



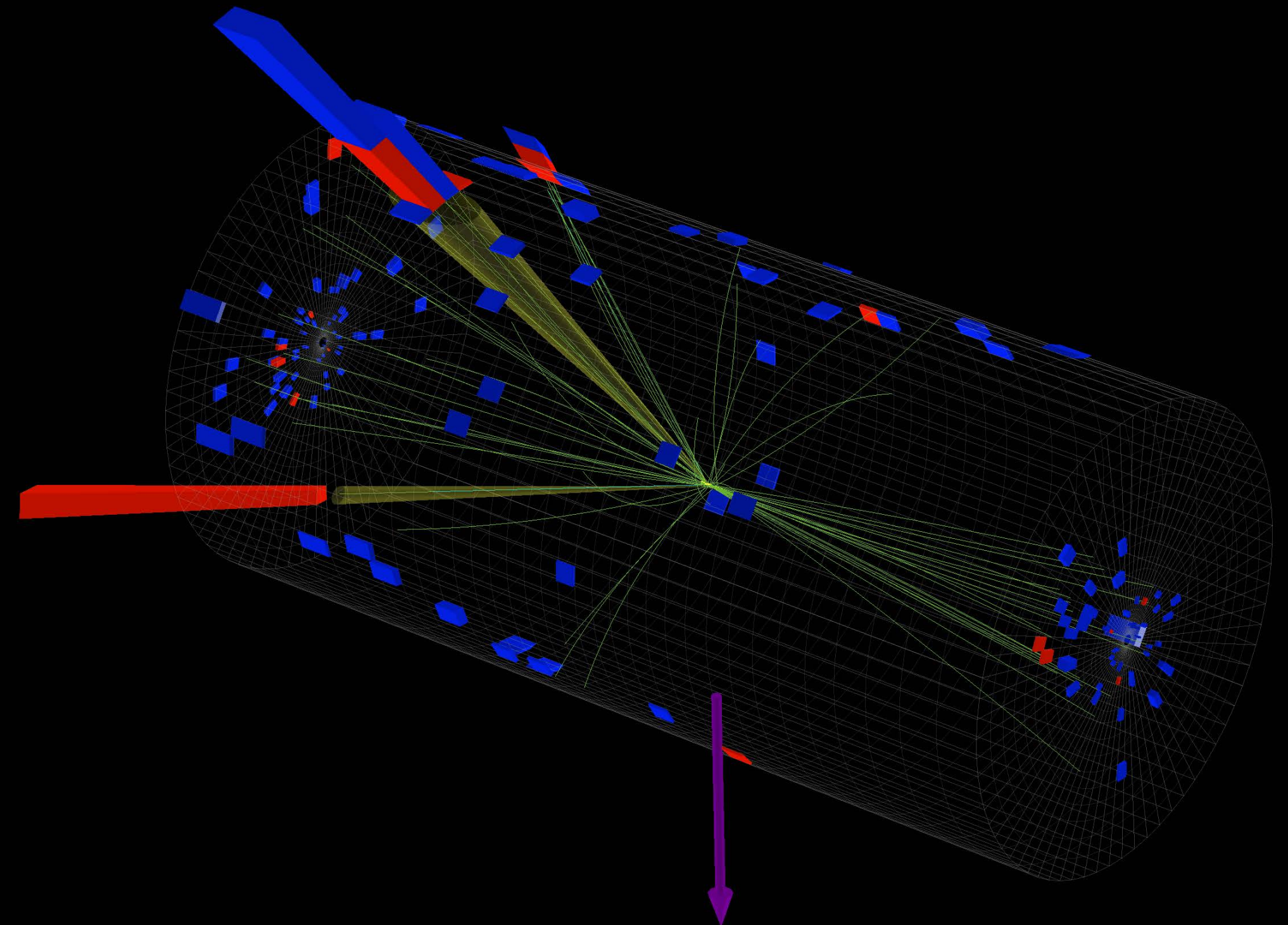
李政道研究所  
TSUNG-DAO LEE INSTITUTE

# *AI For HEP*

*— A Representation Learning Perspective*

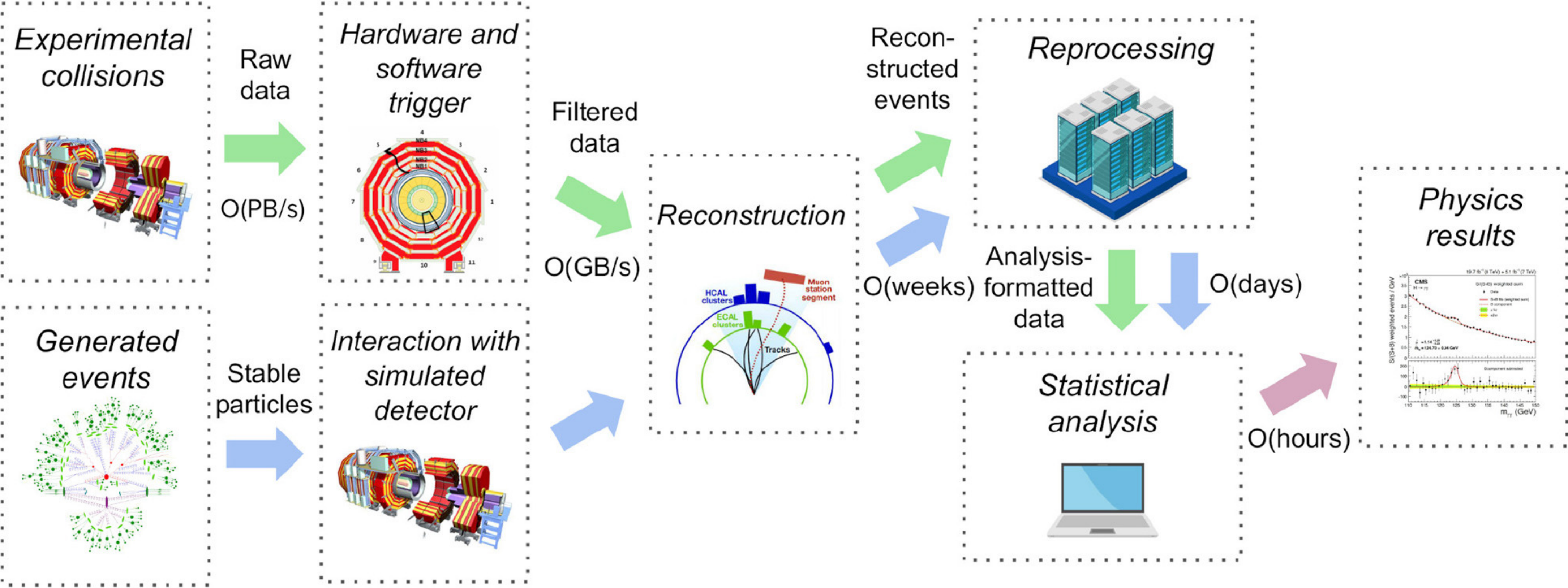
Huilin Qu (曲慧麟)

*New Opportunities for Particle Physics 2026*  
*August 16, 2026*





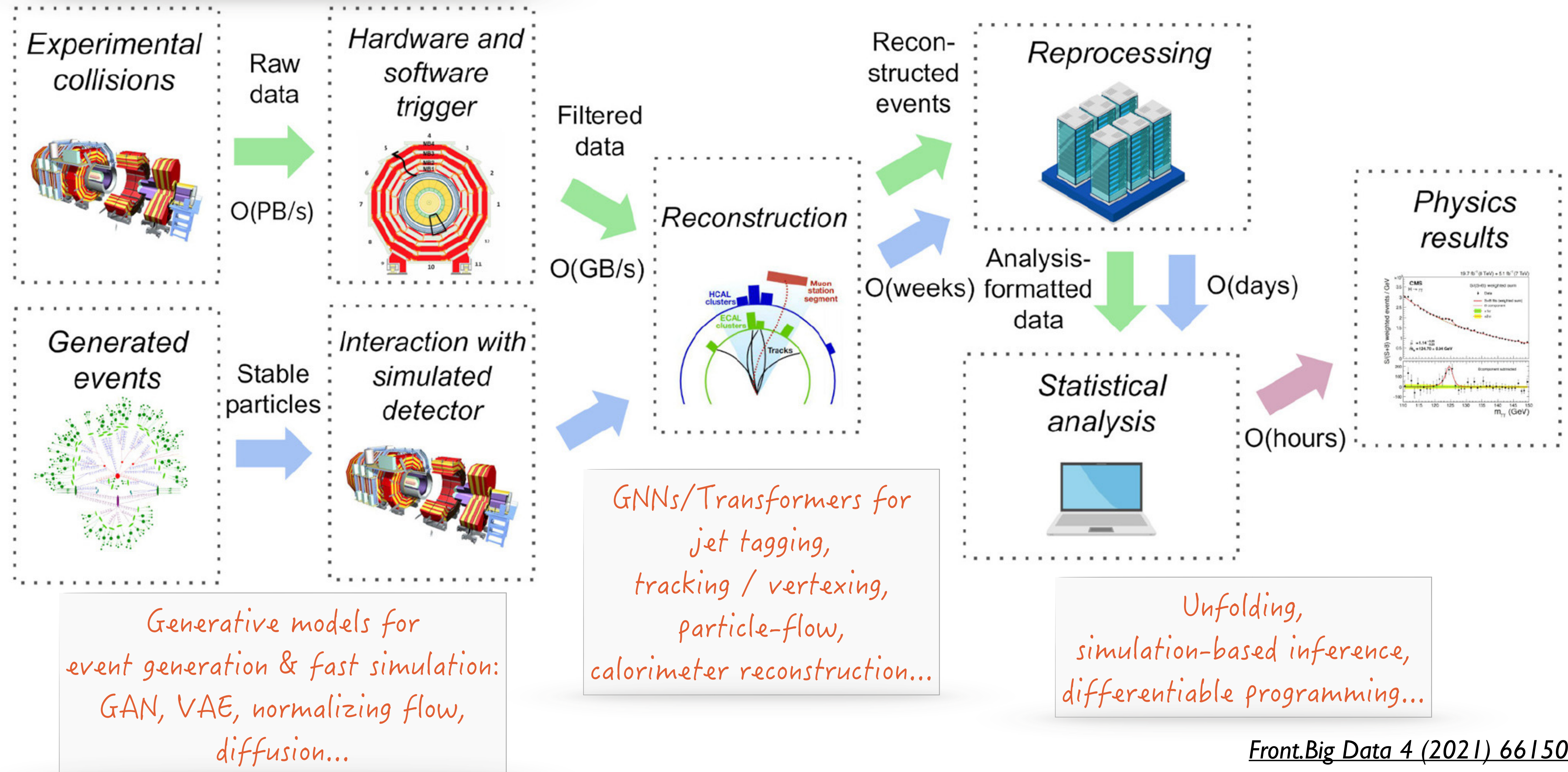
# HEP DATA FLOW...





# HEP DATA FLOW... MATCHED WITH ML!

Ultrafast inference (FPGA/ASIC),  
data compression, anomaly detection...

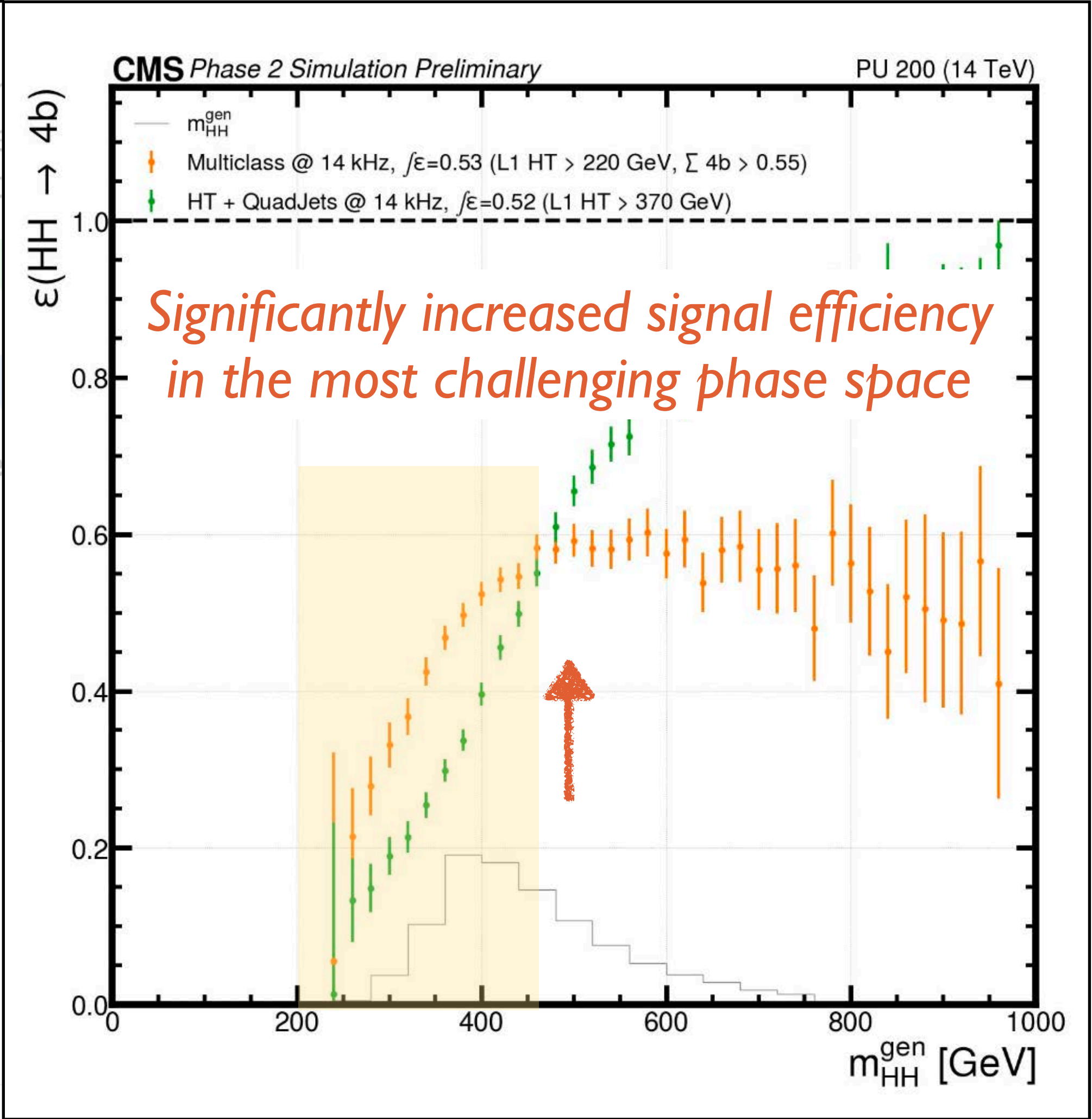
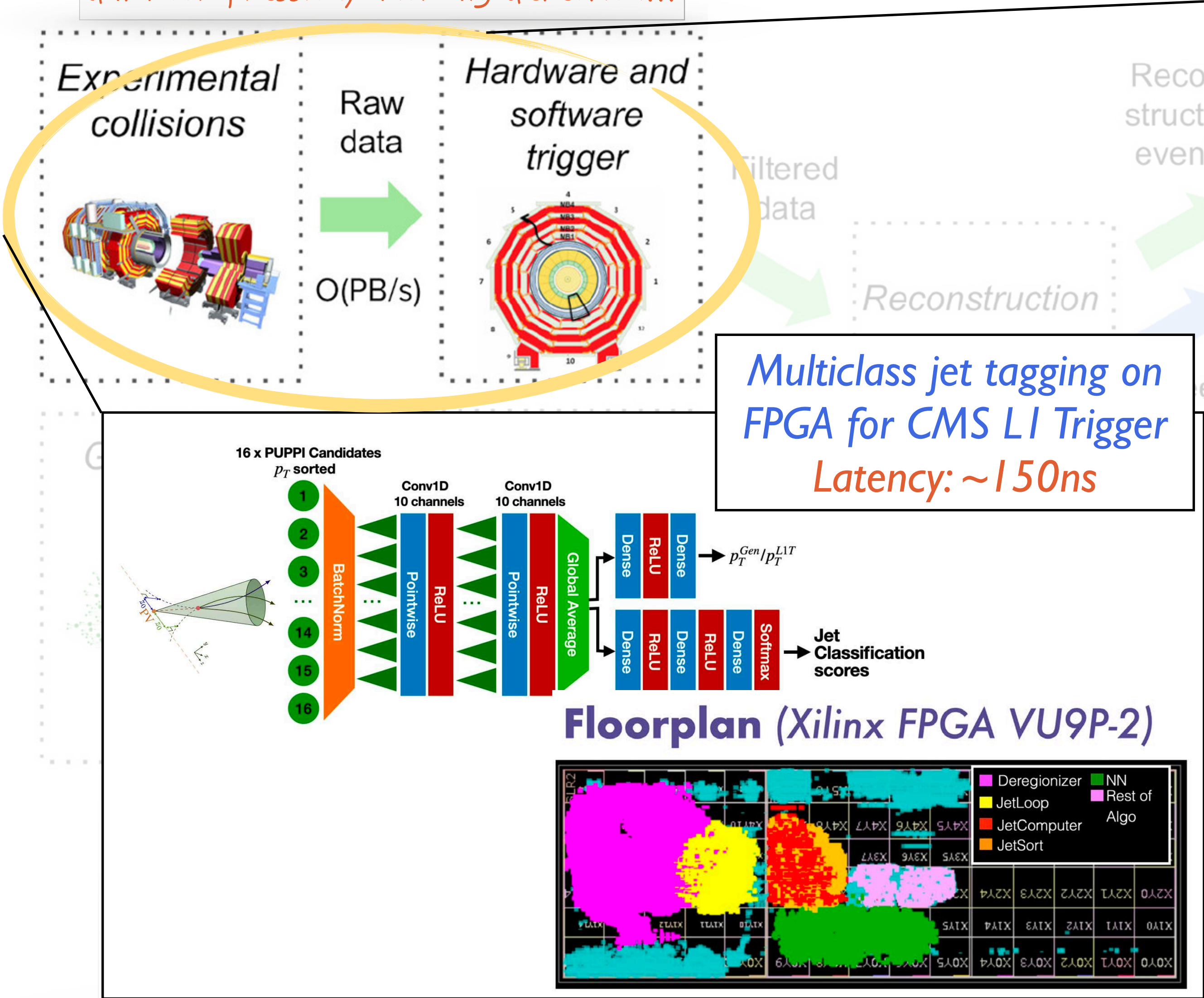




# HEP DATA FLOW... MATCHED WITH ML!

Ultrafast inference (FPGA/ASIC),  
data compression, anomaly detection...

