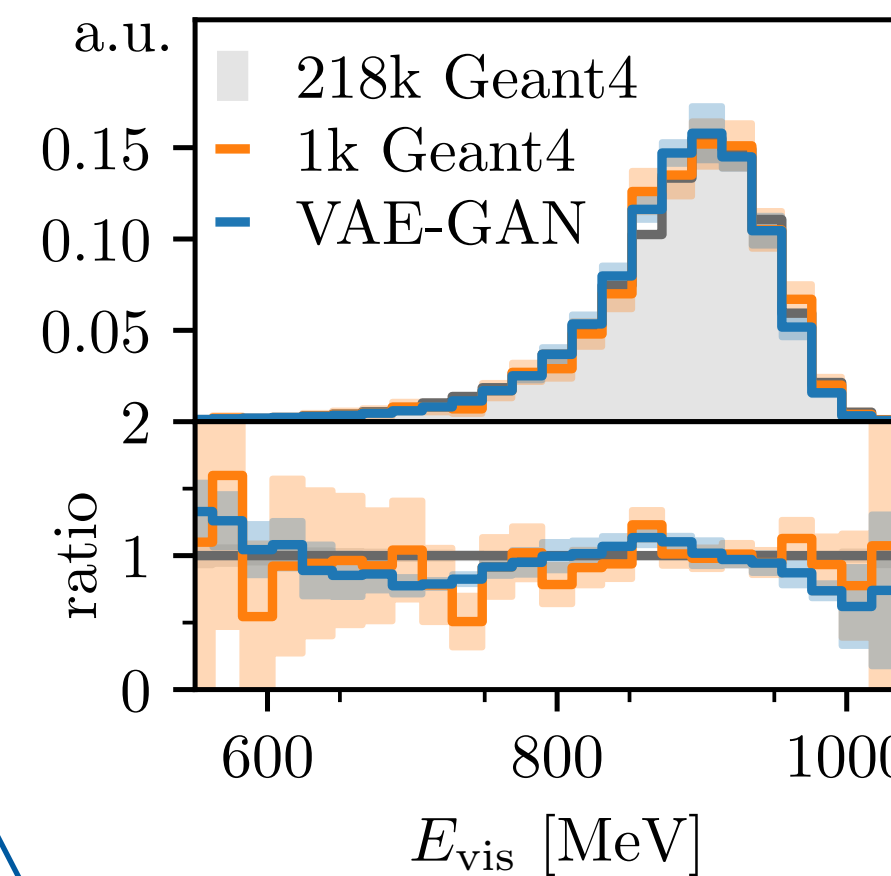


Image-
shaped data

Unknown true
distribution, limited data

Harder learning task $\rightarrow$ training on
multiple training set sizes unfeasible

Calorimeter Simulations: Setup

- Split into **218k validation data points** and **50k evaluation data points**
- Generate quantiles by dividing the validation set into equally populated parts

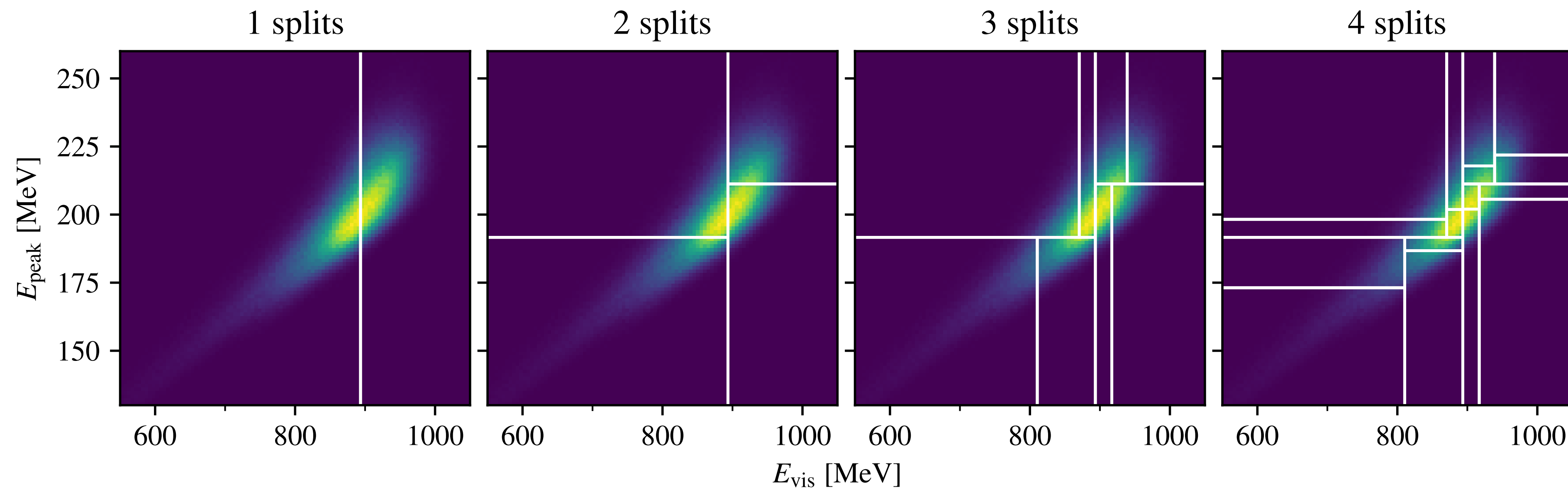


Image-
shaped data

Unknown true
distribution, limited data

Harder learning task → training on
multiple training set sizes unfeasible

Calorimeter Simulations: Setup

Calculate deviation metric $\overline{D}_{\text{JS}}(g || p) = \frac{1}{2} \sum_{Q_i \in \mathbf{Q}} \left(g_i \log \frac{g_i}{\frac{1}{2}(g_i + p_i)} + p_i \log \frac{p_i}{\frac{1}{2}(g_i + p_i)} \right).$

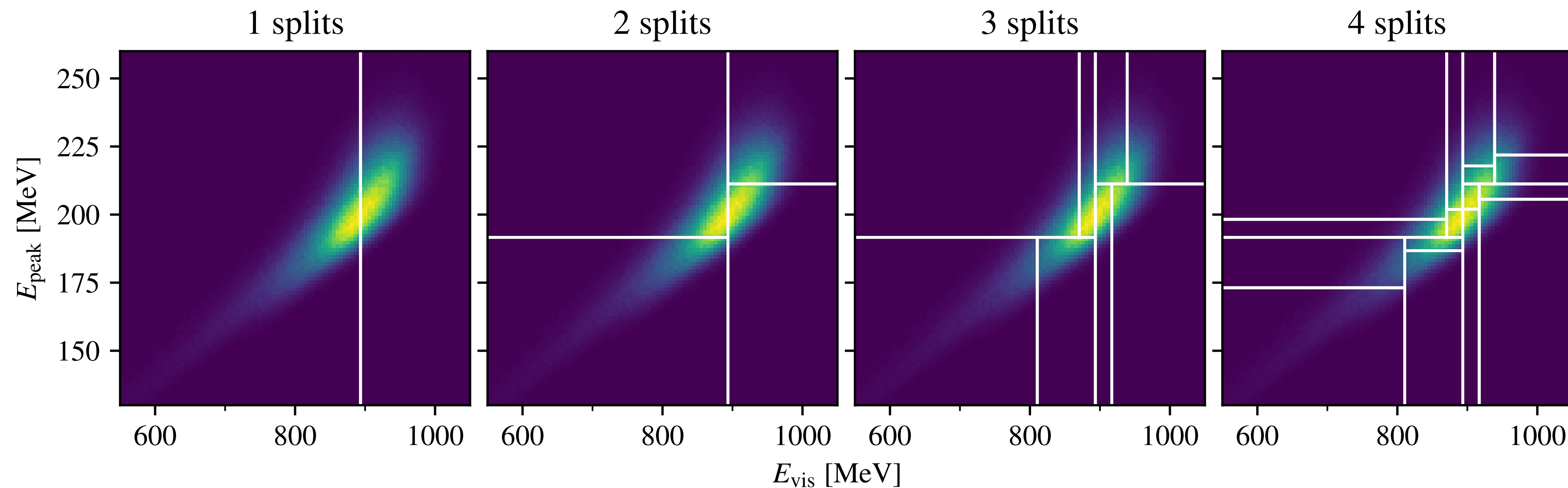


Image-
shaped data

Unknown true
distribution, limited data

Harder learning task → training on
multiple training set sizes unfeasible

Calorimeter Simulations: Results

- Evaluate for fixed training (1k) and evaluation set sizes (5k, 10k, 50k)
- Use less than $n_{\text{data}}/10$ bins

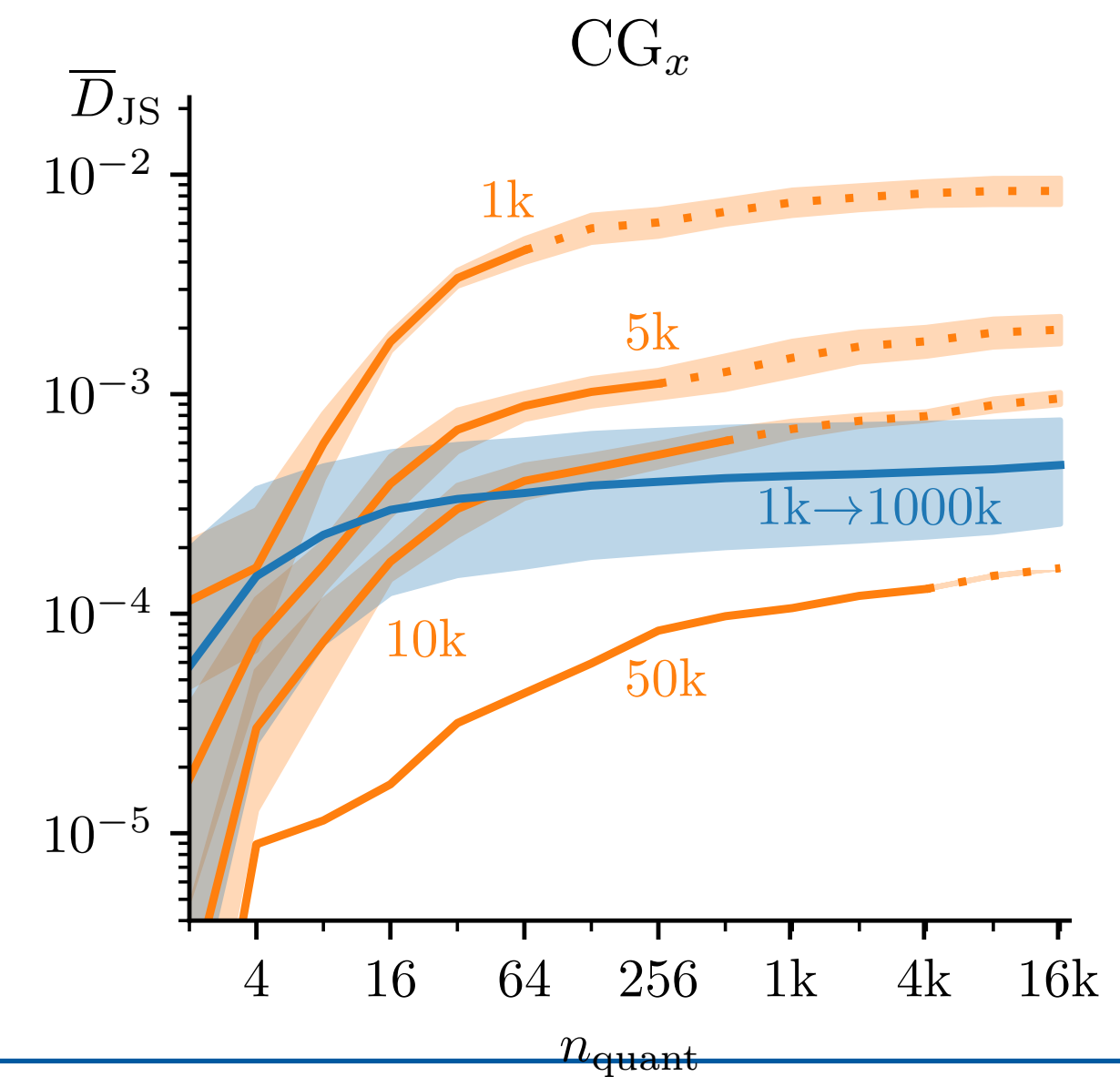
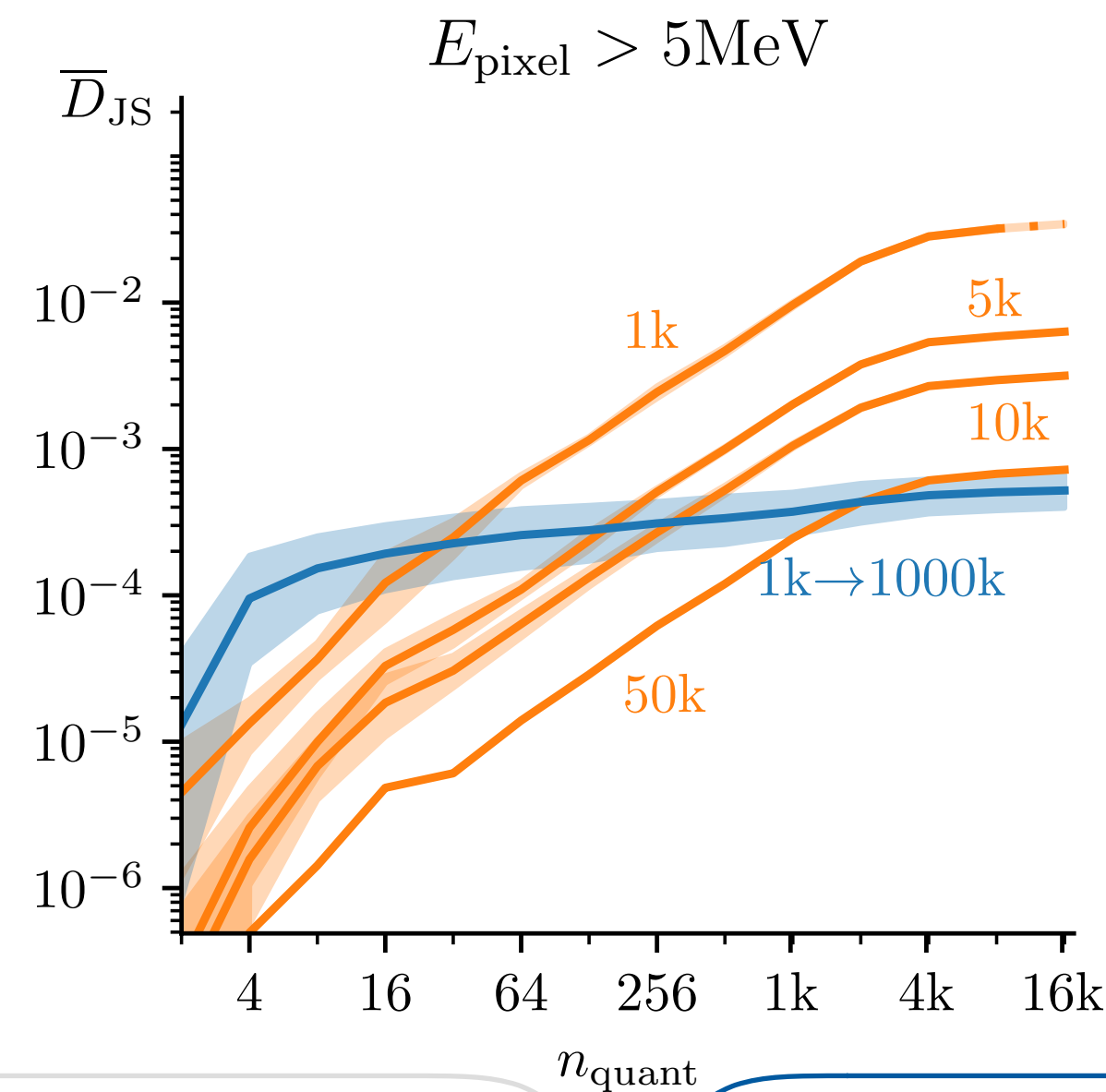
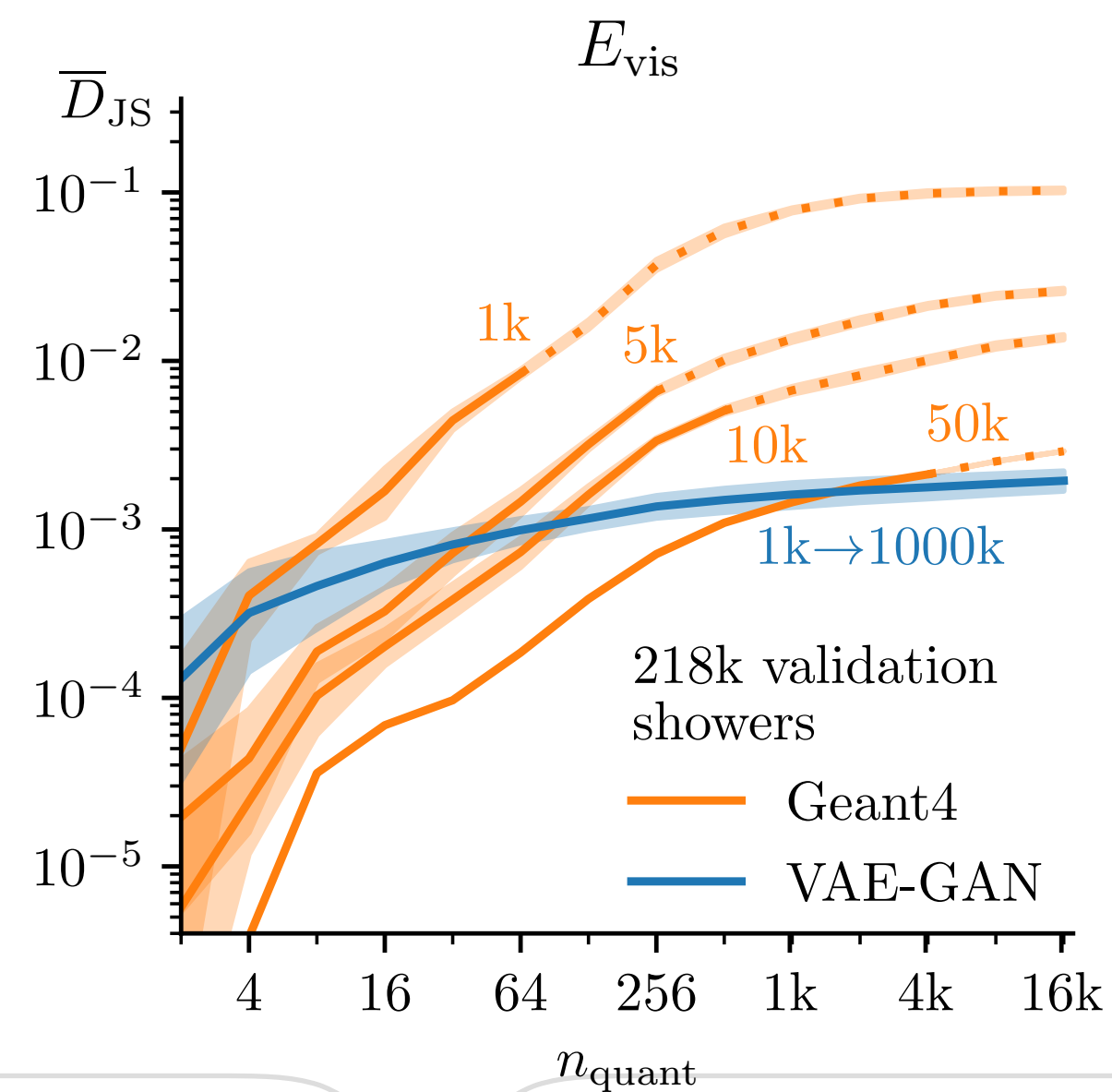