CMS-DP-2025-032

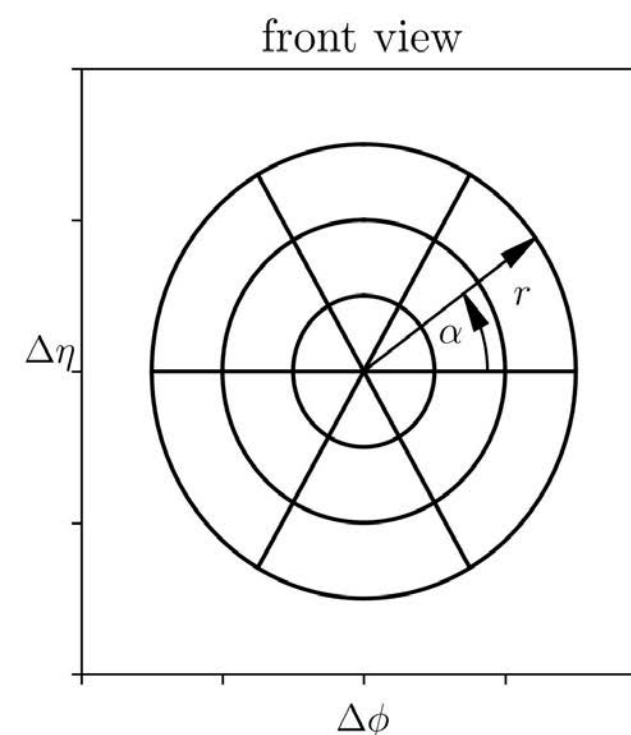
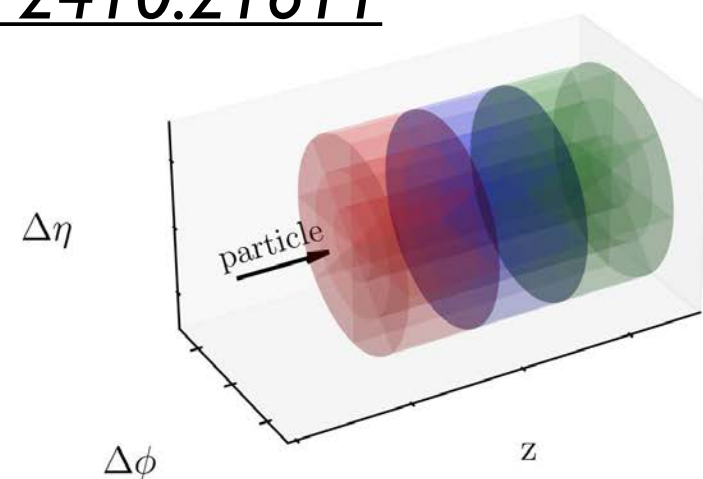




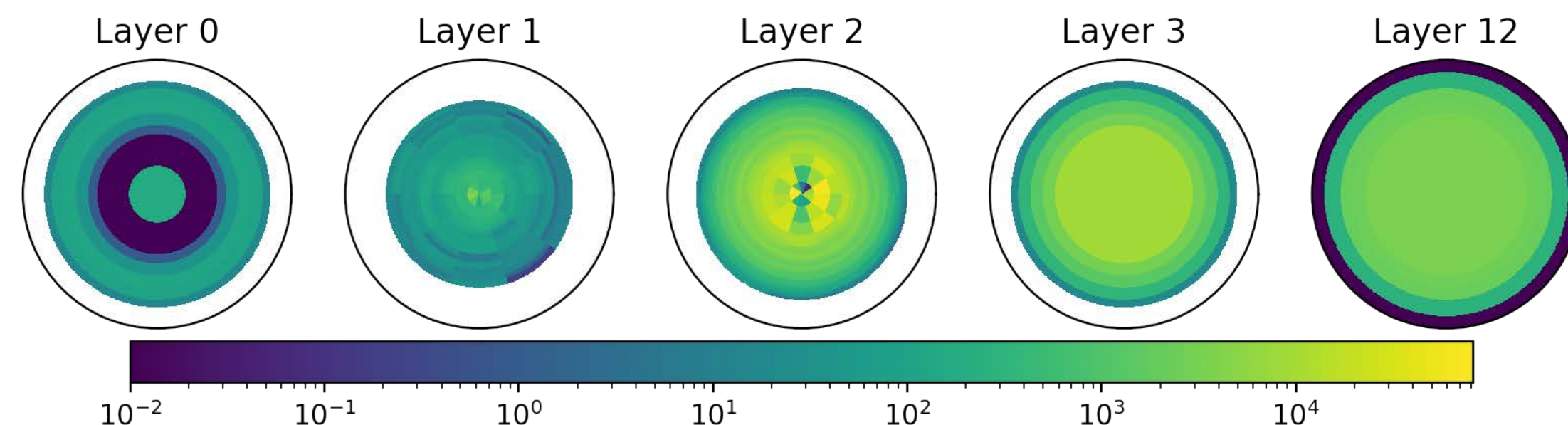
# HEP DATA FLOW... MATCHED WITH ML!

Fast calorimeter simulation  
using generative models

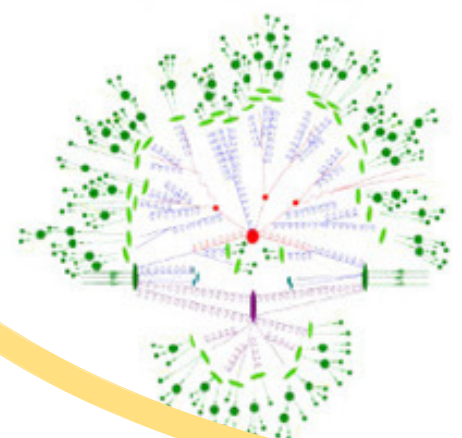
arXiv: 2410.21611



Photon shower at  $E = 1048.6$  GeV



Generated  
events



Stable  
particles



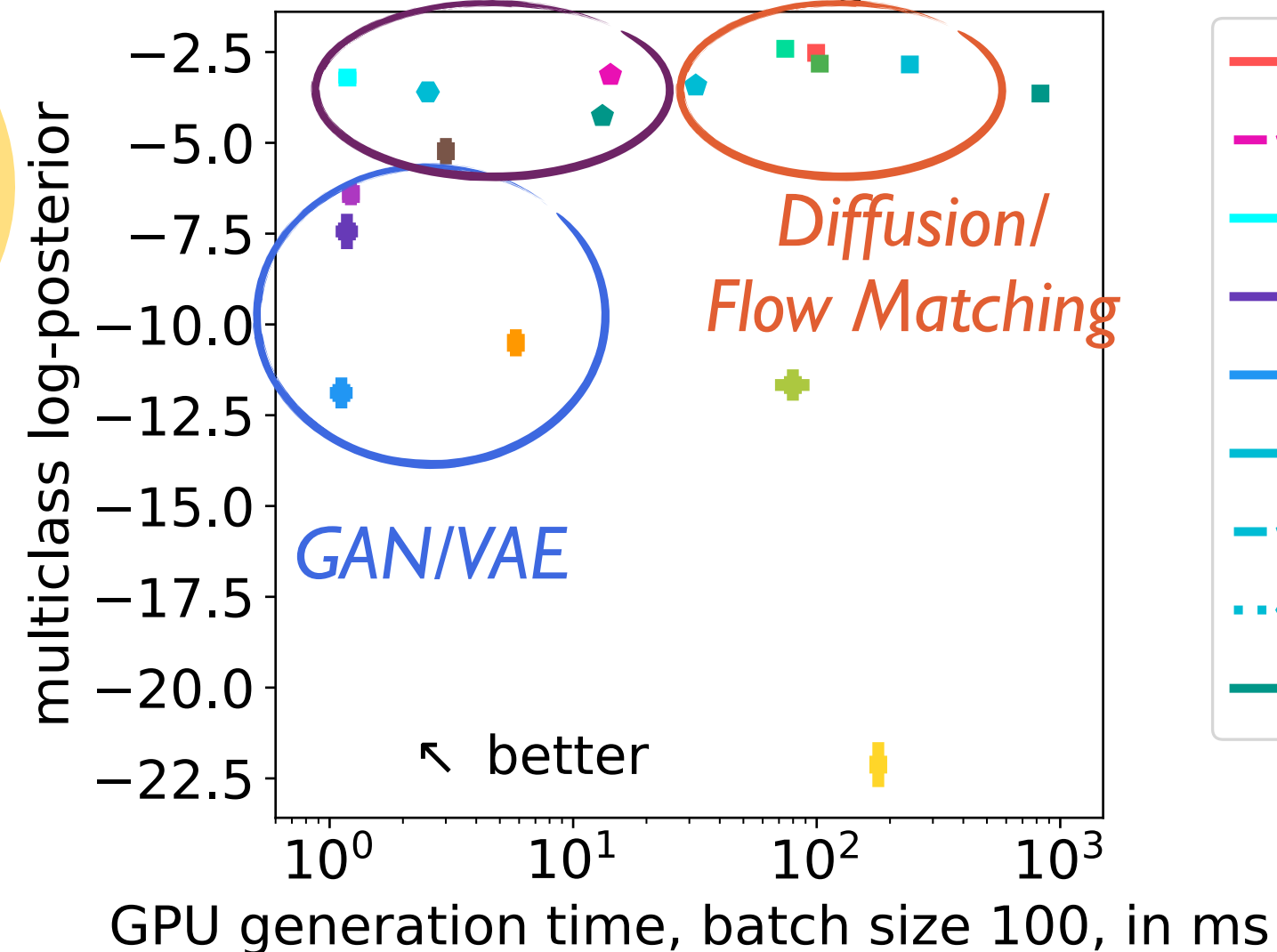
Interaction with  
simulated  
detector



Generative models for  
event generation & fast simulation:  
GAN, VAE, normalizing flow,  
diffusion...

Normalizing  
Flow

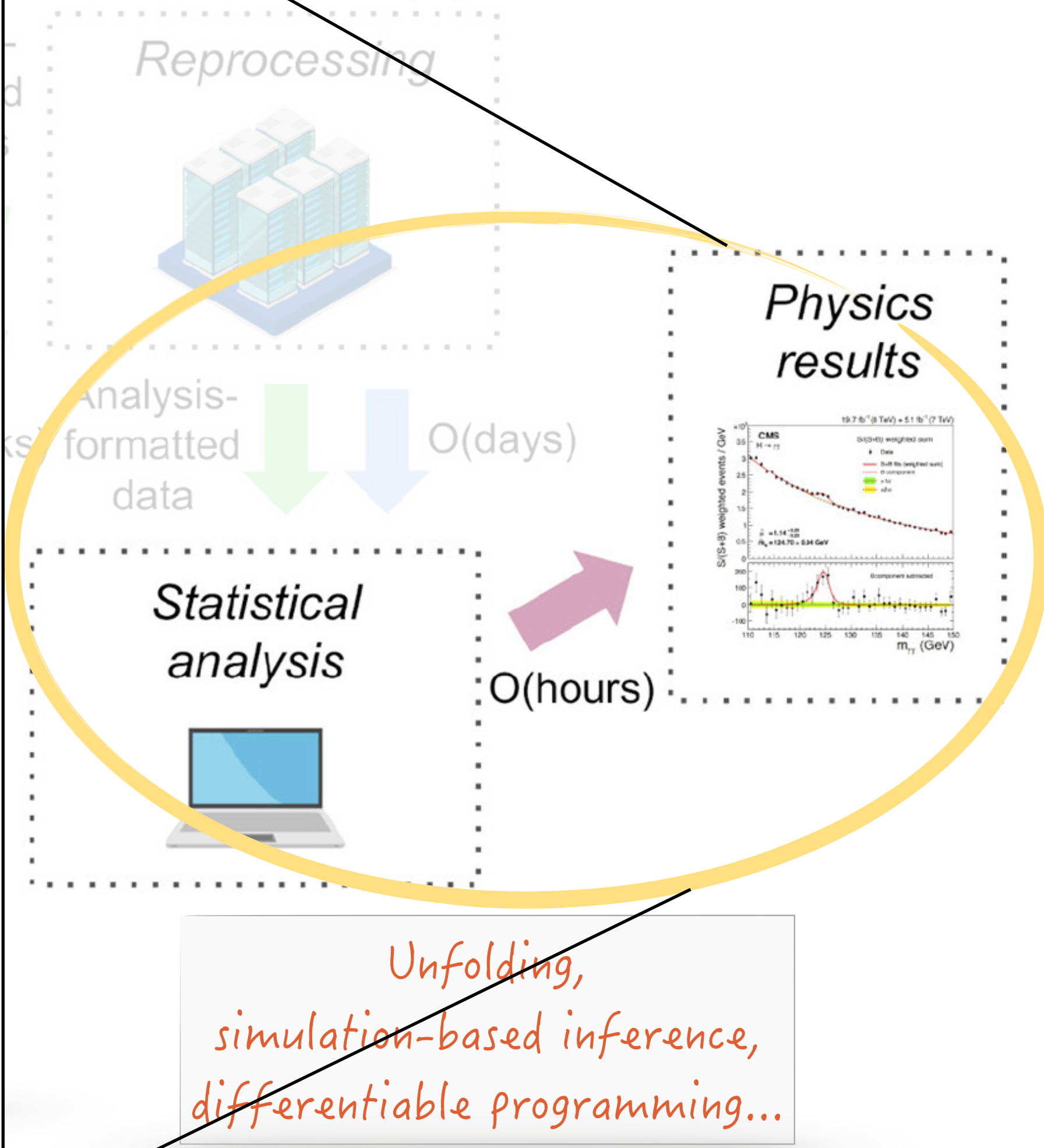
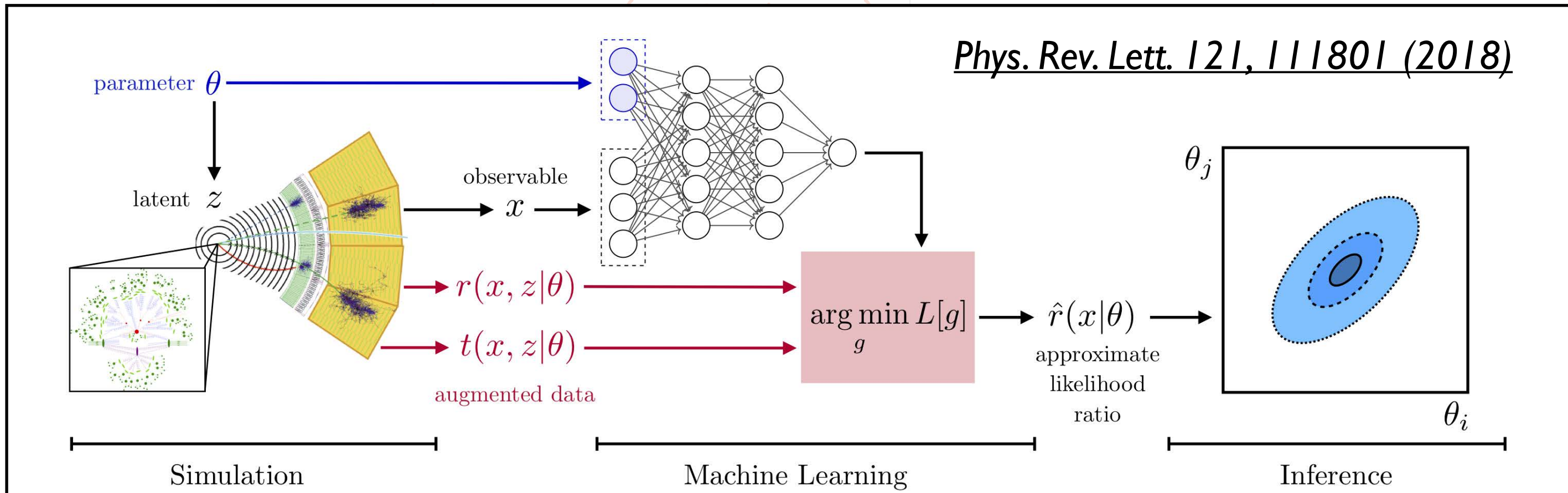
Pareto front: shower generation time on a GPU vs. multiclass log-posterior, dataset 2



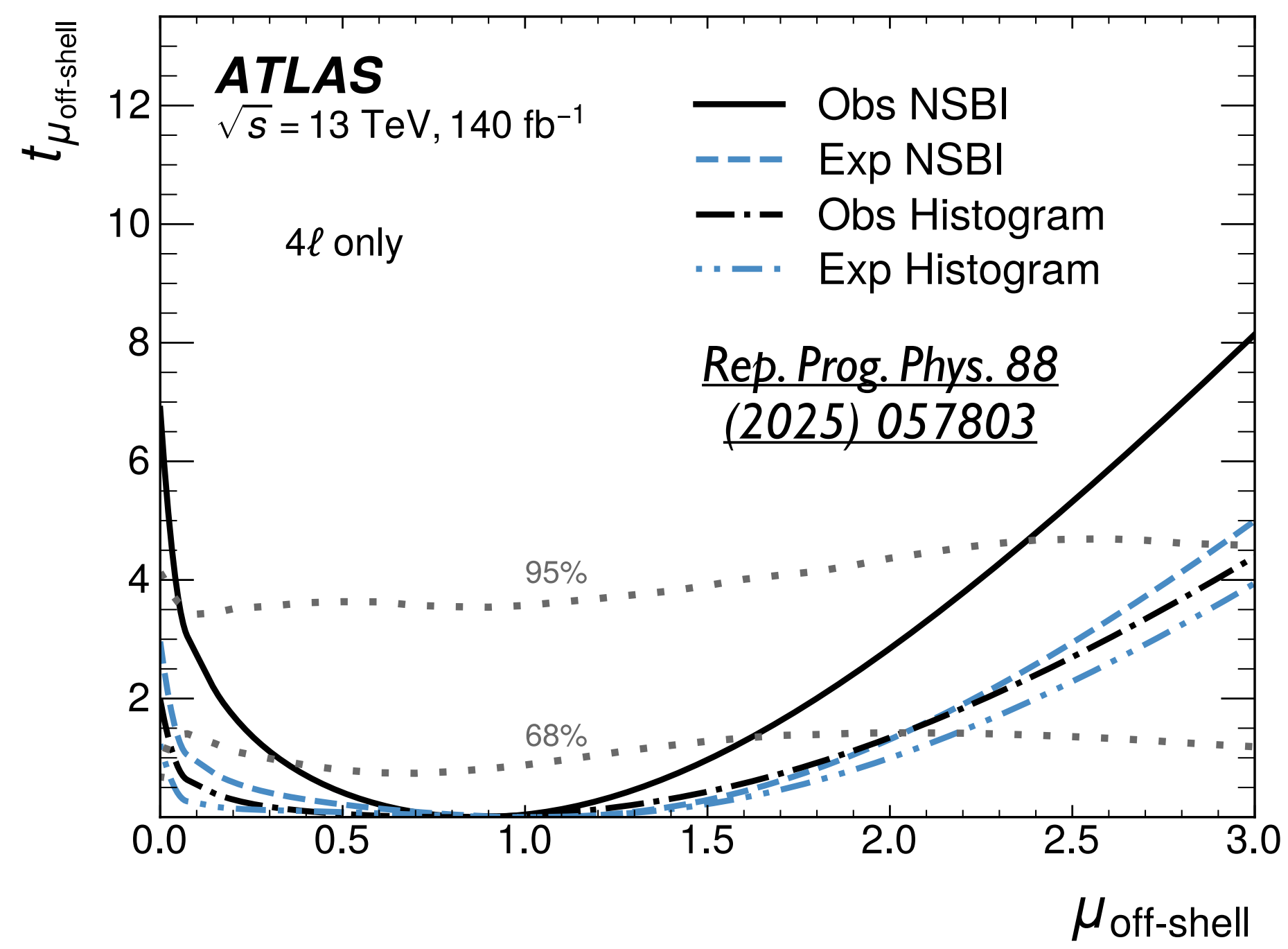
- |                       |                   |
|-----------------------|-------------------|
| CaloDiffusion         | iCaloFlow student |
| conv. L2LFlows        | SuperCalo         |
| CaloINN               | DeepTree          |
| MDMA                  | CaloPointFlow     |
| Calo-VQ               | CaloVAE+INN       |
| CaloScore             | CaloLatent        |
| CaloScore distilled   | CaloDiT           |
| CaloScore single-shot | CaloDREAM         |
| iCaloFlow teacher     |                   |



# HEP DATA FLOW... MATCHED WITH ML!



Approximating complex likelihood using neural networks for faster and more optimal statistical inference.