Image-shaped data

Unknown true distribution, limited data

Harder learning task → training on multiple training set sizes unfeasible

Calorimeter Simulations: Results



- Evaluate for fixed training set size
- Add multiples of the training set size
- Use less than $n_{\text{data}}/10$ bins

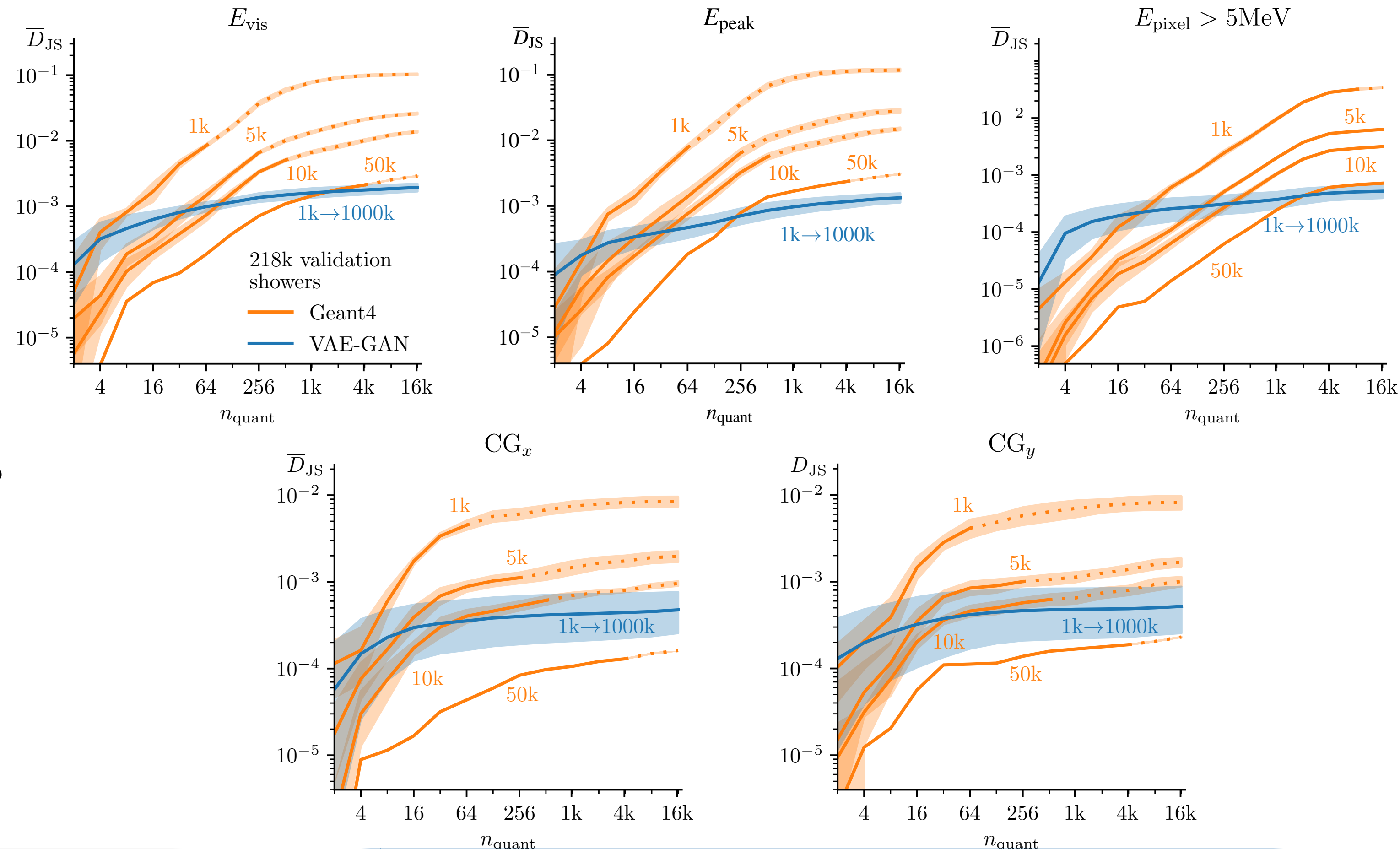


Image-
shaped data

Unknown true
distribution, limited data

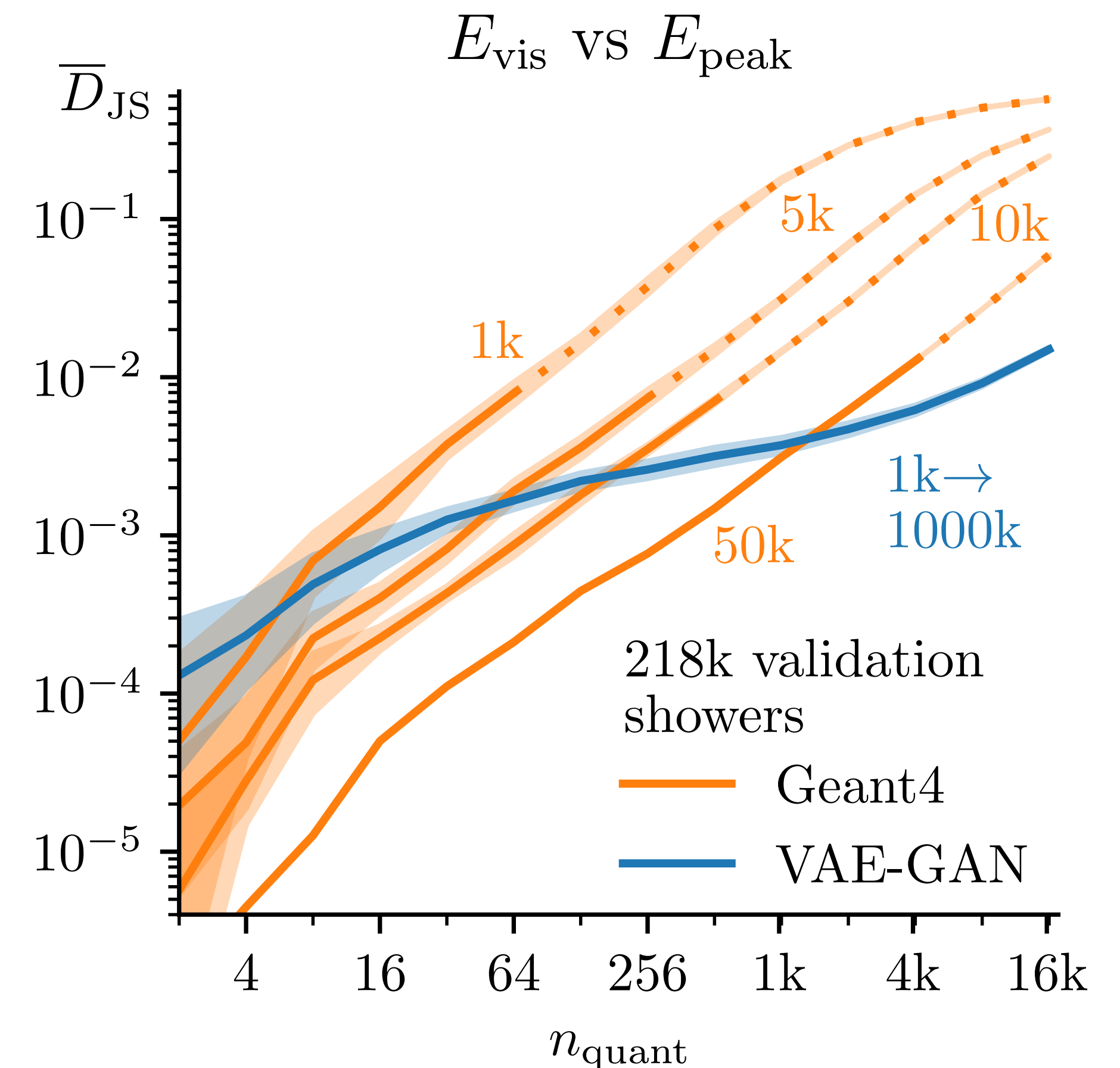
Harder learning task → training on
multiple training set sizes unfeasible

Calorimeter Simulations: Results



- Evaluate for fixed training (1k) and evaluation set sizes (5k, 10k, 50k)
- Use less than $n_{\text{data}}/10$ bins

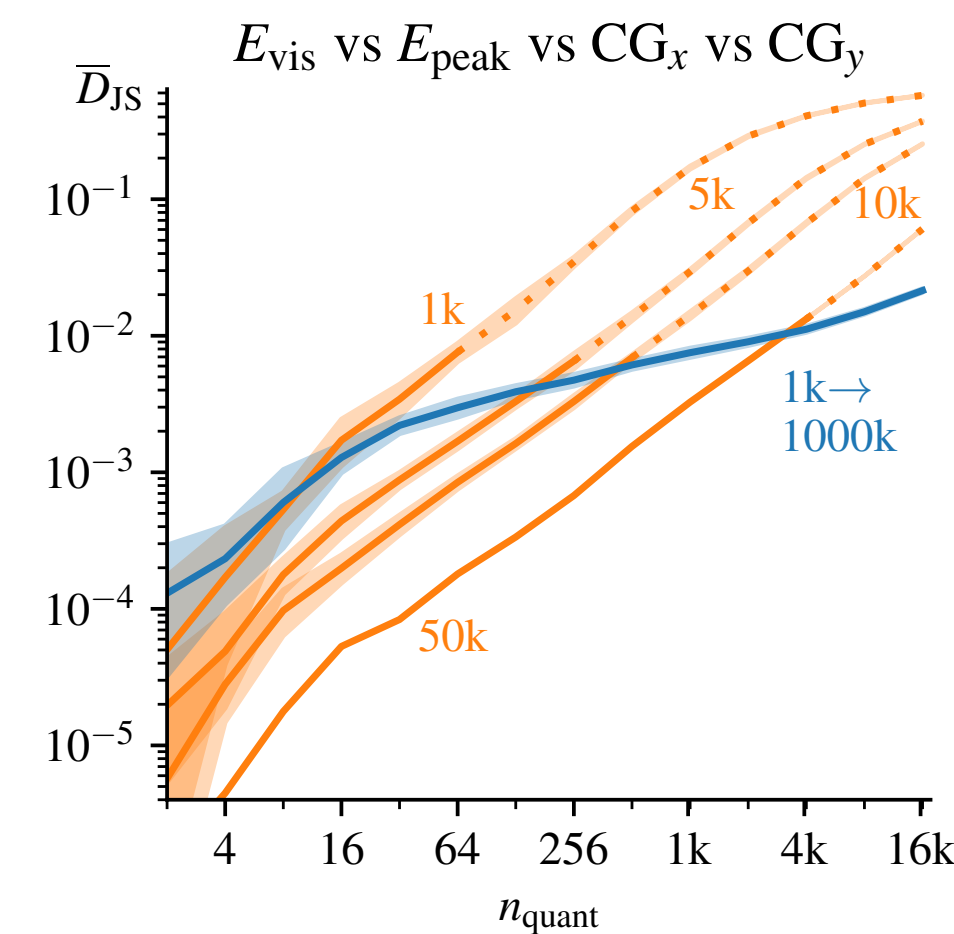
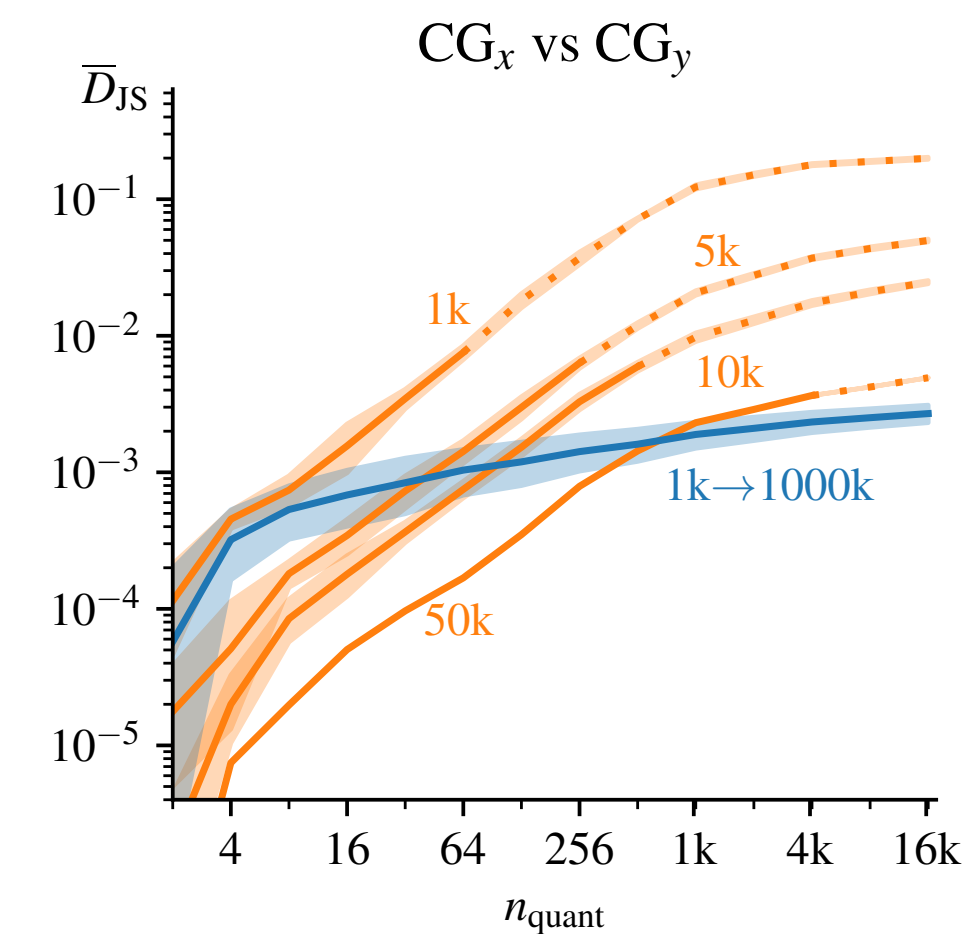
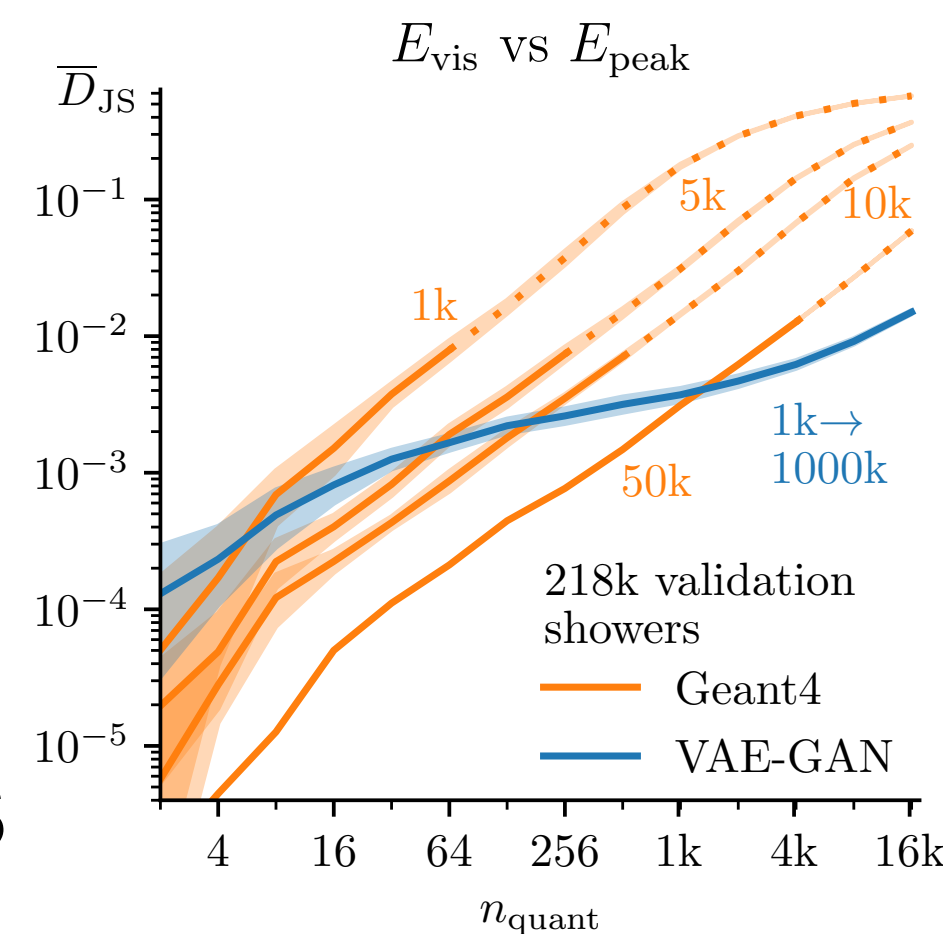
- High-scale features: limited by amount of training data
- Low-scale features: GAN estimation can not be matched by adding more data



Calorimeter Simulations: Results



- Evaluate for fixed training set size
- Add multiples of the training set size
- Use less than $n_{\text{data}}/10$ bins



- High-scale features: limited by amount of training data
- Low-scale features: GAN estimation can not be matched by adding more data

Image-shaped data

Unknown true distribution, limited data

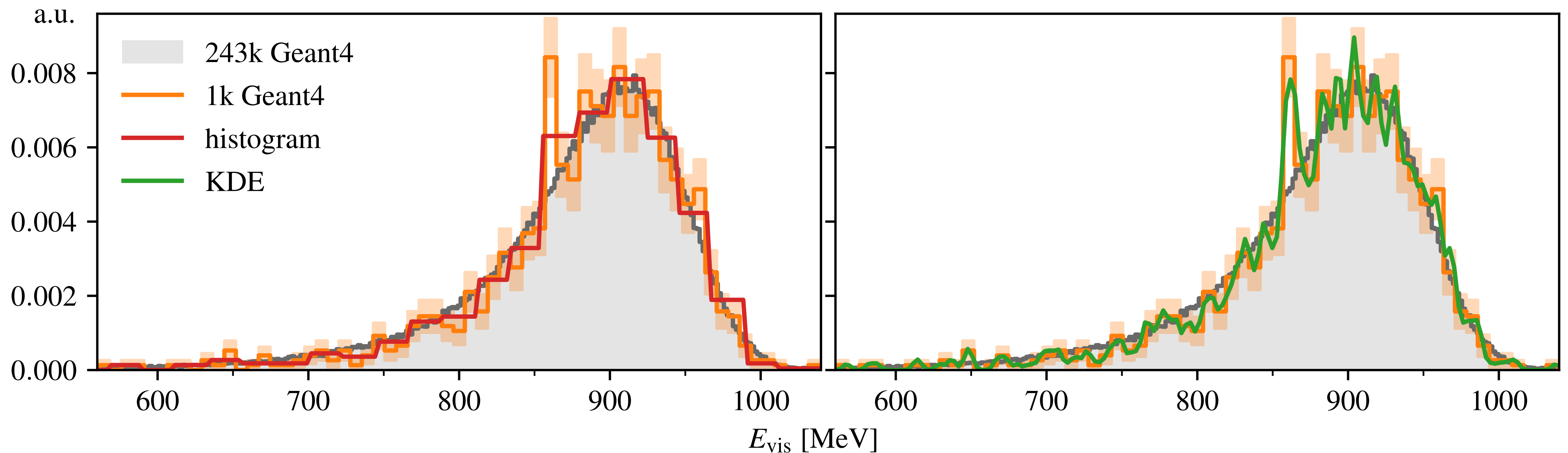
Harder learning task → training on multiple training set sizes unfeasible

Calorimeter Simulations: Results



How good is the density estimation actually?