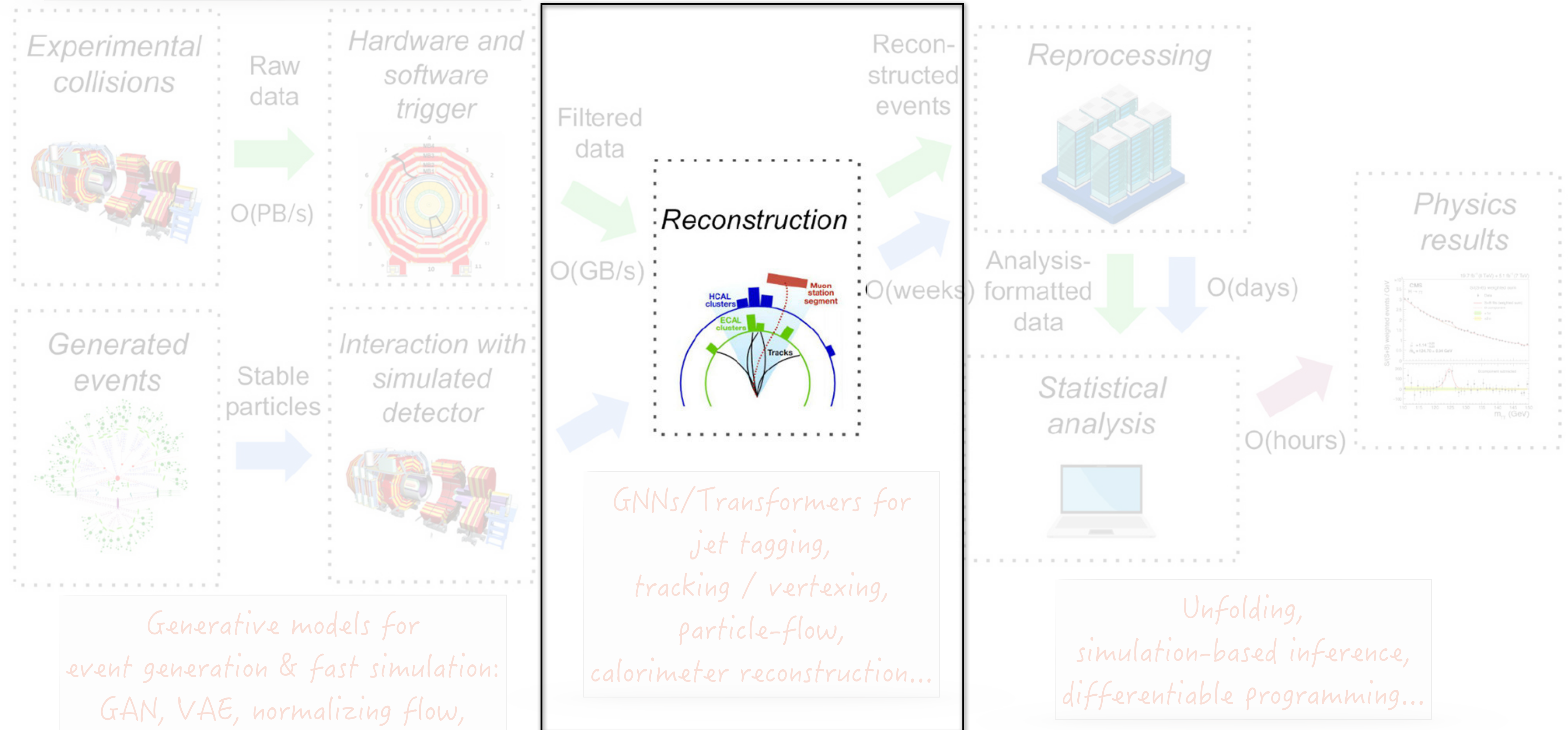
Unfolding, simulation-based inference, differentiable programming...



# HEP DATA FLOW... MATCHED WITH ML!

Ultrafast inference (FPGA/ASIC),  
data compression, anomaly detection...

Focus of today

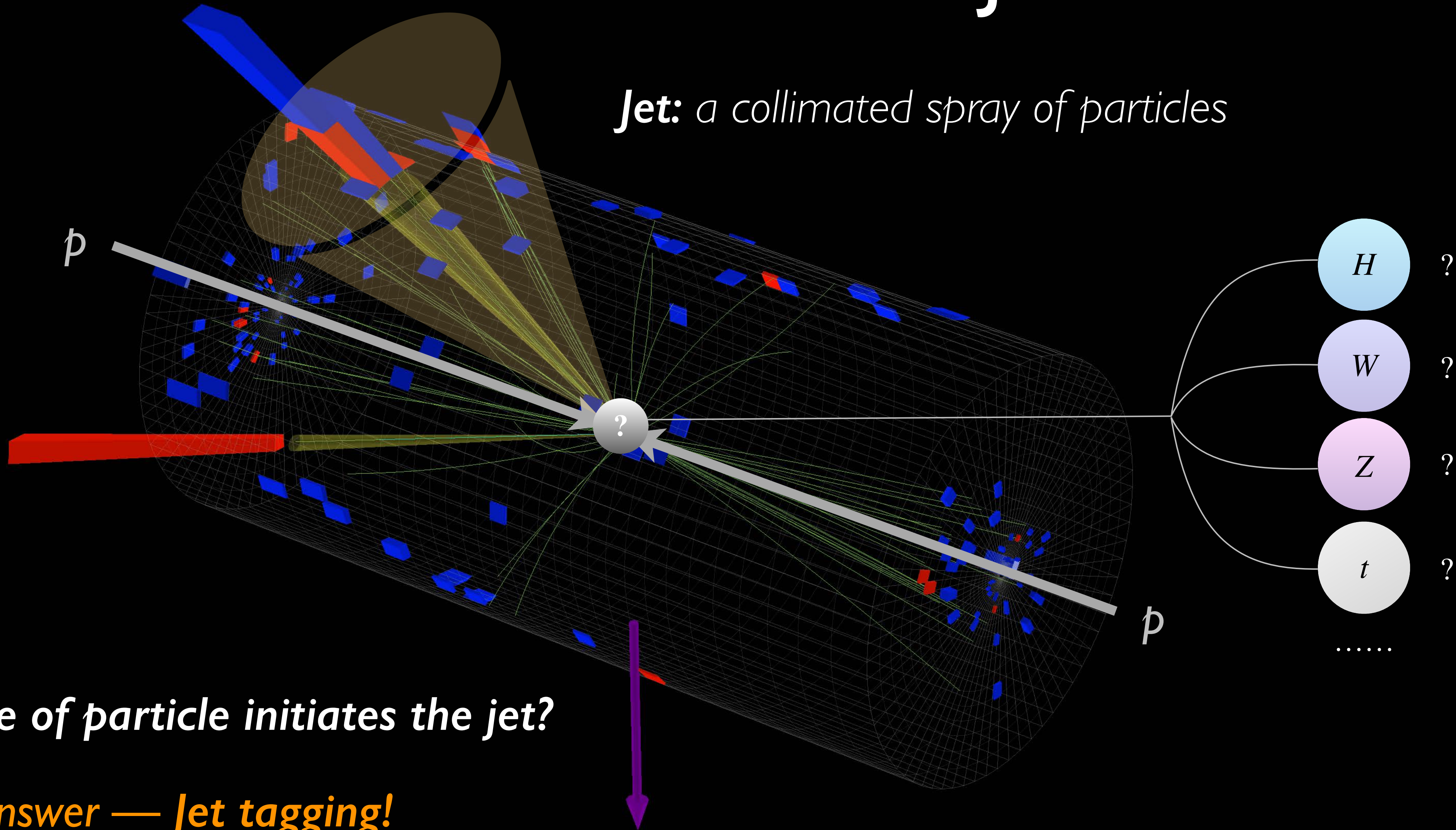






# EXAMPLE: JET TAGGING

*Jet: a collimated spray of particles*



Key question:  
***What type of particle initiates the jet?***

***The answer — Jet tagging!***

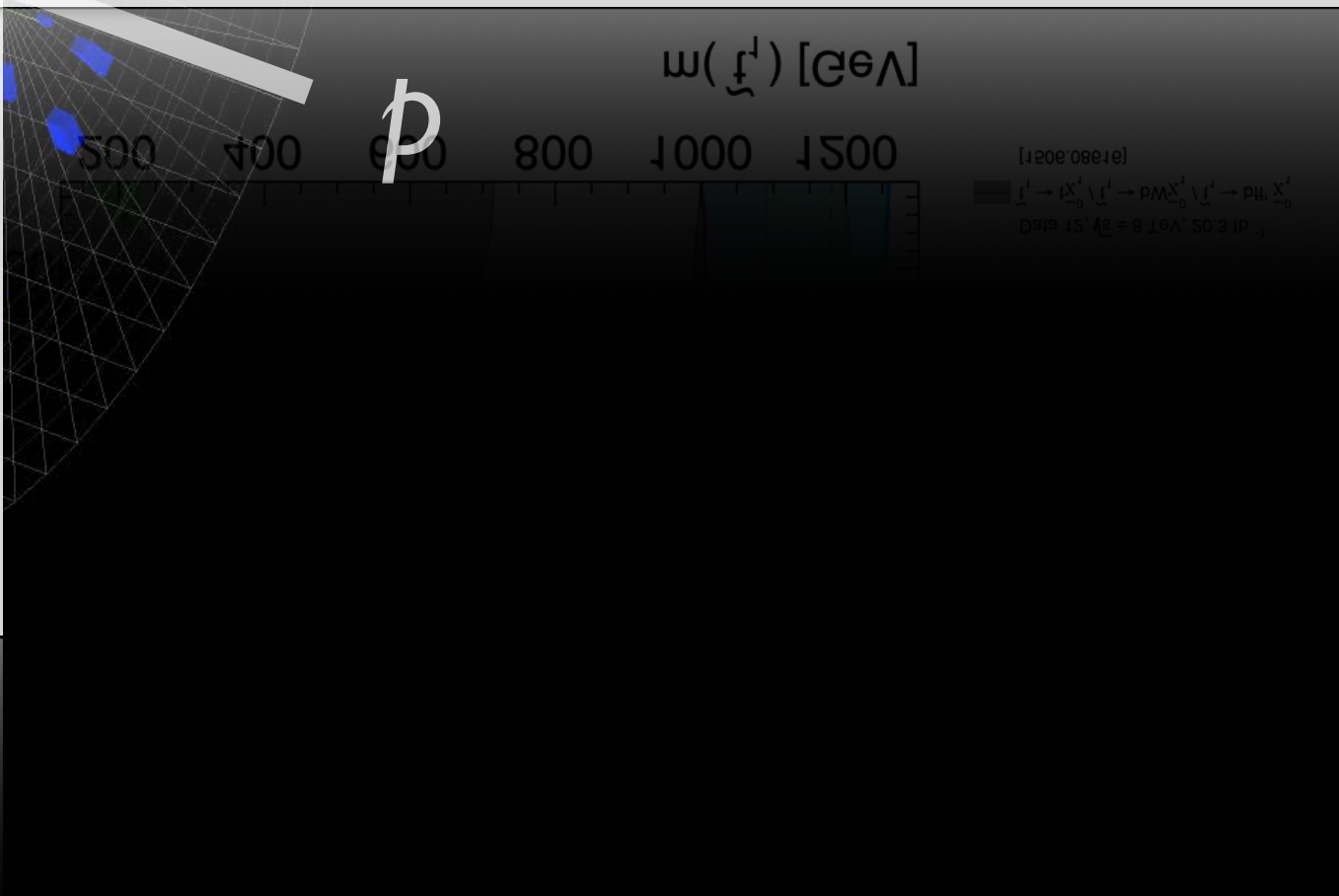
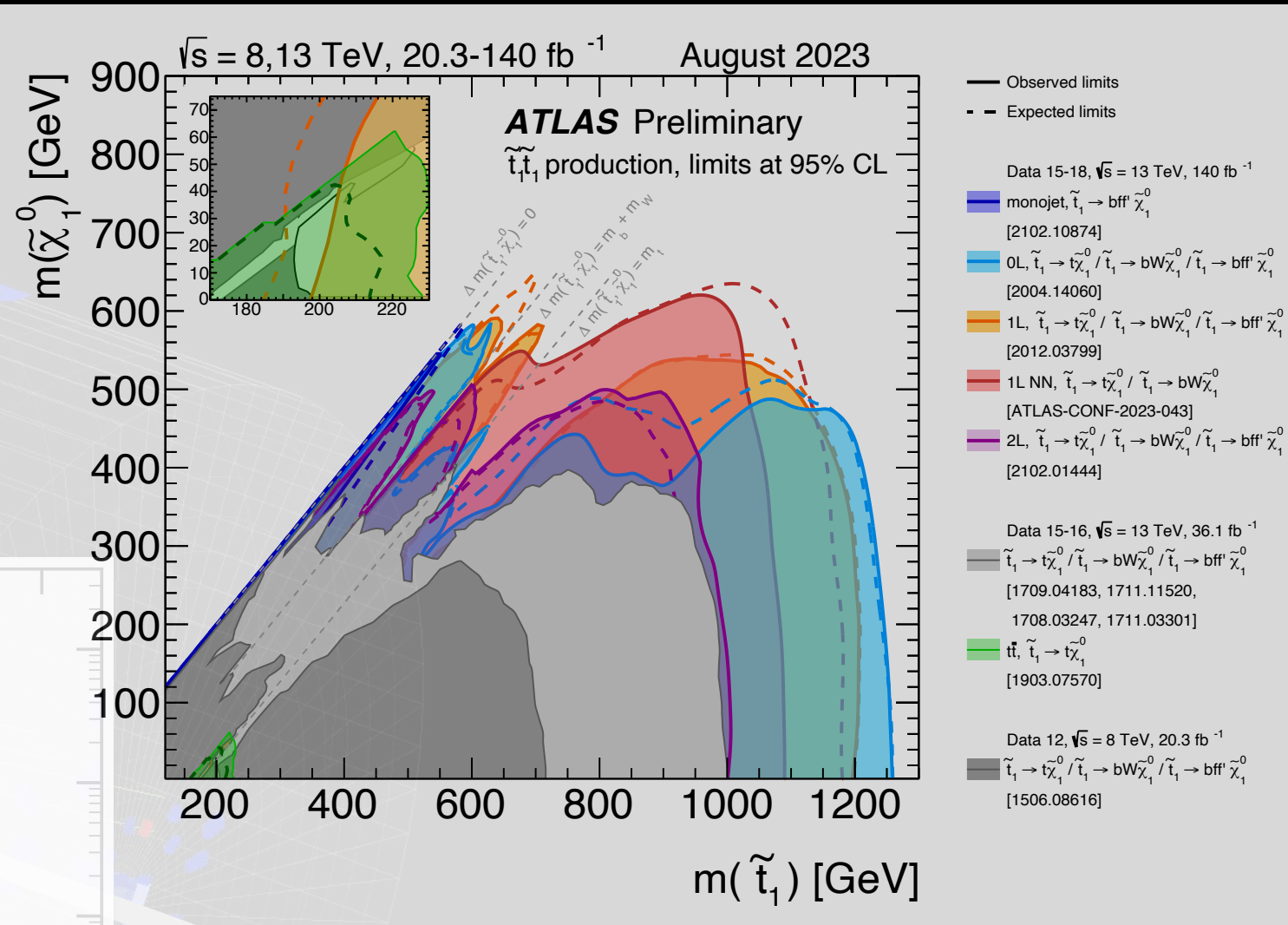
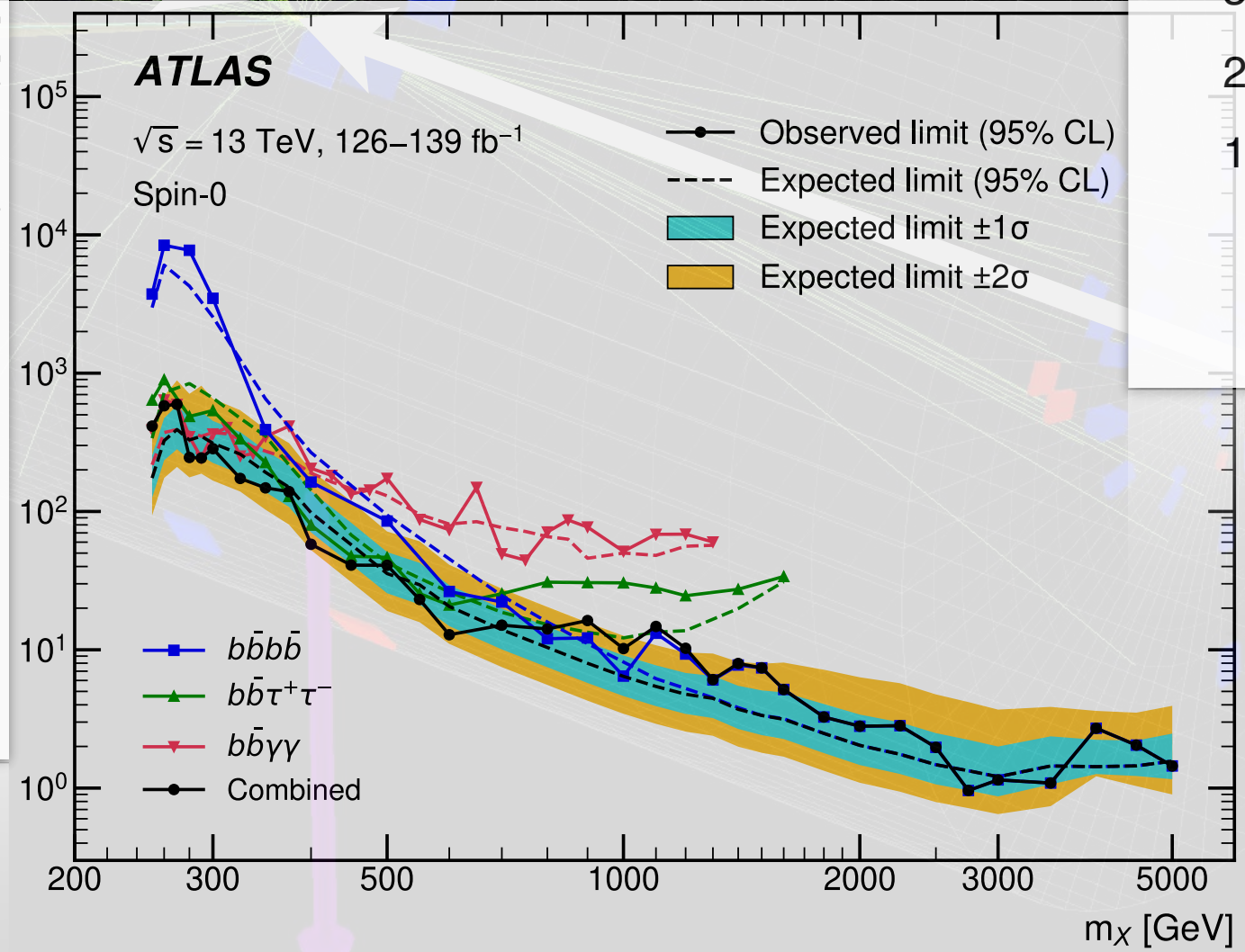
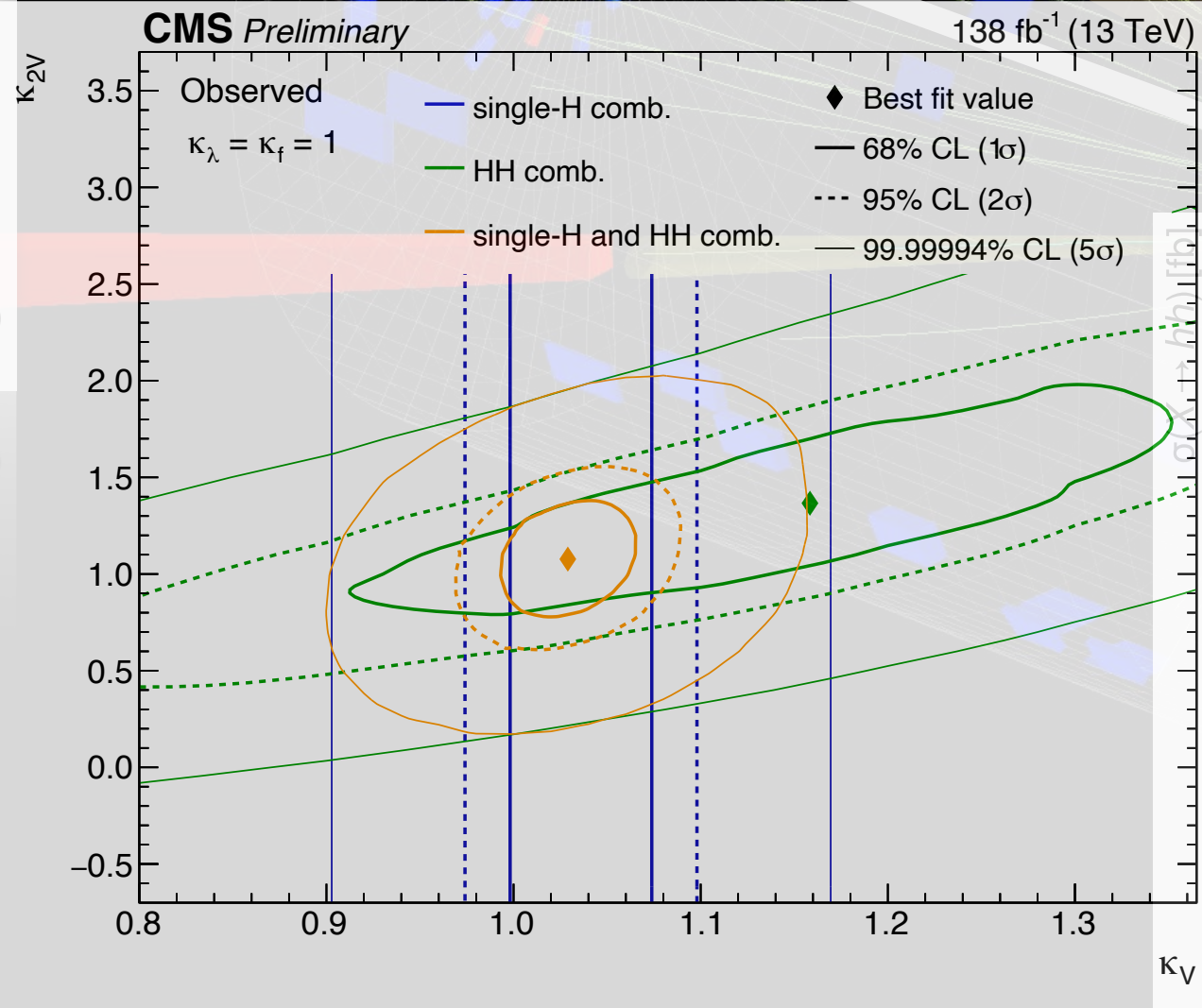
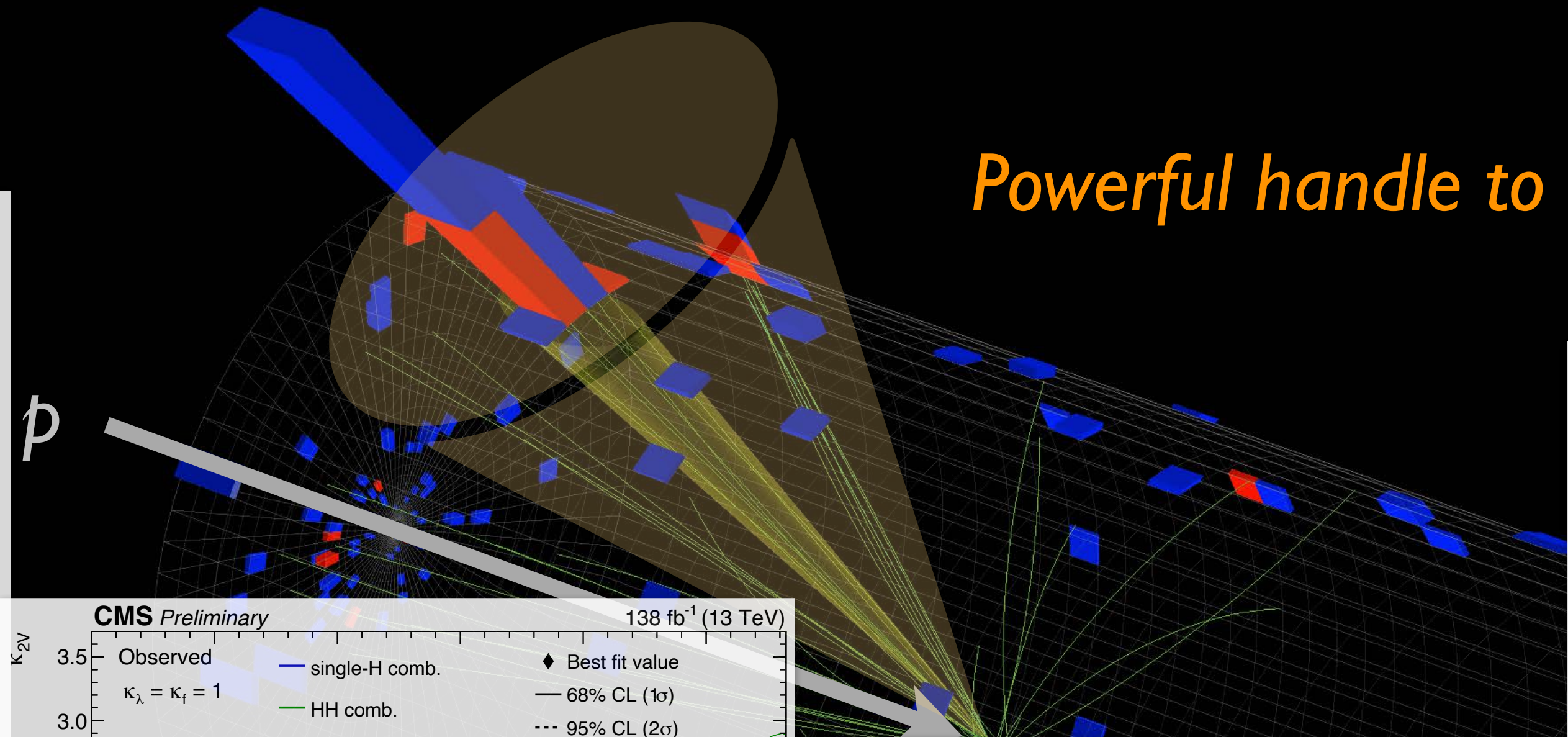
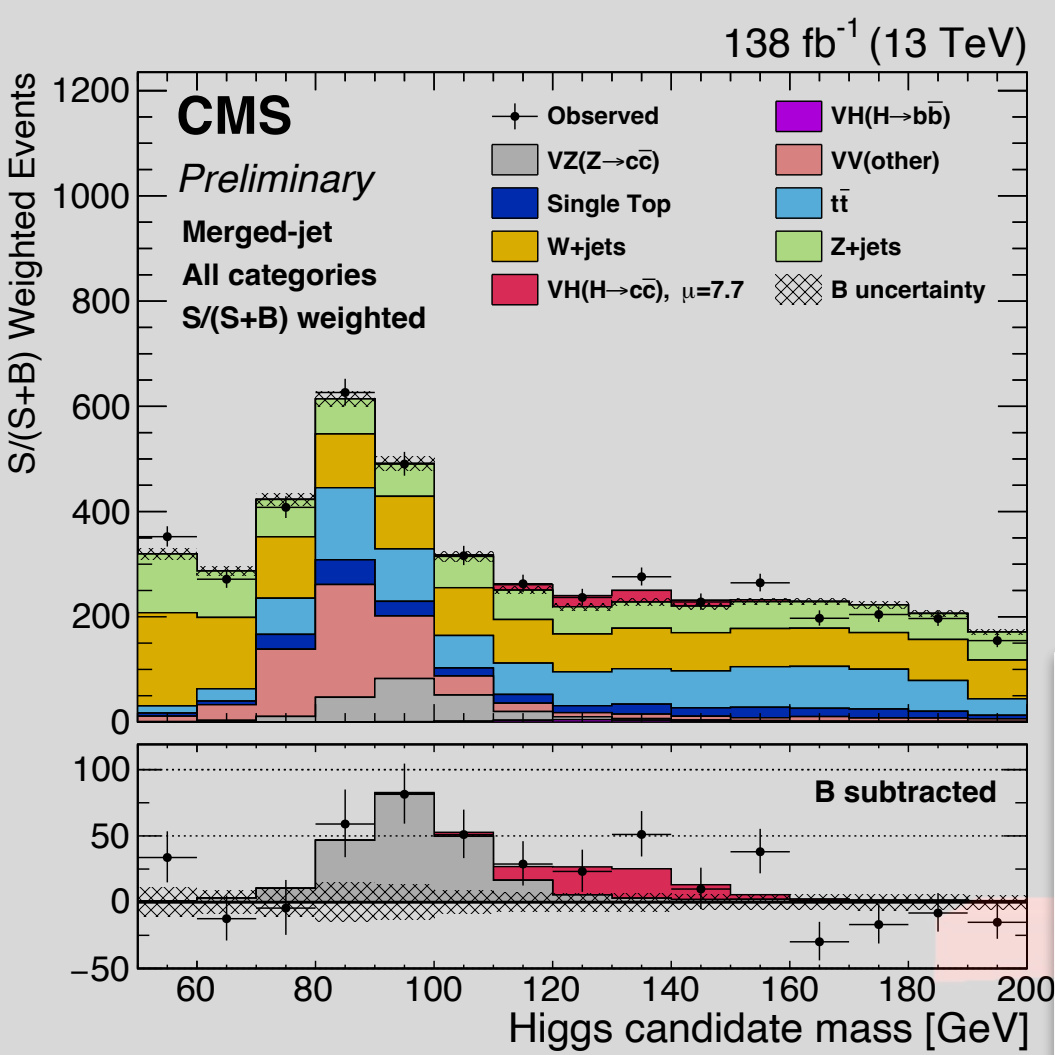