- Compare to KDE and histogram estimators (maximizing log-likelihood of cross-validation sets)

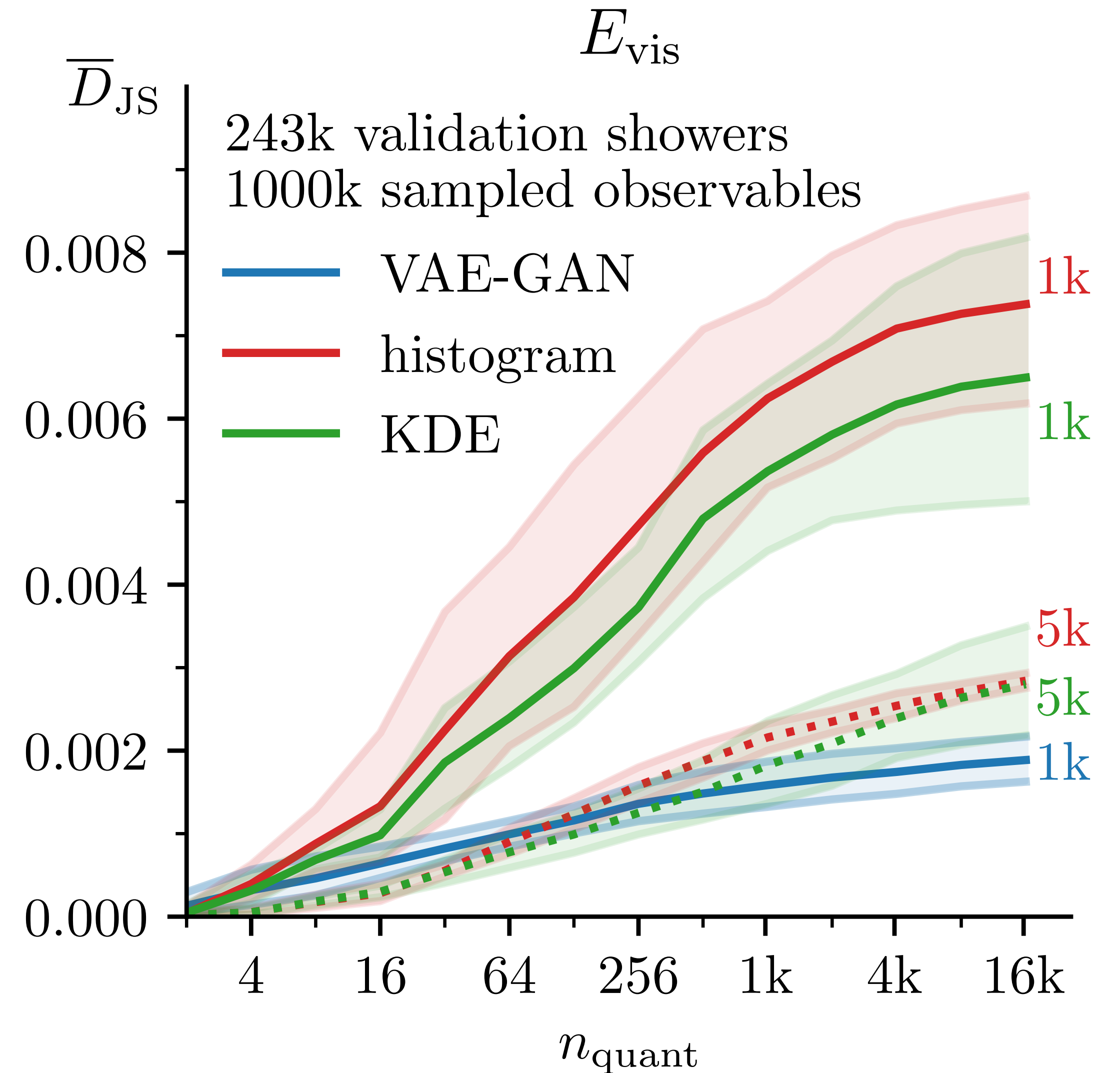


Calorimeter Simulations: Results



- Generate 10^6 samples from every density estimator