CMS Experiment at LHC, CERN  
Data recorded: Sat Aug 5 15:32:22 2017 CEST  
Run/Event: 300515 / 205888132

# EXAMPLE: JET TAGGING

Powerful handle to search for new phenomena



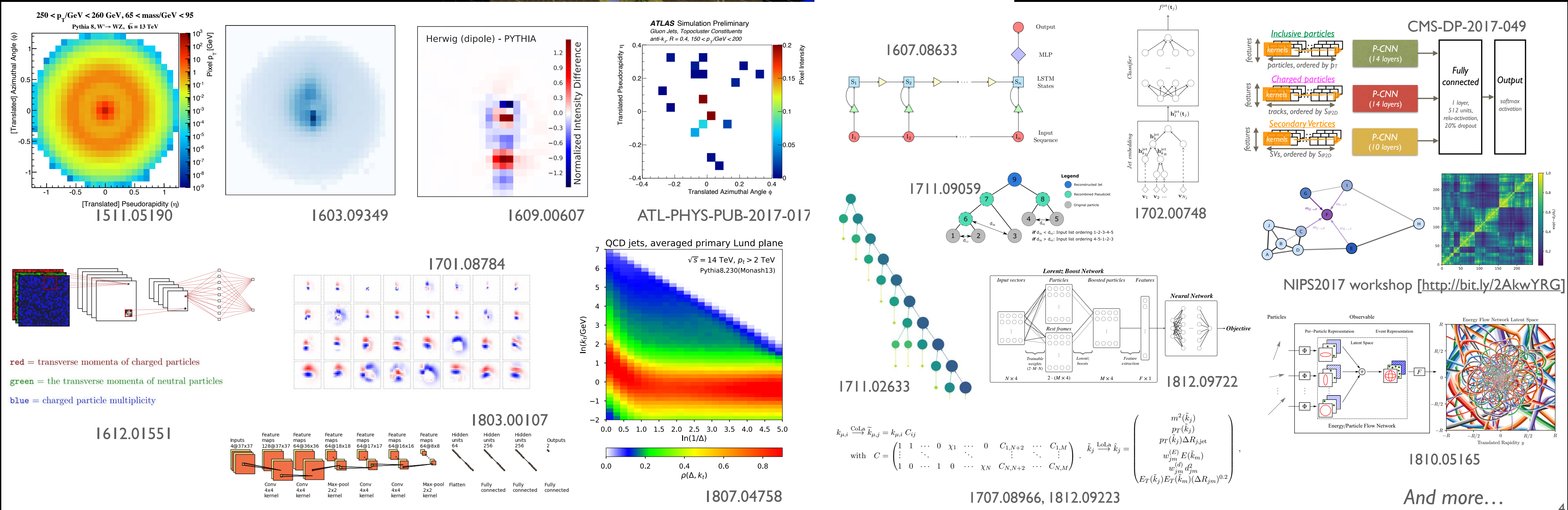




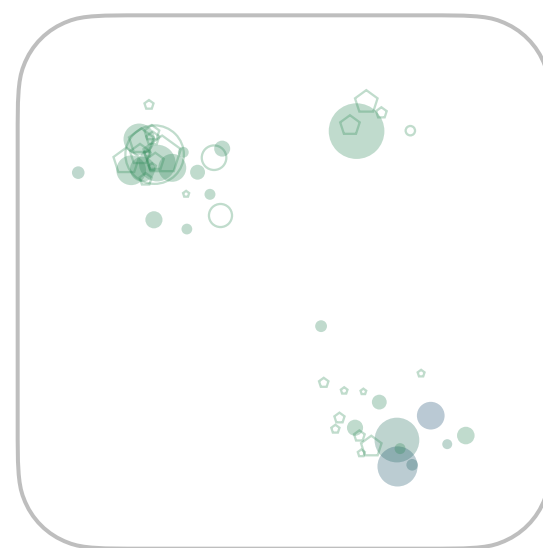
CMS Experiment at LHC, CERN  
Data recorded: Sat Aug 5 15:32:22 2017 CEST  
Run/Event: 300515 / 205888132

# EXAMPLE: JET TAGGING

But *event playground for machine learning*

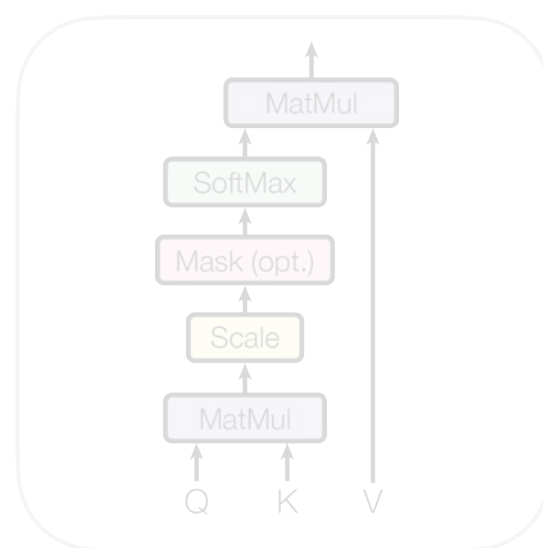






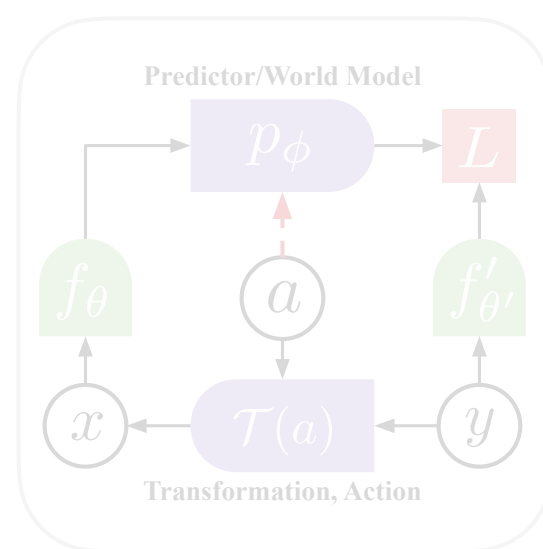
## Evolution of Jet Representations

*image, sequence, point cloud, ...*



## Attention Is All You Need?

*or how physics can help...*



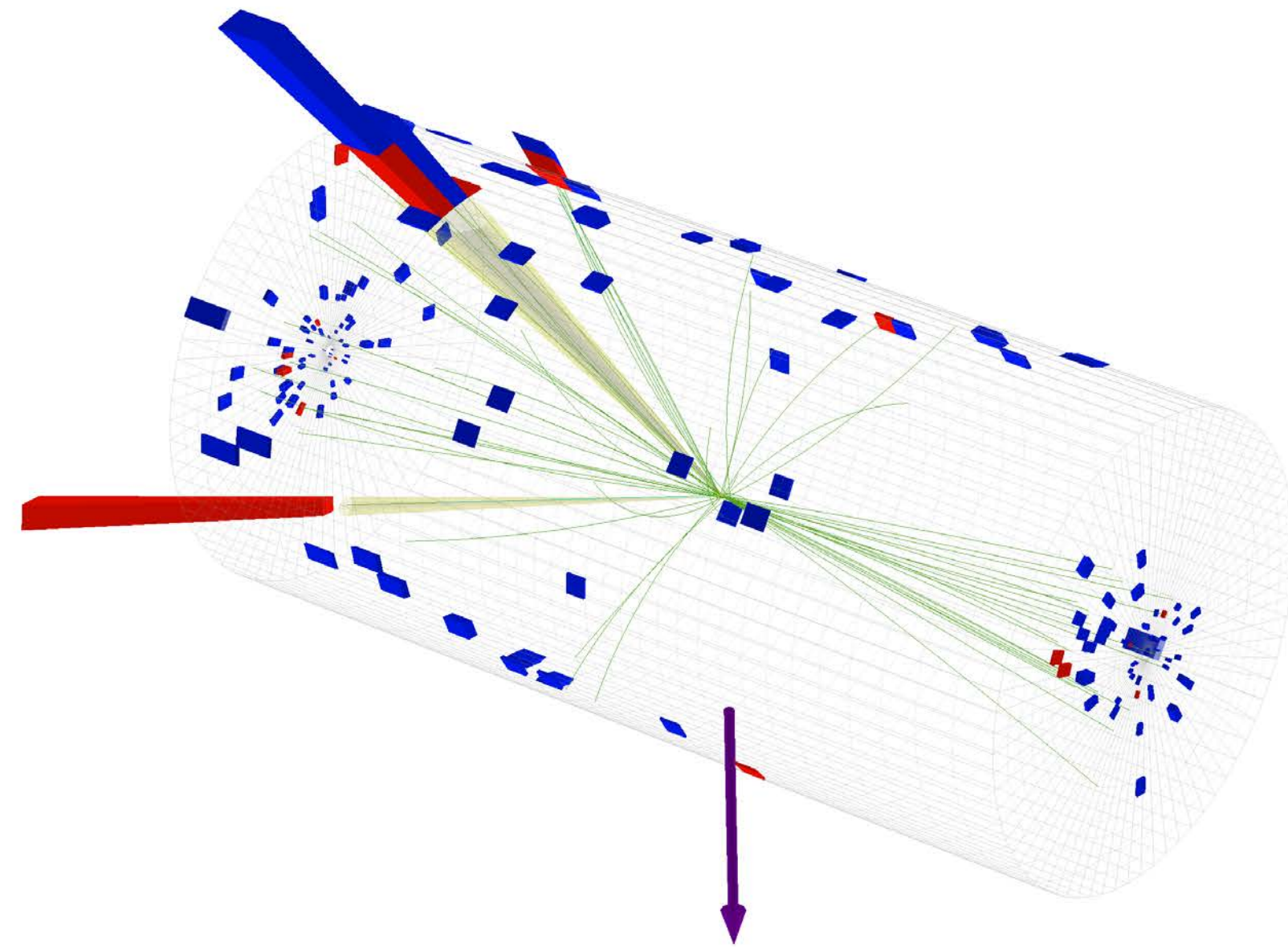
## Towards a Foundation Model

*scaling, self-supervision, and more...*



# DATA REPRESENTATION

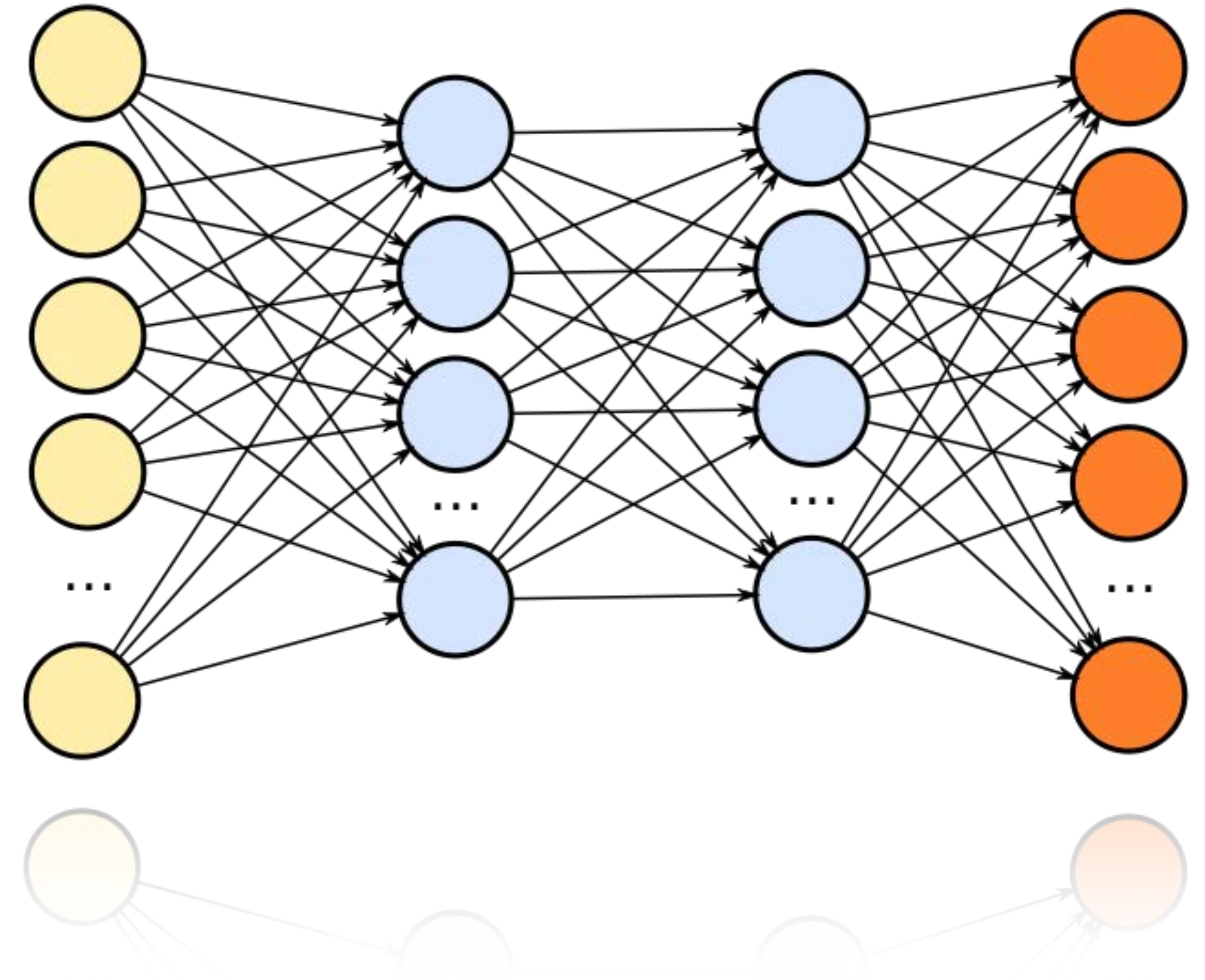
*HEP*



*Collision events, detector hits, sensor arrays, ...*

**X**

*ML*

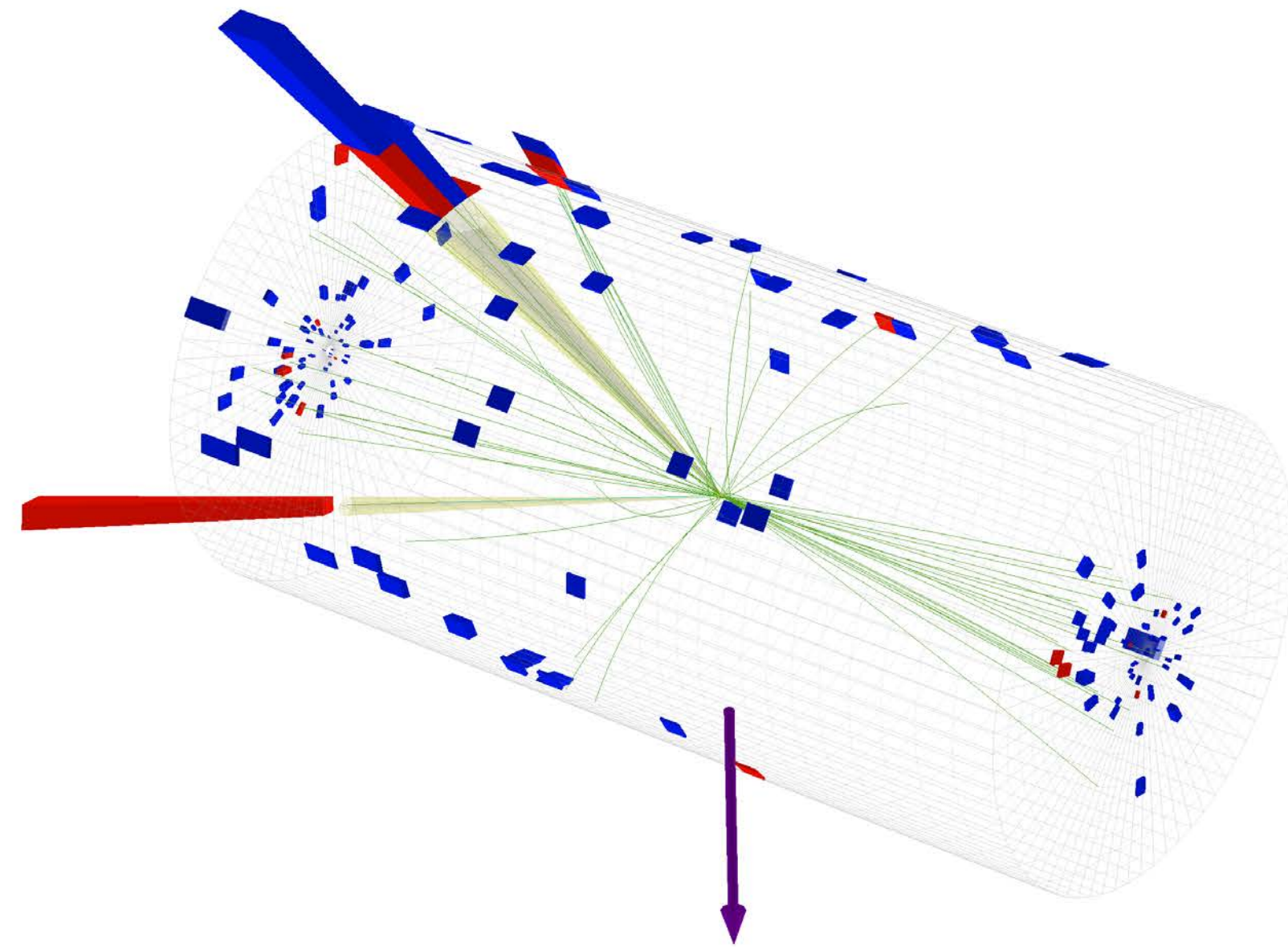


*First and foremost:  
How to represent the data?*



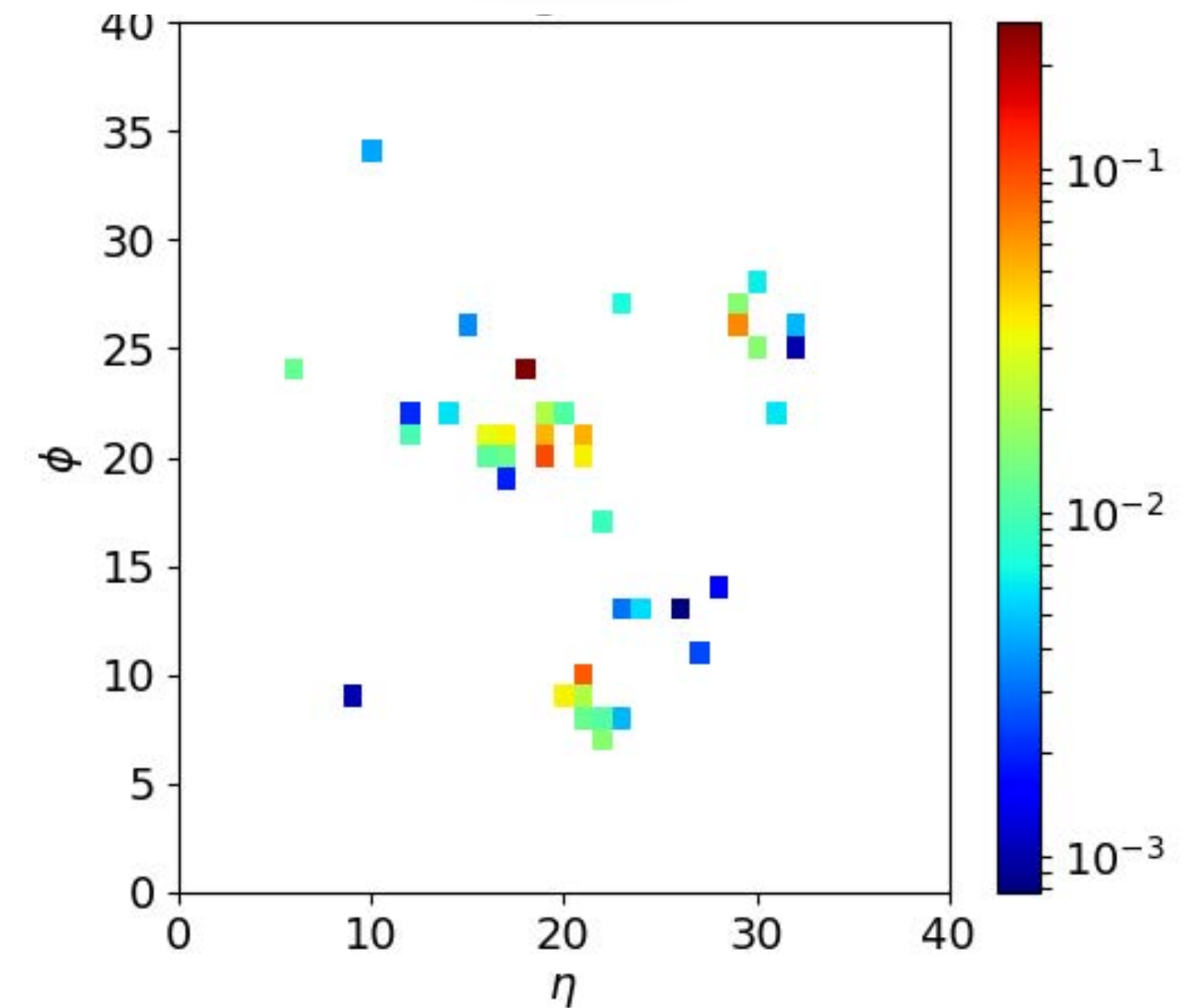
# DATA REPRESENTATION: IMAGE

*HEP*



*Collision events, detector hits, sensor arrays, ...*

*Image*



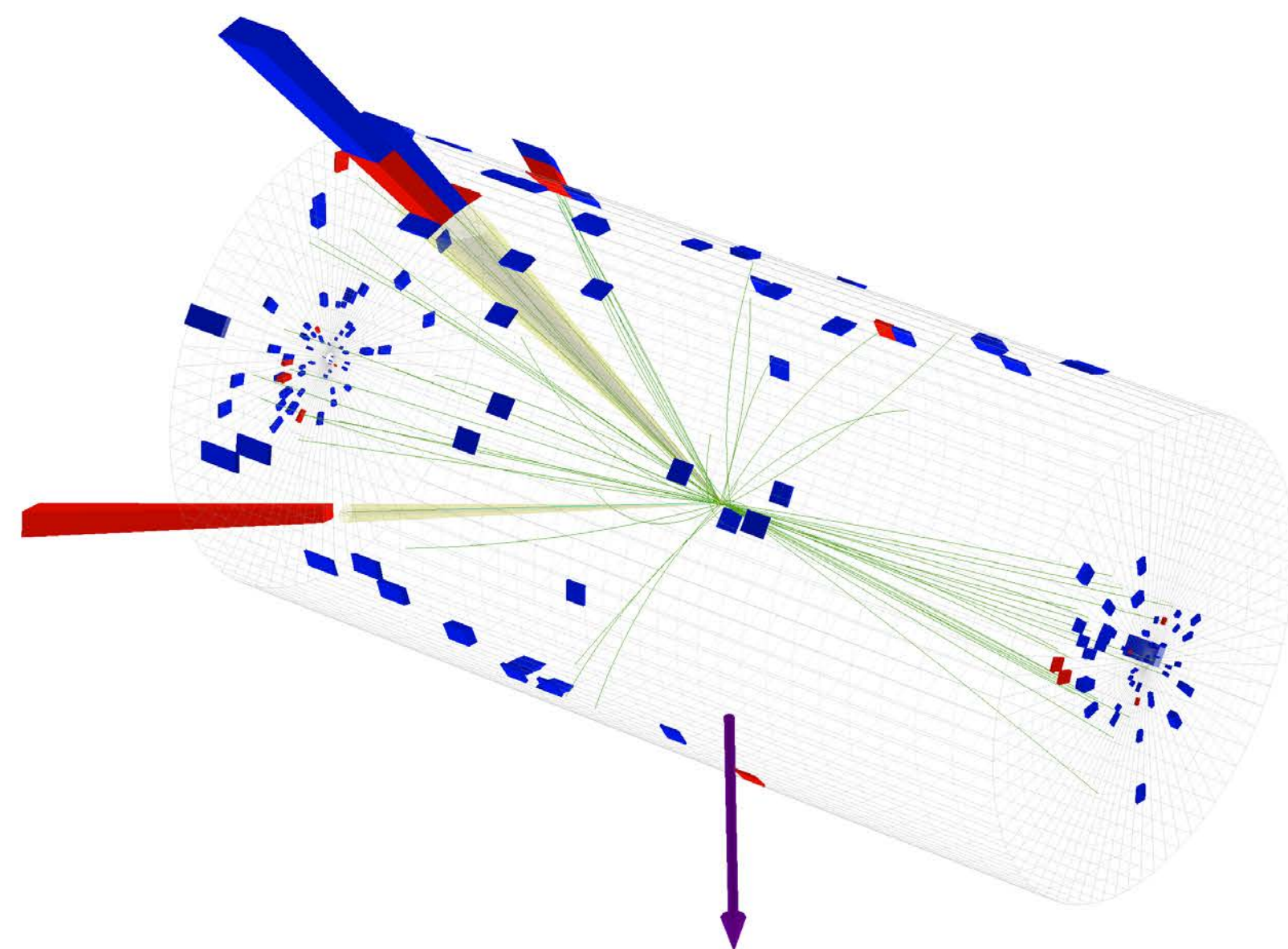
e.g., L. de Oliveira, M. Kagan, L. Mackey, B. Nachman and A. Schwartzman, arXiv: 1511.05190  
(see also the review by Kagan, arXiv:2012.09719)

- Convert to 2D/3D image => **Computer vision**
  - then use convolutional neural networks (CNNs)
  - but: inhomogeneous geometry, high sparsity, ...