- GAN outperforms standard density estimators



Calorimeter Simulations: Results

DASHH

- Generate 10^6 samples from every density estimator

- GAN outperforms standard density estimators

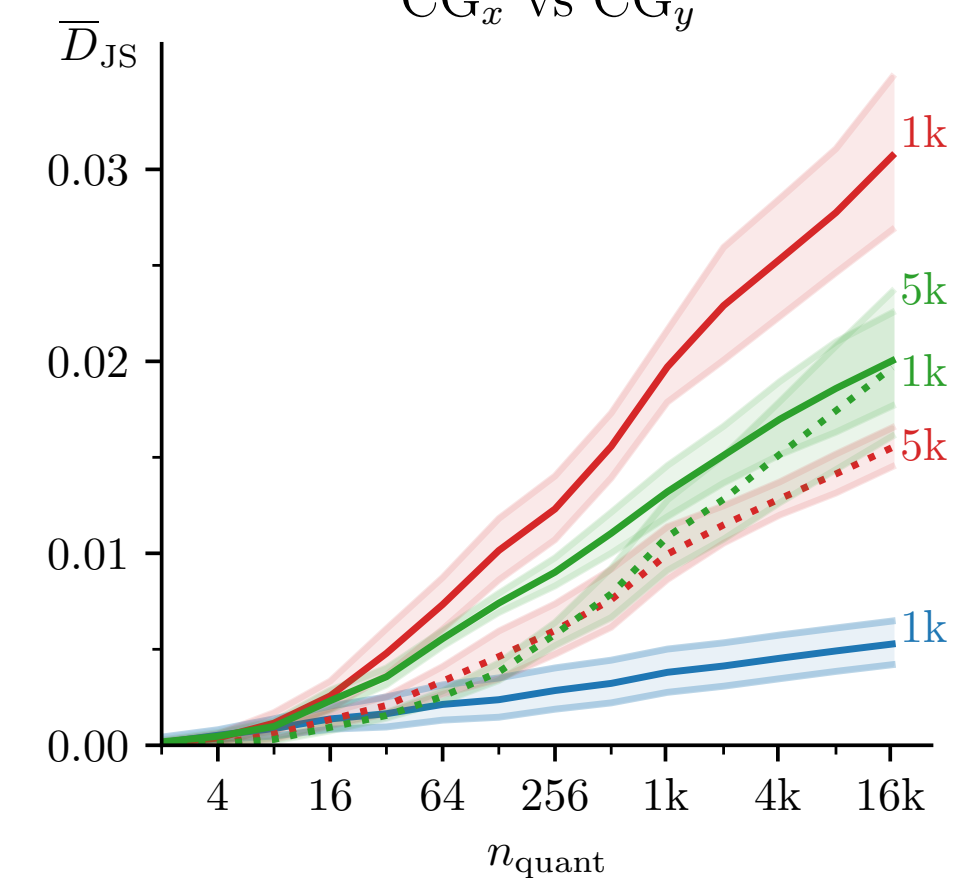