# DATA REPRESENTATION: SEQUENCE

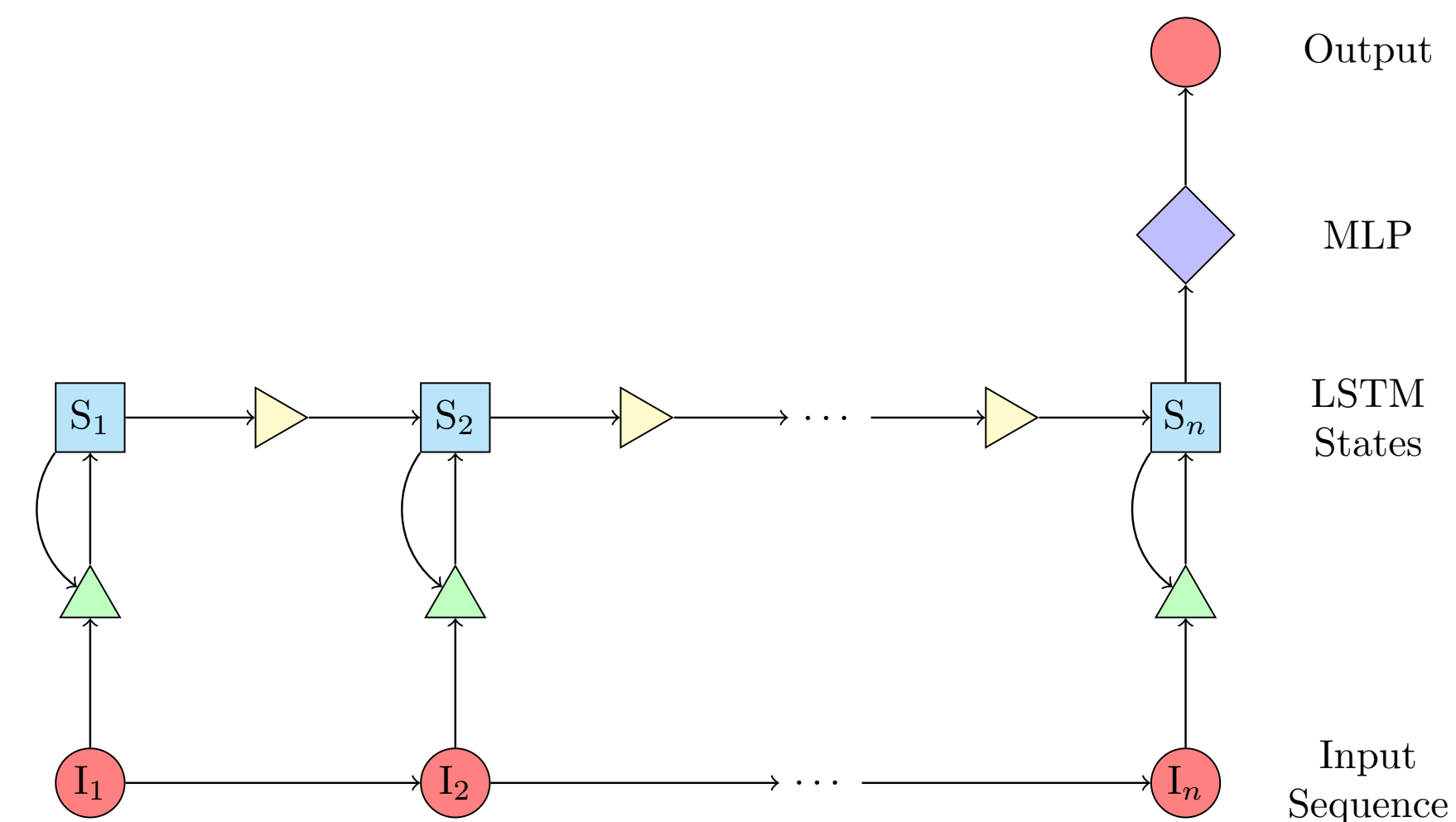
*HEP*



*Collision events, detector hits, sensor arrays, ...*



*Sequence*

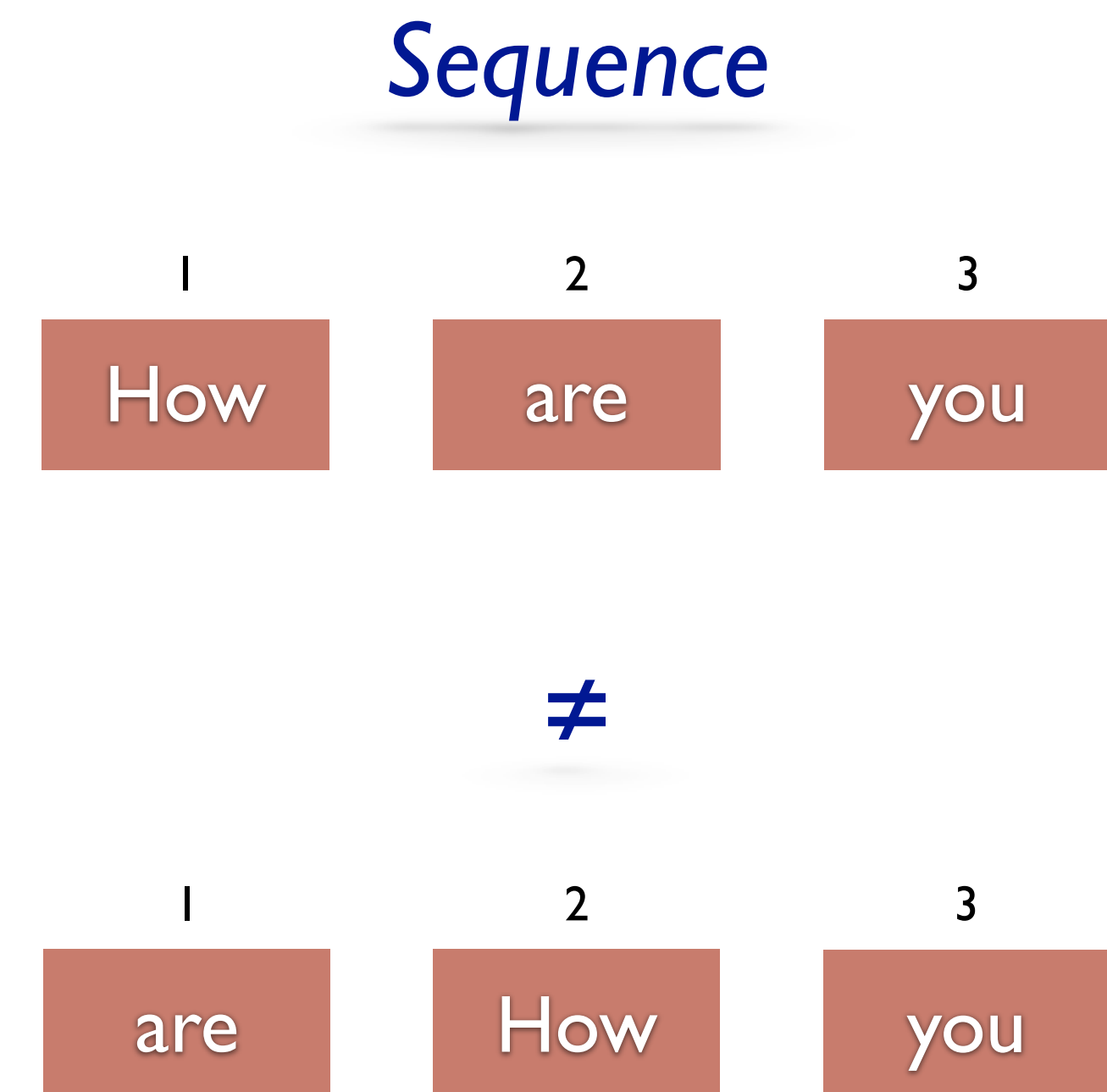
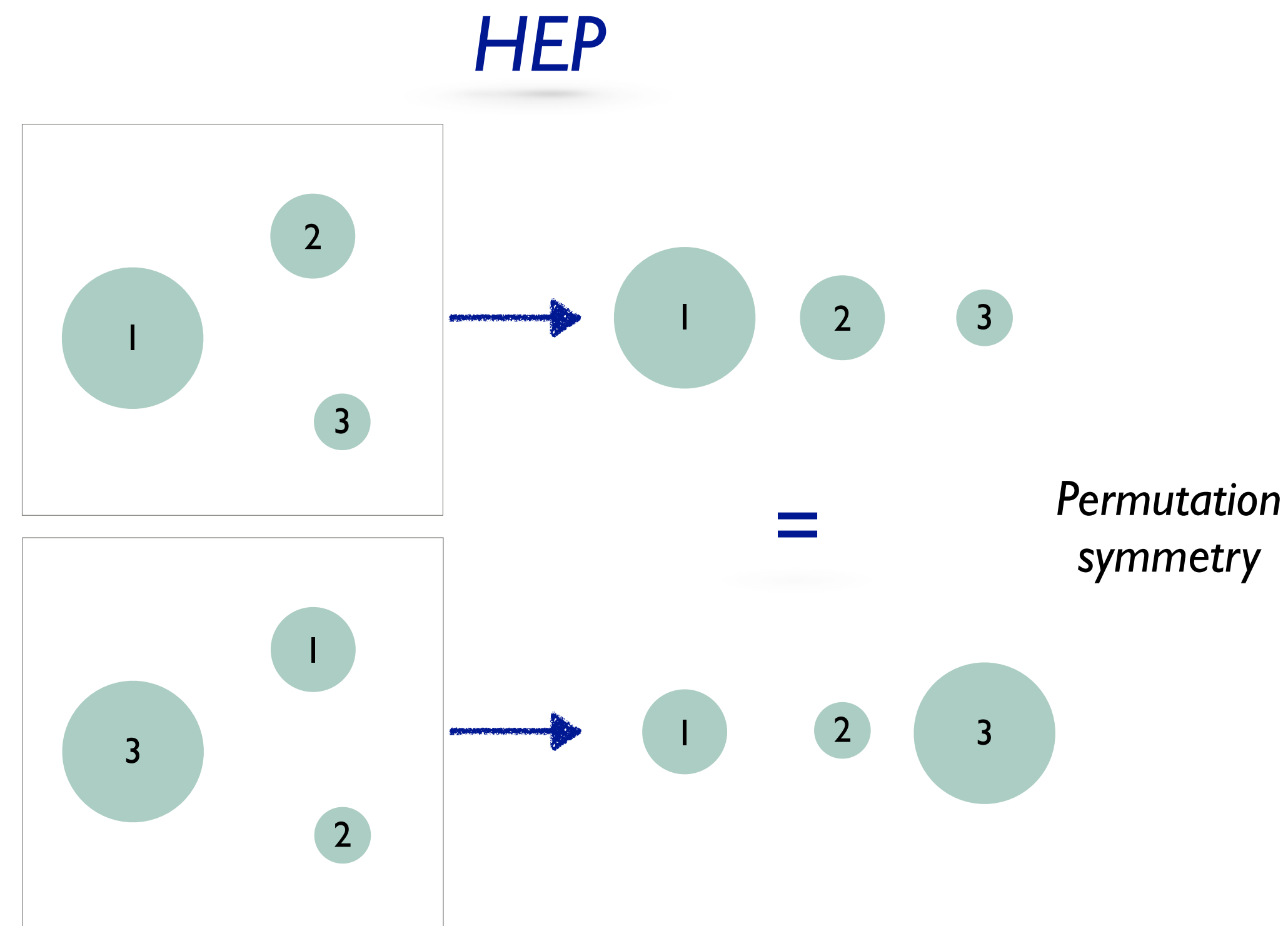


e.g., Guest, Collado, Baldi, Hsu, Urban, Whiteson  
arXiv: 1607.08633

- Convert to a sequence => **Natural language processing (NLP)**
  - recurrent neural network (RNN), e.g., GRU/LSTM; 1D CNNs; etc.



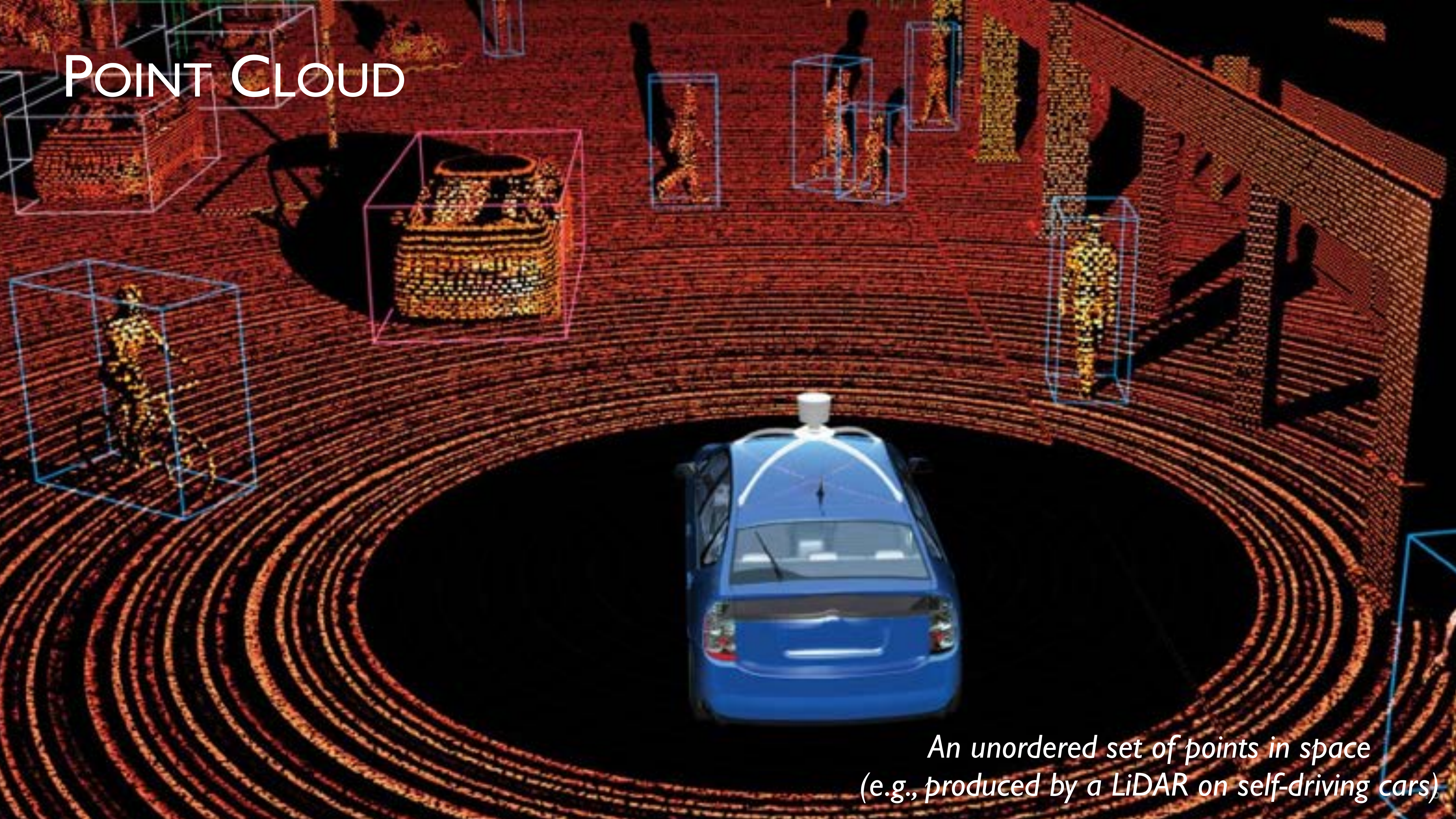
# DATA REPRESENTATION: SEQUENCE?



- Convert to a sequence => **Natural language processing (NLP)**
  - recurrent neural network (RNN), e.g., GRU/LSTM; 1D CNNs; etc.
  - but: must impose an **ordering** on the particles/hits, which can limit the learning performance



# POINT CLOUD

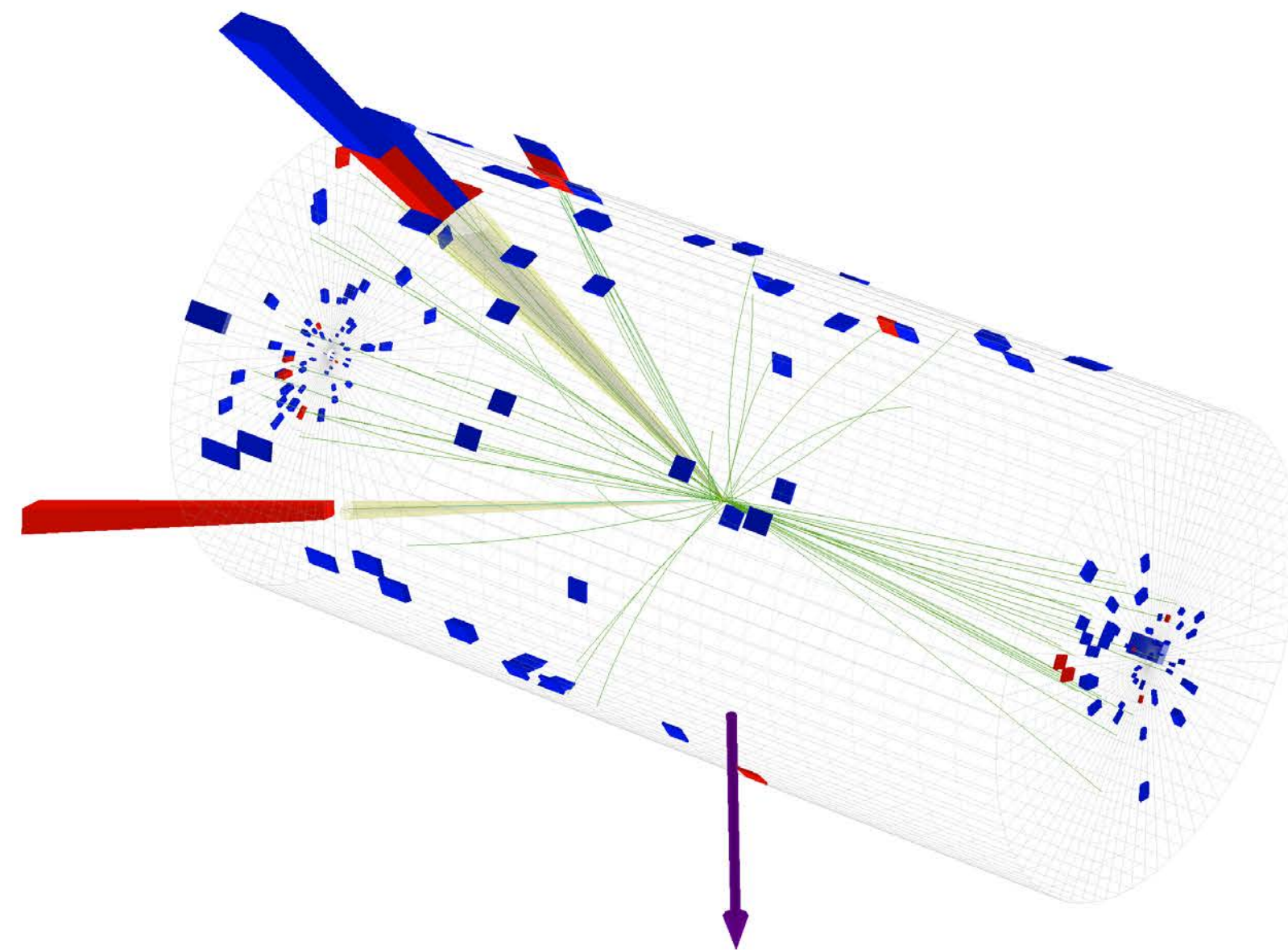


*An unordered set of points in space  
(e.g., produced by a LiDAR on self-driving cars)*



# DATA REPRESENTATION: POINT CLOUD

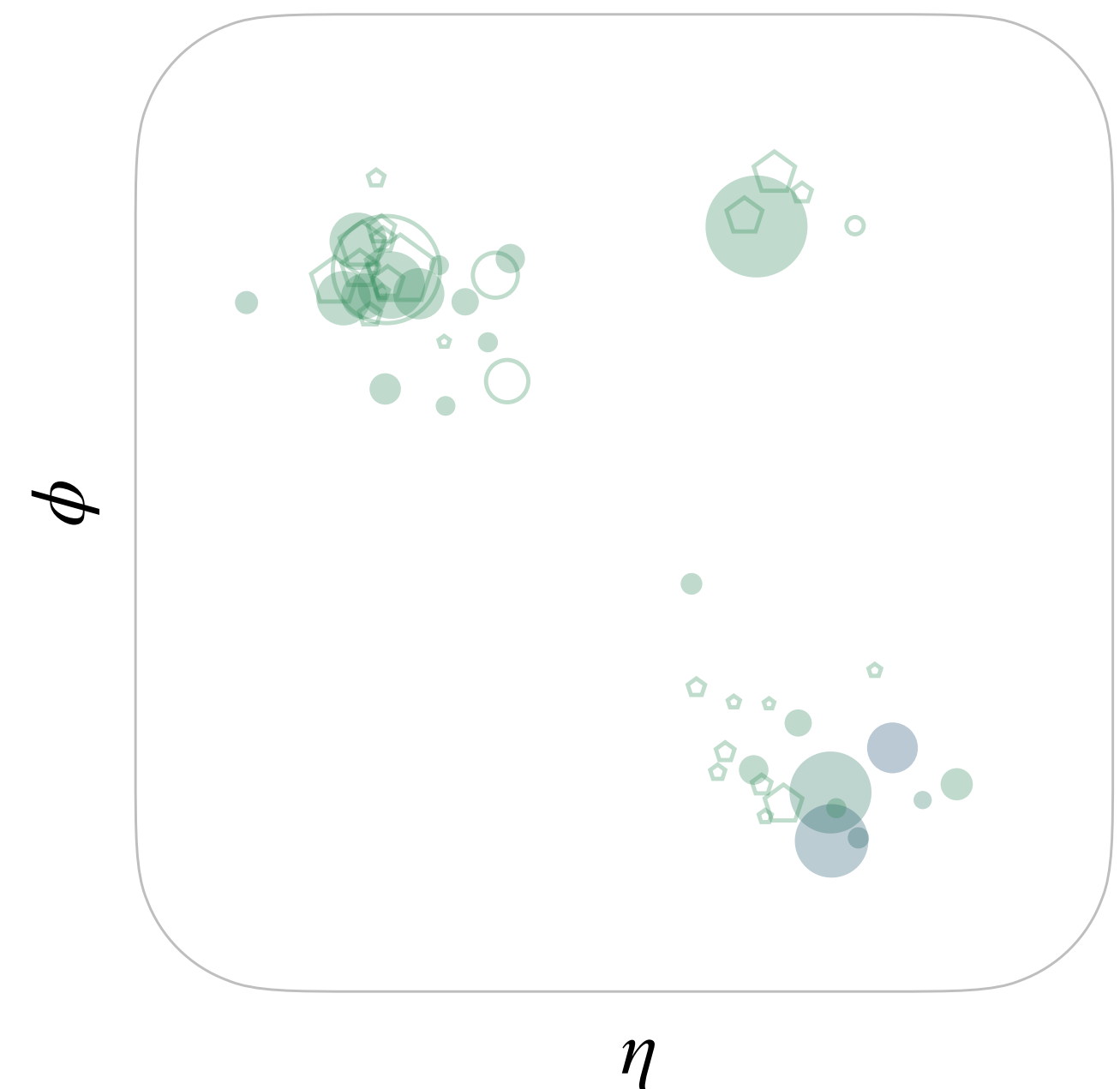
*HEP*



*Collision events, detector hits, sensor arrays, ...*



*Point cloud*



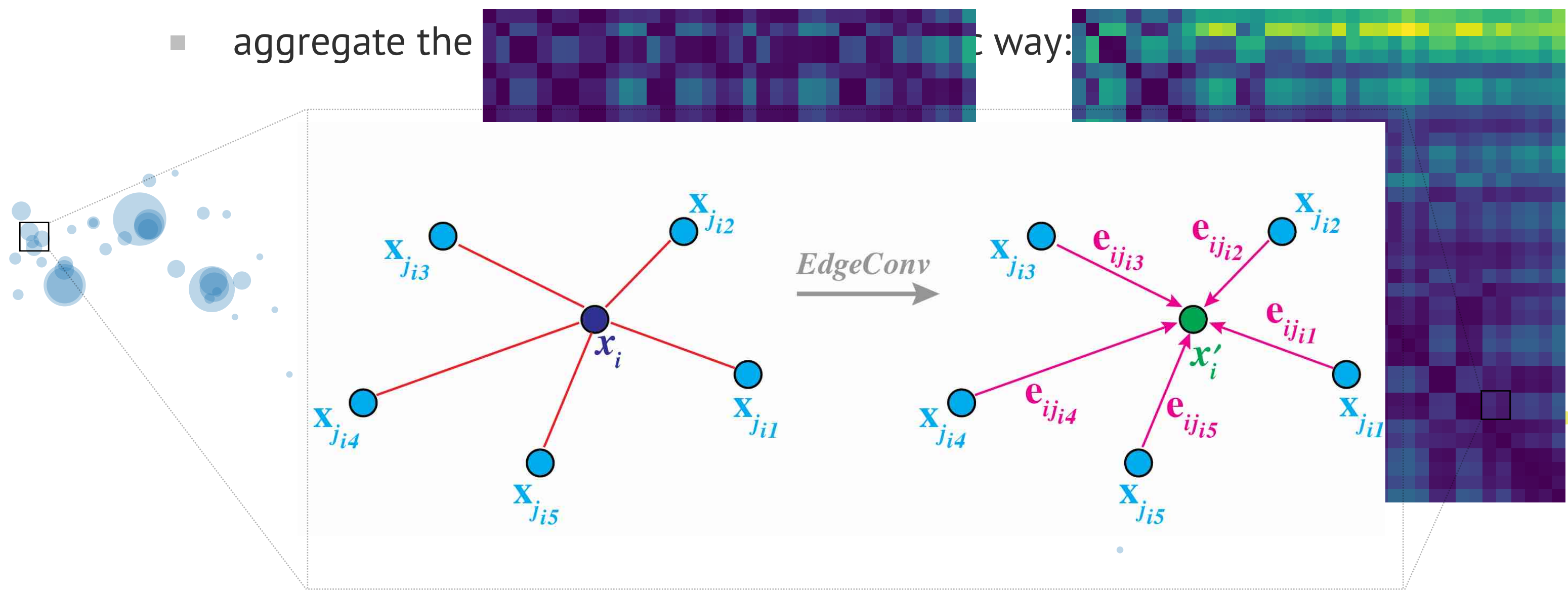
- HEP data as a point cloud

- each particle / detector cell is a point in the cloud
  - for each point: (spatial) coordinates + any additional properties (energy/momentum, detector response, ...)
- key feature: ***permutation symmetry*** – *algorithm output should not depend on the input ordering*

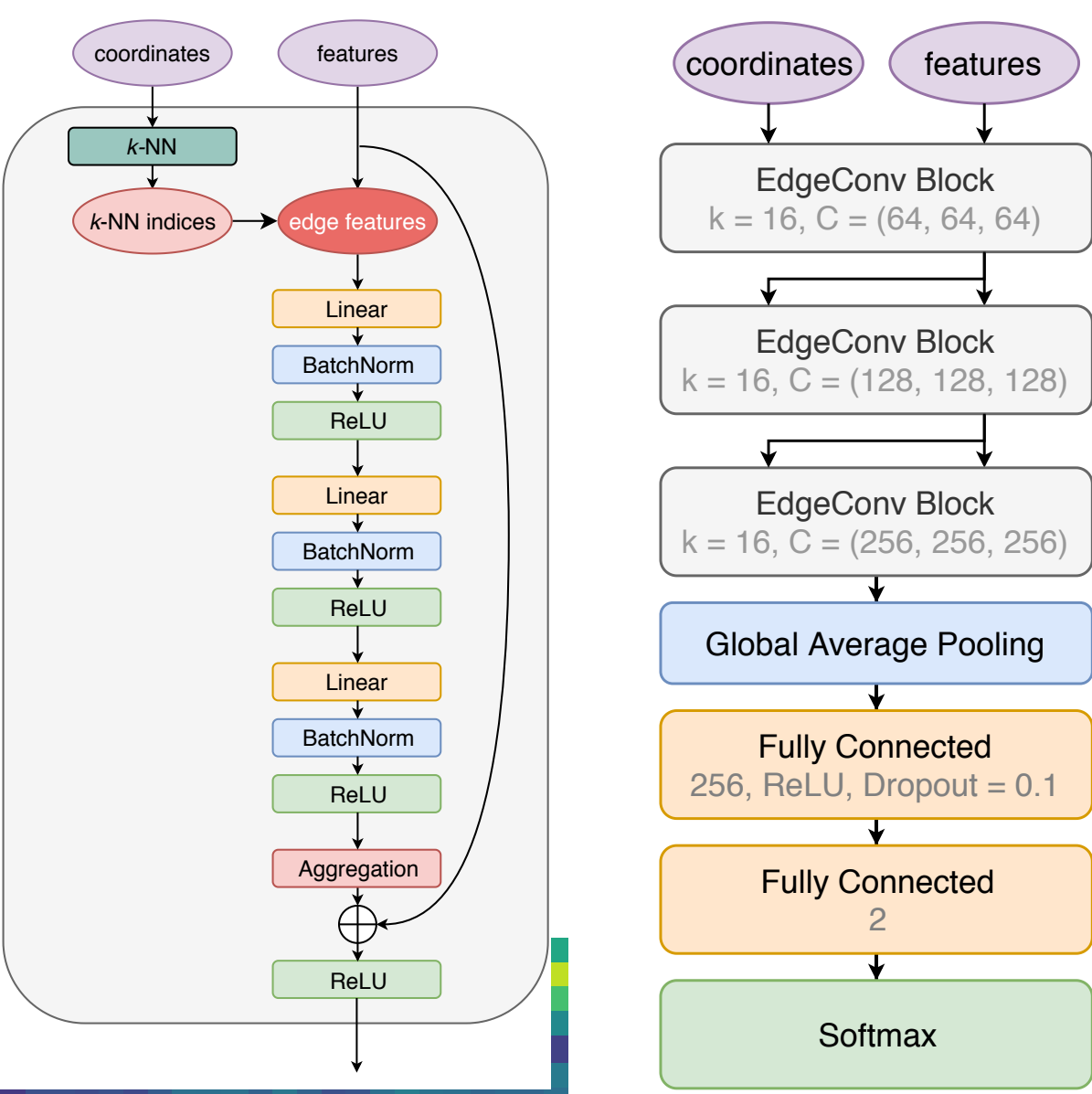


# PARTICLENET

- ParticleNet: jet tagging via particle clouds
  - treating a jet as an **unordered set of particles**, distributed in the  $\eta - \phi$  space
  - **graph neural network architecture**, adapted from DGCNN [arXiv:1801.07829]
    - treating a point cloud as a graph: each point is a vertex
      - for each point, a local patch is defined by finding its k-nearest neighbors
    - designing a permutation-invariant “convolution” function
      - define “edge feature” for each center-neighbor pair:  $e_{ij} = \text{MLP}(x_i, x_j)$
      - aggregate the



HQ and L. Gouskos  
*Phys.Rev.D 101 (2020) 5, 056019*



Performance on top quark tagging benchmark  
[SciPost Phys. 7, 014 (2019)]

|                  | $1/\epsilon_b$ at $\epsilon_s = 30\%$ |
|------------------|---------------------------------------|
| ResNeXt-50       | $1147 \pm 58$                         |
| P-CNN            | $759 \pm 24$                          |
| PFN              | $888 \pm 17$                          |
| ParticleNet-Lite | $1262 \pm 49$                         |
| ParticleNet      | <b><math>1615 \pm 93</math></b>       |

See also P.T. Komiske, E. M. Metodiev and J. Thaler, *JHEP 01 (2019) 121*



# PARTICLENET

HQ and L. Gouskos  
*Phys.Rev.D 101 (2020) 5, 056019*

- ParticleNet: jet tagging via particle clouds

- treating a jet as an **unordered**
- graph neural network** architecture
  - treating a point cloud as a graph
    - for each point, a local patch is defined
  - designing a permutation-invariant pooling
  - defining a “distance” between points
  - aggregating information



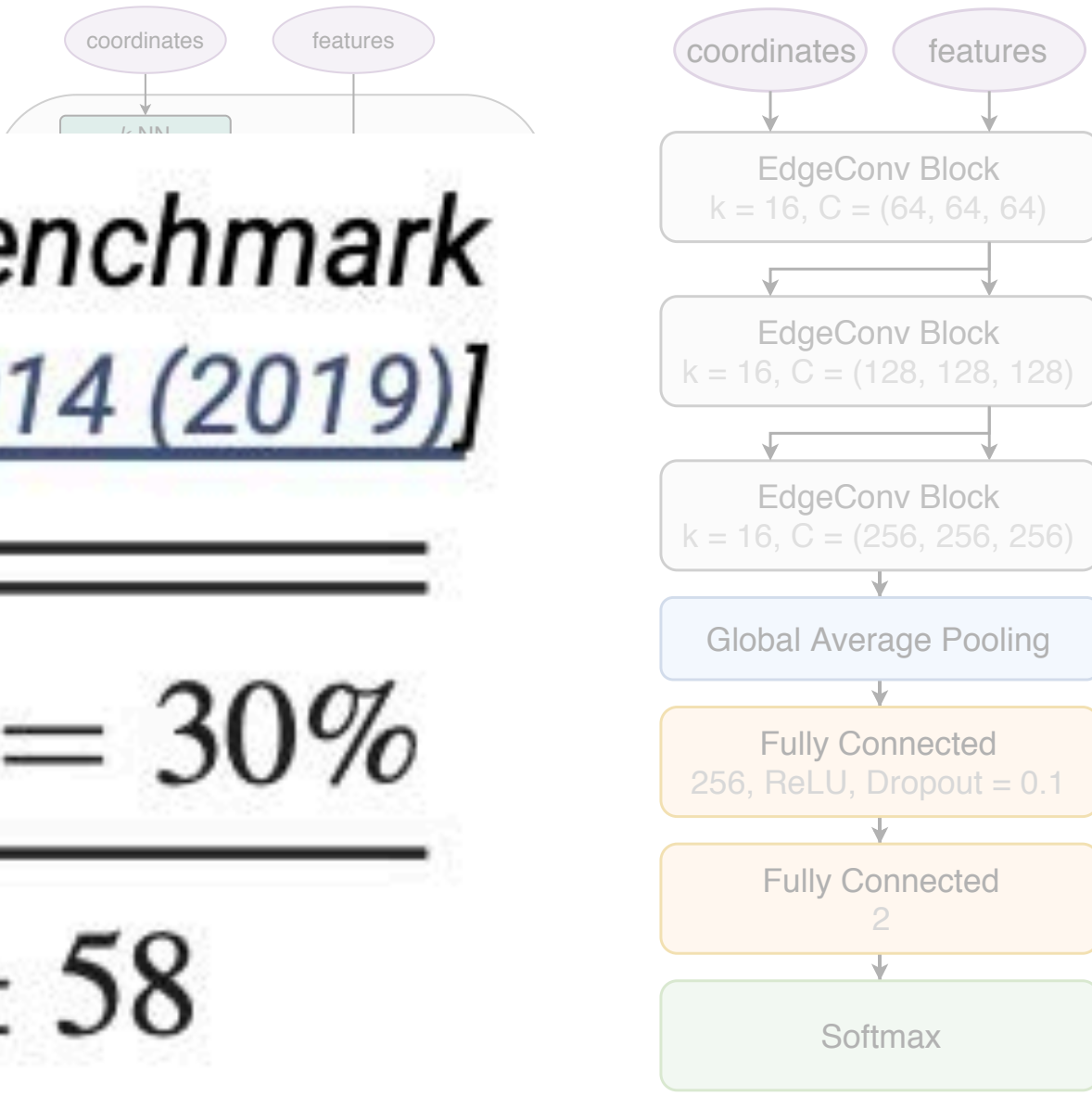
Image

Sequence

Point cloud

## Performance on top quark tagging benchmark *[SciPost Phys. 7, 014 (2019)]*

| $1/\epsilon_b$ at $\epsilon_s = 30\%$ |                                 |
|---------------------------------------|---------------------------------|
| ResNeXt-50                            | $1147 \pm 58$                   |
| P-CNN                                 | $759 \pm 24$                    |
| PFN                                   | $888 \pm 17$                    |
| ParticleNet-Lite                      | $1262 \pm 49$                   |
| <b>ParticleNet</b>                    | <b><math>1615 \pm 93</math></b> |



top quark tagging benchmark  
*[SciPost Phys. 7, 014 (2019)]*

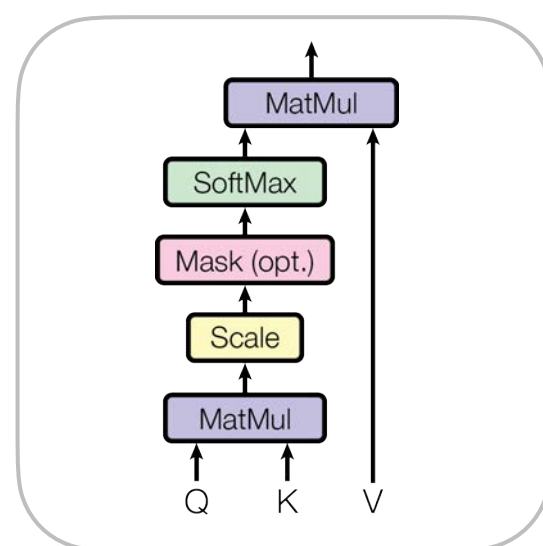
| $1/\epsilon_b$ at $\epsilon_s = 30\%$ |  |
|---------------------------------------|--|
| 1147 ± 58                             |  |
| 759 ± 24                              |  |
| 888 ± 17                              |  |
| 1262 ± 49                             |  |
| <b>1615 ± 93</b>                      |  |





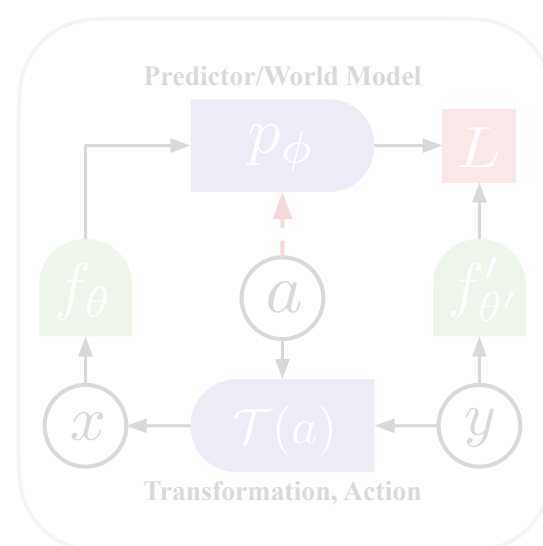
## Evolution of Jet Representations

*image, sequence, point cloud, ...*



## Attention Is All You Need?

*or how physics can help...*



## Towards a Foundation Model

*scaling, self-supervision, and more...*



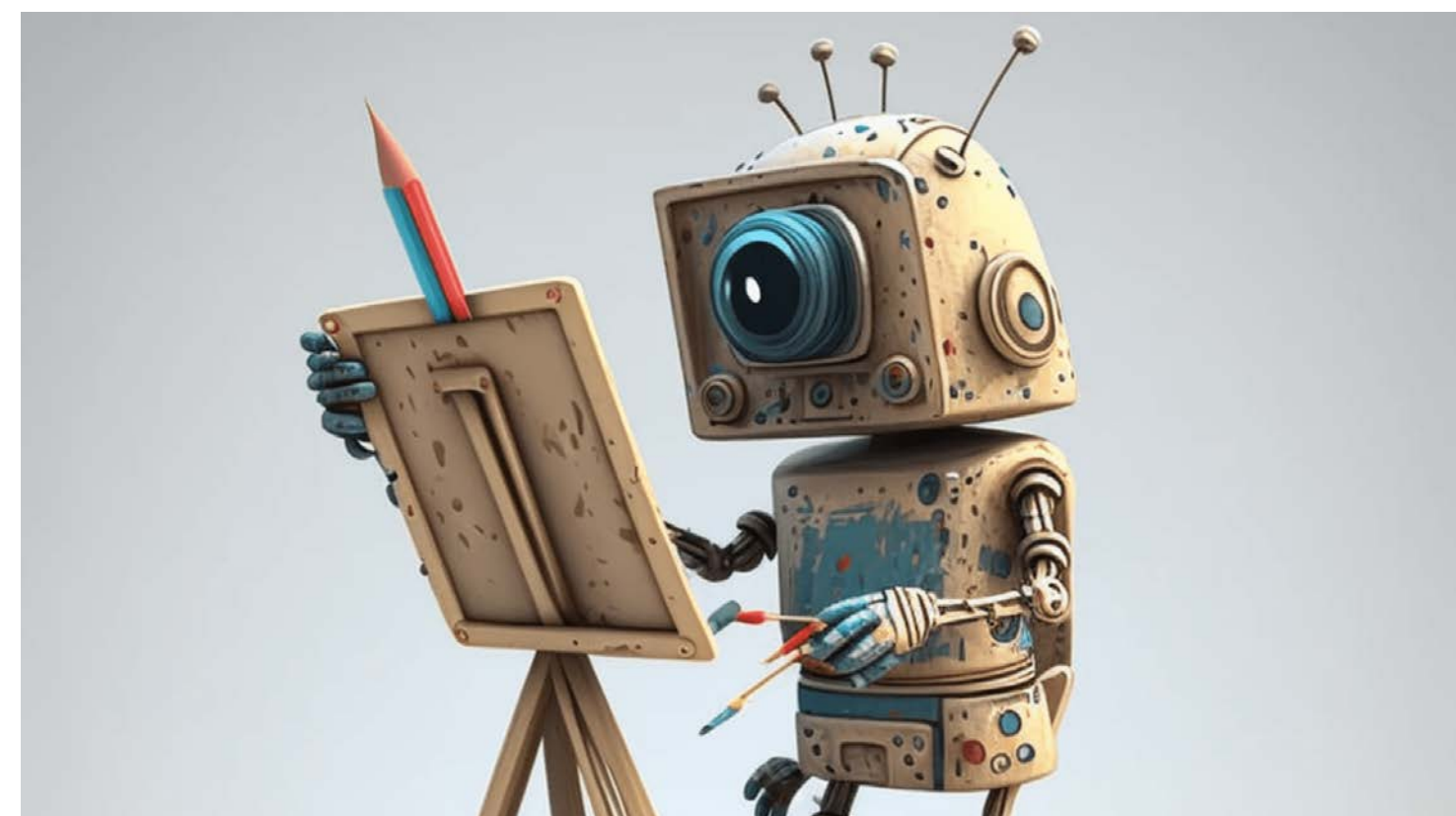
# ATTENTION IS ALL YOU NEED?