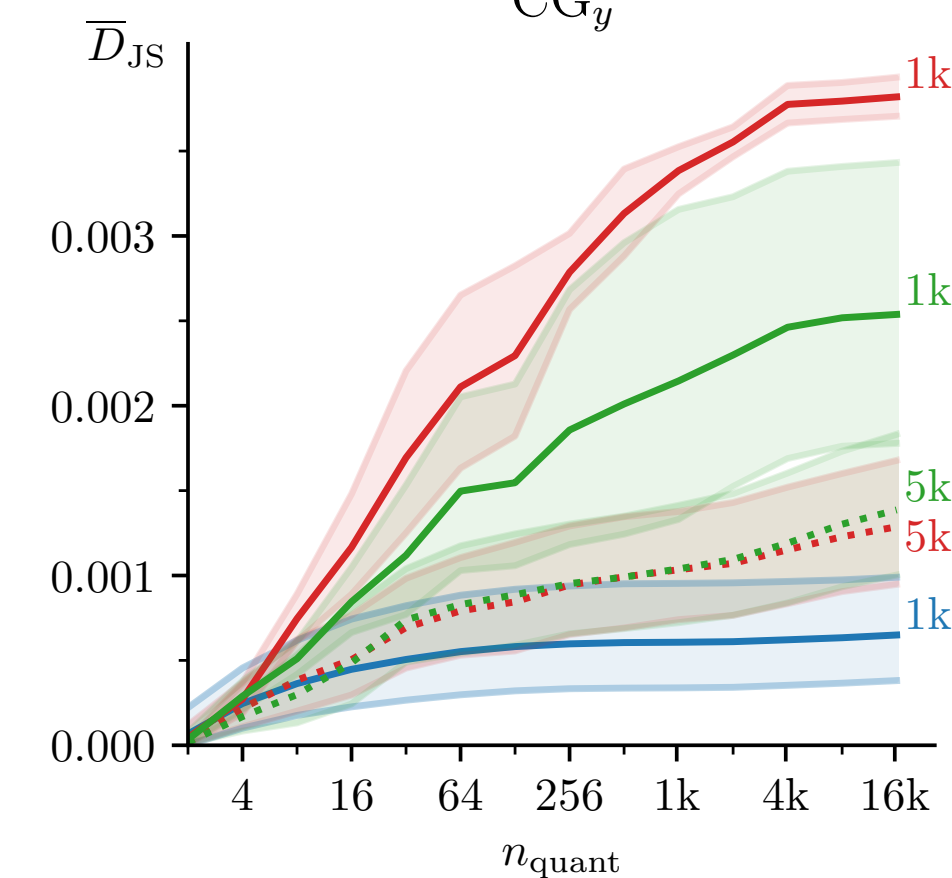
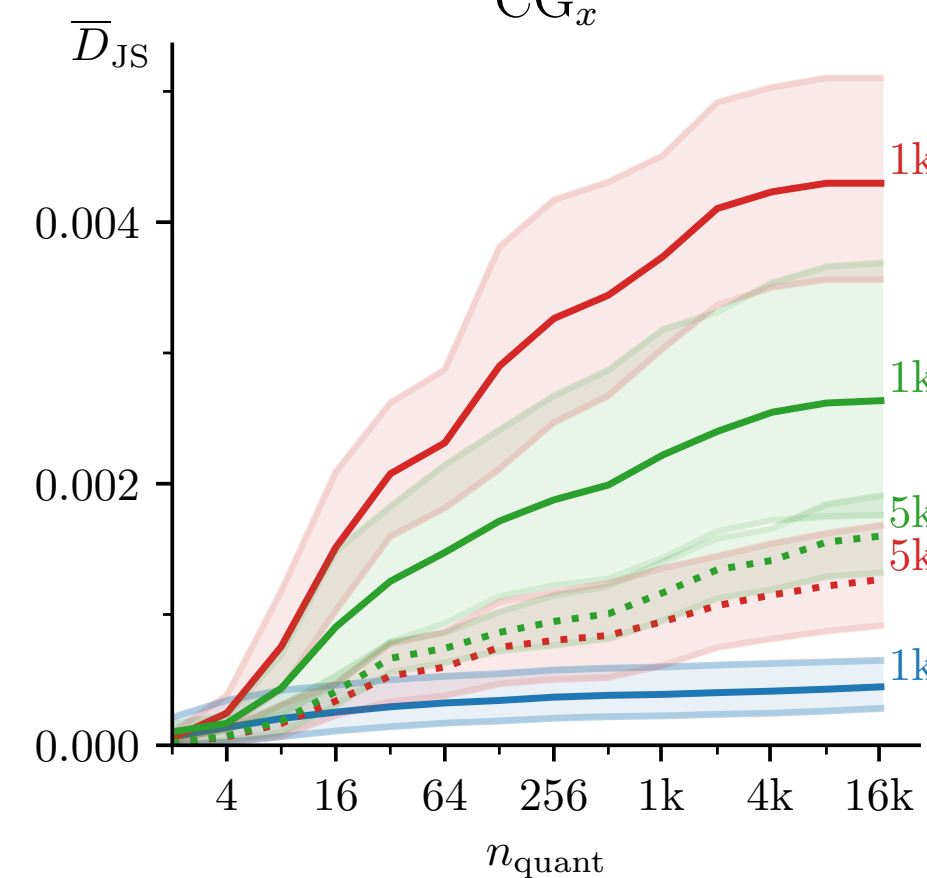
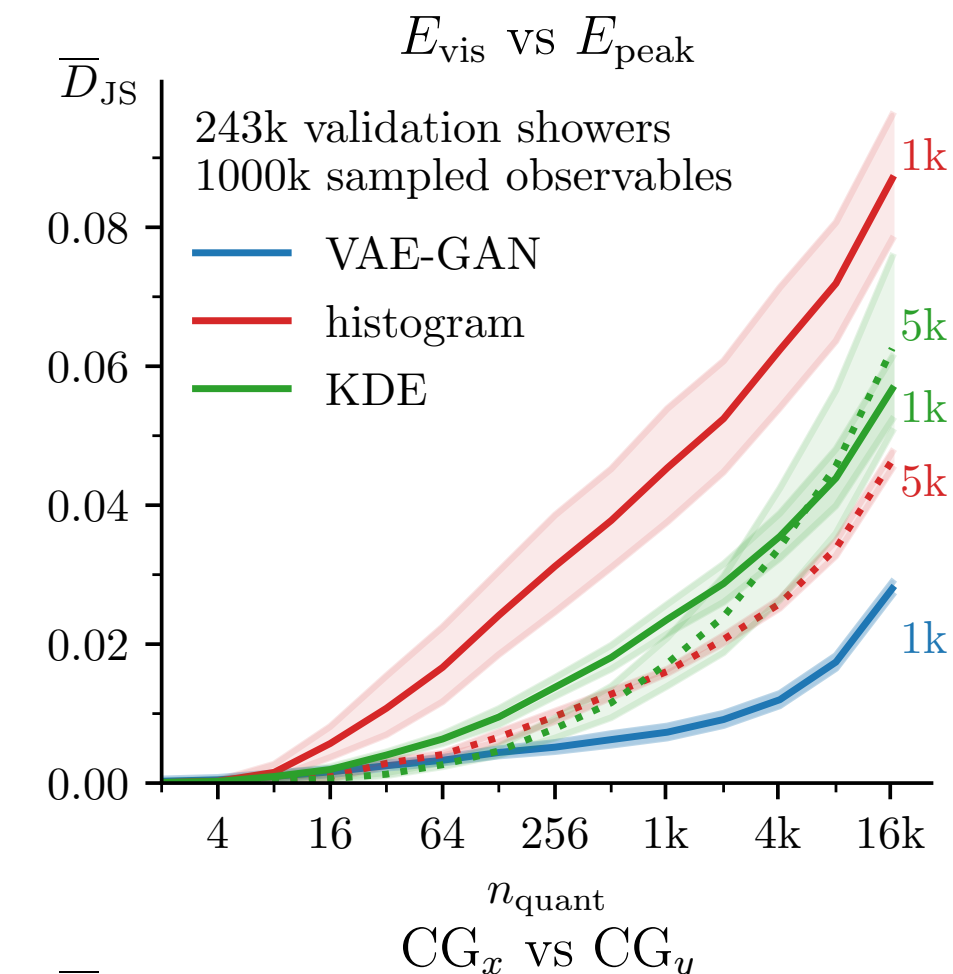
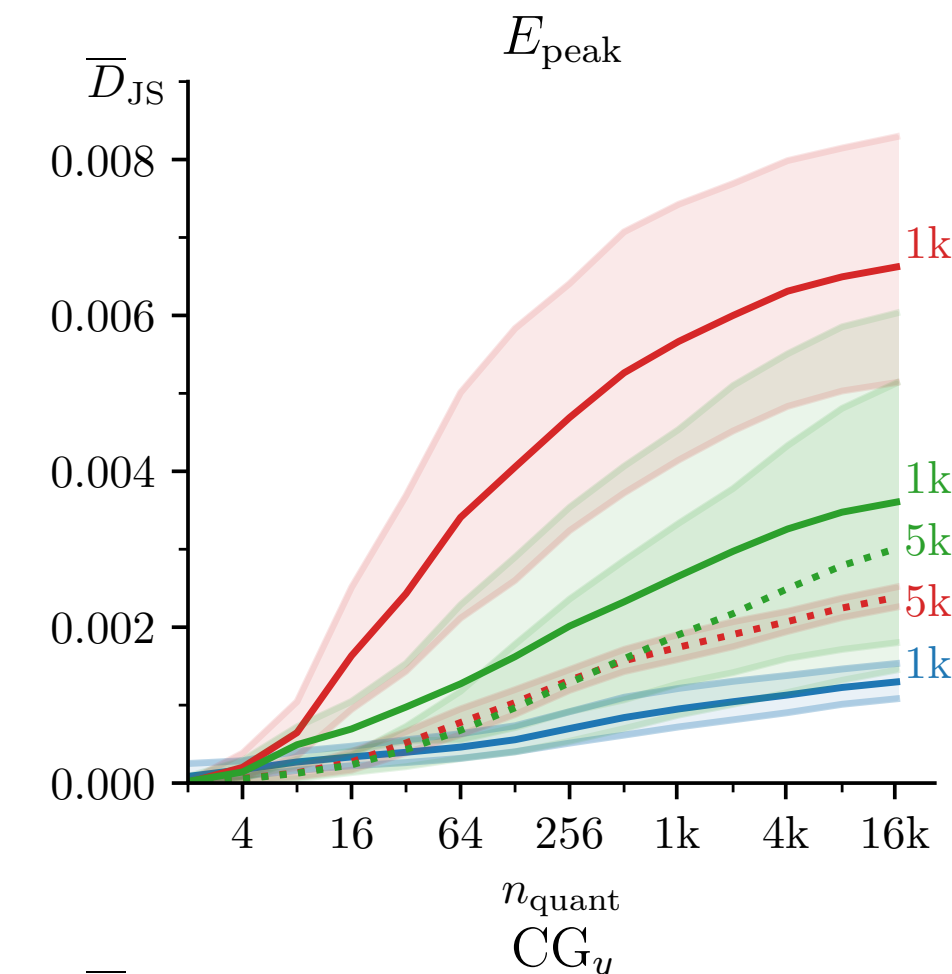
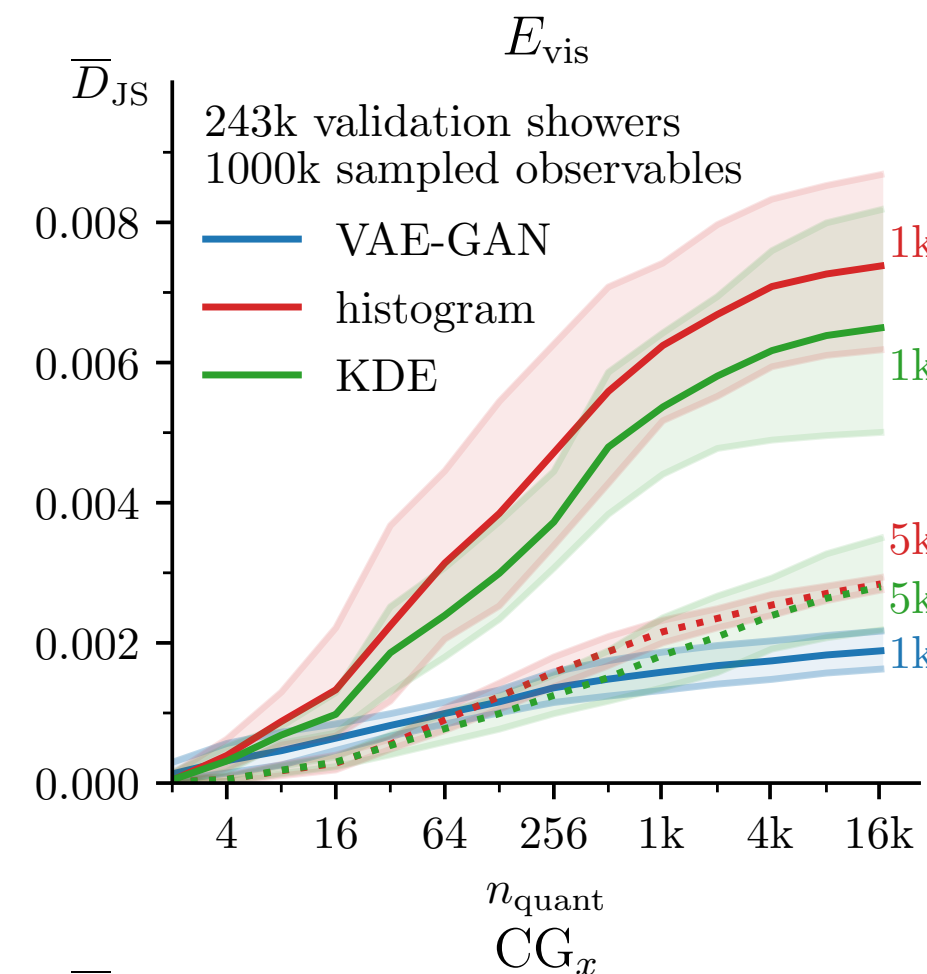


Image-shaped data

Unknown true distribution, limited data

Harder learning task → training on multiple training set sizes unfeasible

Conclusion

- What about # samples? How many new points should we generate from a generative model?
 - Depends on GAN setup and problem

- For high-scale observables (e.g. *mean, standard deviation, low moments*) generative network limited to the amount of training data
- For a smooth interpolation (e.g. *segments of the distribution, integrated quantities*) a generative networks outperform even higher numbers of data

References



- [0]: P. Calafiura, J. Catmore, D. Costanzo, and A. Di Girolamo, “ATLAS HLLHC Computing Conceptual Design Report,” CERN, Geneva, Tech. Rep., Sep 2020. [Online]. Available: <https://cds.cern.ch/record/2729668>
- [1]: ILD Concept Group, H. Abramowicz et al., *International Large Detector: Interim Design Report*, 3, 2020.
- [2]: L. de Oliveira, M. Paganini, and B. Nachman, “Learning particle physics by example: Location-aware generative adversarial networks for physics synthesis,” *Computing and Software for Big Science*, vol. 1, no. 1, Sep 2017. [Online]. Available: <http://dx.doi.org/10.1007/s4178101700046>
- [3]: M. Paganini, L. de Oliveira, and B. Nachman, “Calogan: Simulating 3d high energy particle showers in multilayer electromagnetic calorimeters with generative adversarial networks,” *Physical Review D*, vol. 97, no. 1, Jan 2018. [Online]. Available: <http://dx.doi.org/10.1103/PhysRevD.97.014021>
- [4]: A. B. L. Larsen, S. K. Sønderby, H. Larochelle, and O. Winther, “Autoencoding beyond pixels using a learned similarity metric,” in *Proceedings of the 33rd International Conference on International Conference on Machine Learning Volume 48*. JMLR.org, 2016, p. 1558–1566.

Why do we need N^2_{quant} datapoints?

$$\hat{\chi}_N^2 = N \sum_{j=0}^N \left(x_j - \frac{1}{N} \right)^2 = N \sum_{j=1}^N \left(\hat{F}\left(F^{-1}\left(\frac{j}{N}\right)\right) - \hat{F}\left(F^{-1}\left(\frac{j-1}{N}\right)\right) - \frac{1}{N} \right)^2,$$

where we can use the CDF $\hat{F}$ of the underlying histogram

$$\hat{f}_n(x) = \frac{1}{hn} \sum_{i=0}^n 1_{(x_-, x_+]}(x_i) \text{ with } h = F^{-1}\left(\frac{j}{N}\right) - F^{-1}\left(\frac{j-1}{N}\right) \text{ and } x_- = F^{-1}\left(\frac{j-1}{N}\right), x_+ = F^{-1}\left(\frac{j}{N}\right) \text{ the edges of the bin } x \text{ is in. For}$$

$h \rightarrow 0$ and $hn, n \rightarrow \infty$ the integrated square error between the true underlying function and the histogram converges to 0. In the same limit the measure fulfills

$$\begin{aligned} \hat{\chi}_N^2 &= \sum_{j=1}^N \left(\int_{F^{-1}((j-1)/N)}^{F^{-1}(j/N)} (\hat{f}(x) - f(x)) dx \right)^2 = \sum_{j=1}^N \int_{F^{-1}((j-1)/N)}^{F^{-1}(j/N)} ((\hat{f}(x) - f(x)) N \underbrace{\int_{F^{-1}((j-1)/N)}^{F^{-1}(j/N)} (\hat{f}(y) - f(y)) dy}_{\approx (\hat{f}(x) - f(x)) \frac{F^{-1}(j/N) - F^{-1}((j-1)/N)}{1/N}}) dx \\ &\approx \sum_{j=1}^N \int_{F^{-1}((j-1)/N)}^{F^{-1}(j/N)} ((\hat{f}(x) - f(x)) \underbrace{2 \frac{F^{-1}(j/N) - F^{-1}((j-1)/N)}{1/N}}_{\rightarrow (F^{-1})'(j/N) = 1/f(F(j/N))}) dx \approx \sum_{j=1}^N \int_{F^{-1}((j-1)/N)}^{F^{-1}(j/N)} \frac{((\hat{f}(x) - f(x))^2}{f(x)} dx = \int \frac{((\hat{f}(x) - f(x))^2}{f(x)} dx. \end{aligned}$$