- **Transformers** are taking over every task

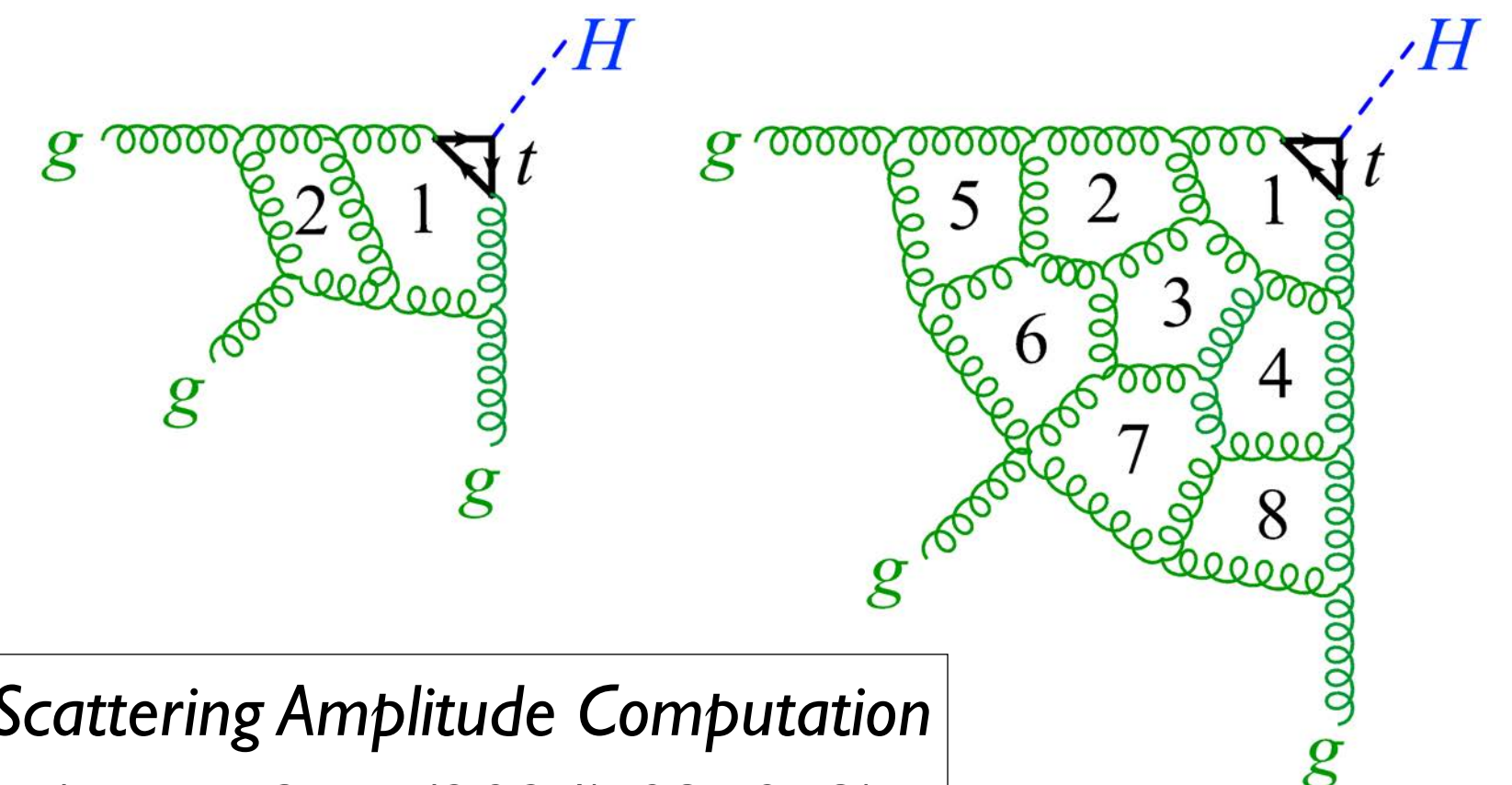
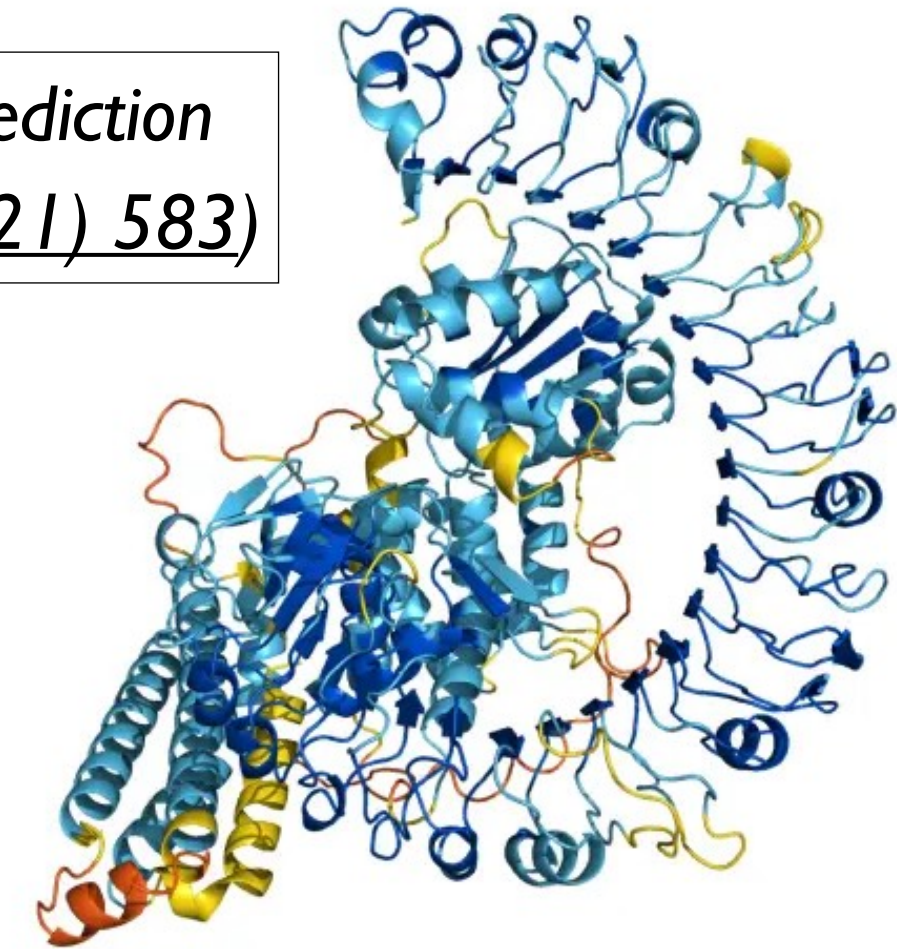
*Large Language Models*



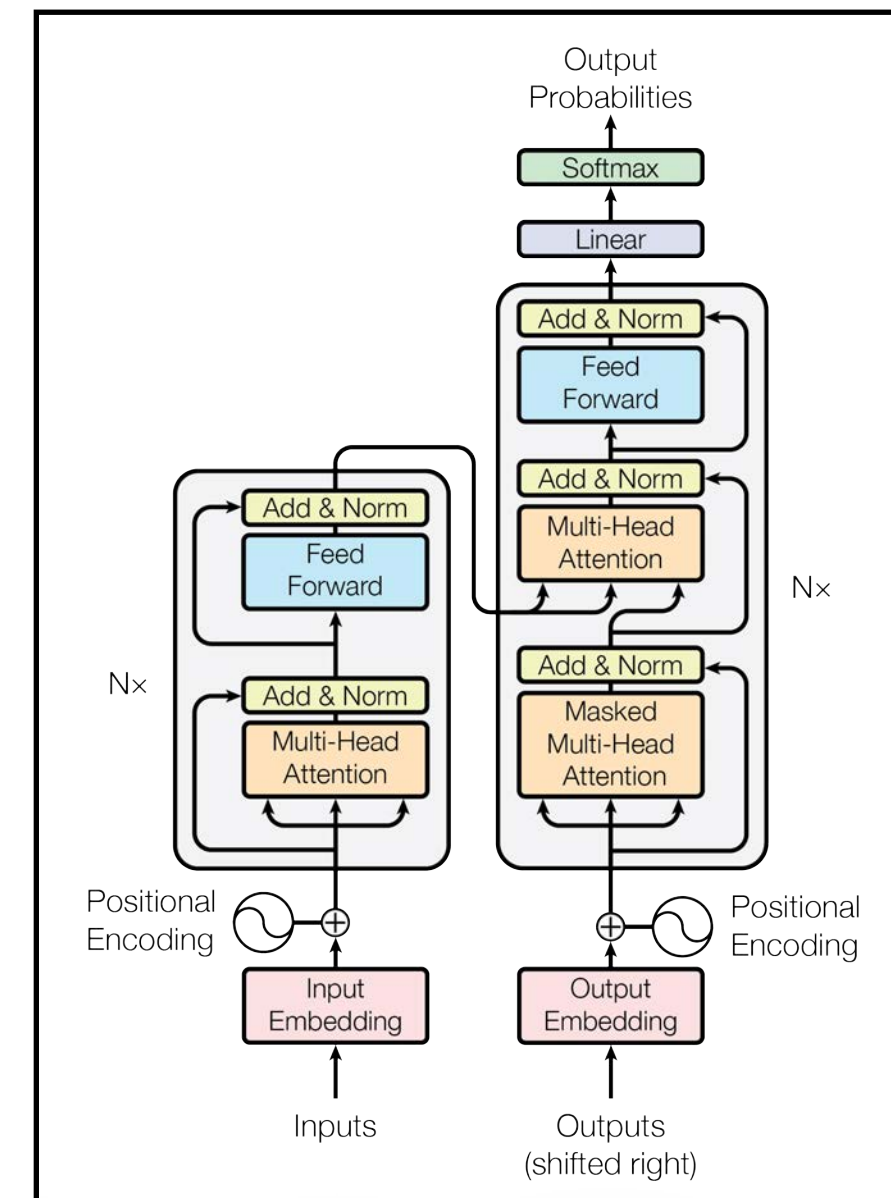
*Image/Video Generation*



*Protein Structure Prediction  
(e.g. Nature 596 (2021) 583)*



*Scattering Amplitude Computation  
(e.g. MLST 5 (2024) 035073)*





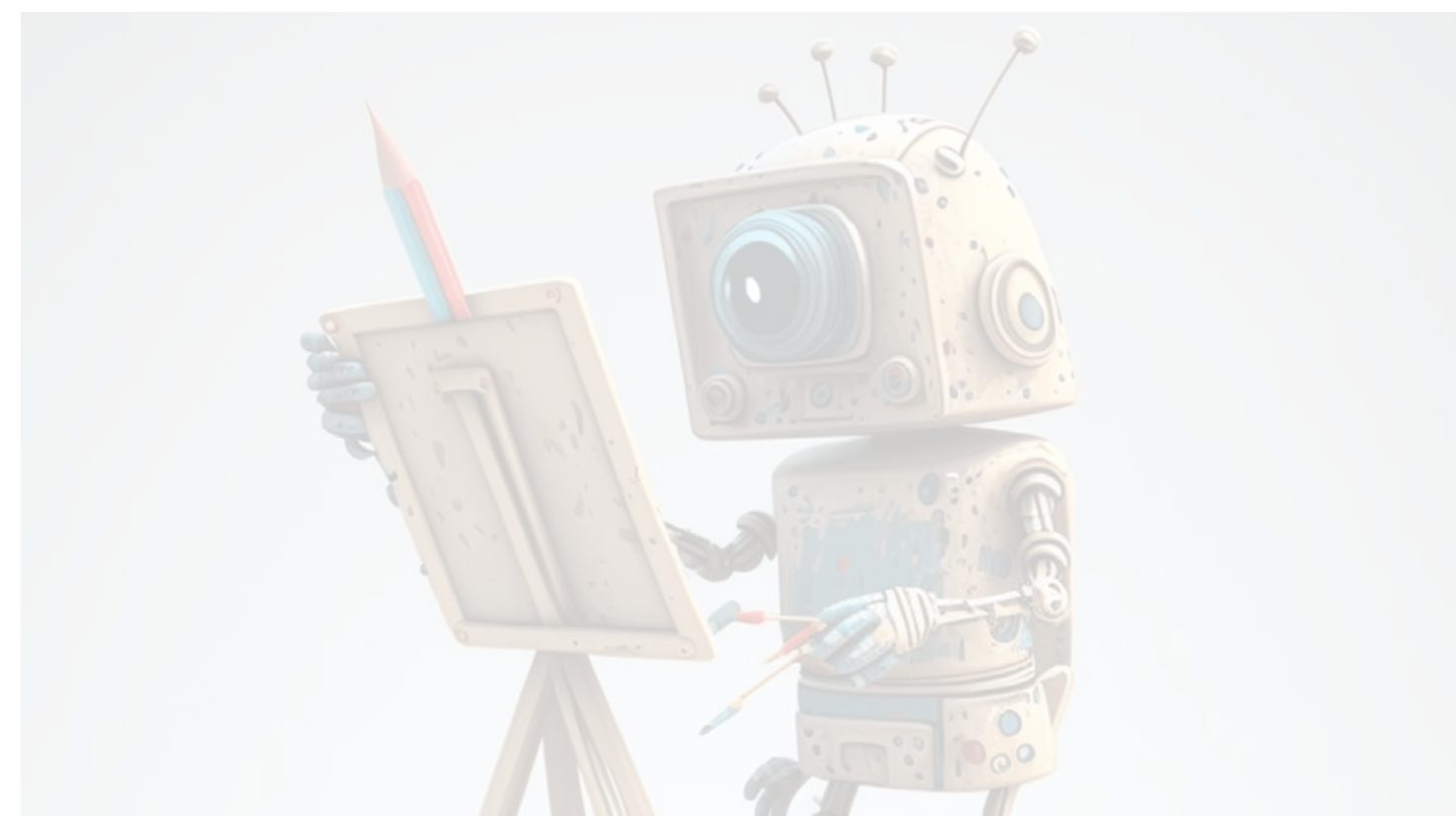
# ATTENTION IS ALL YOU NEED?

- Transformers are taking over every task

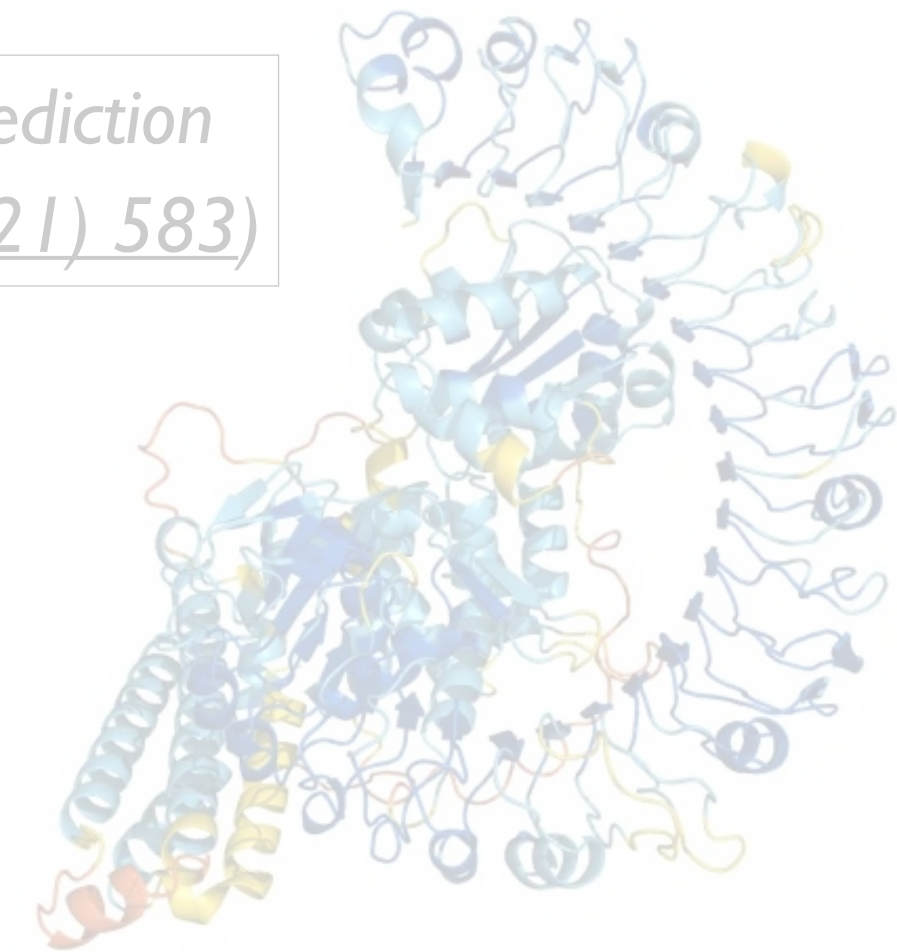
Large Language Models



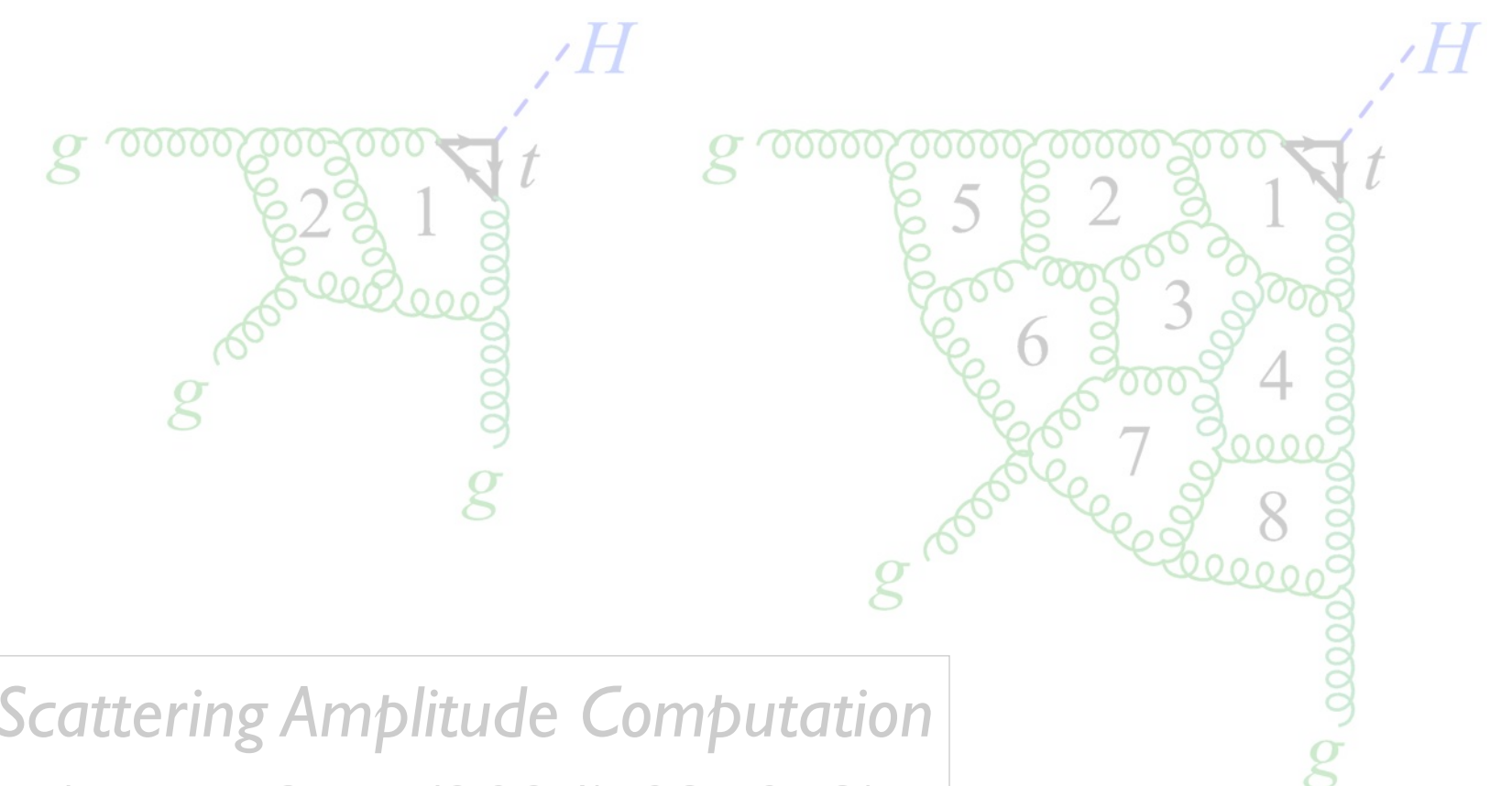
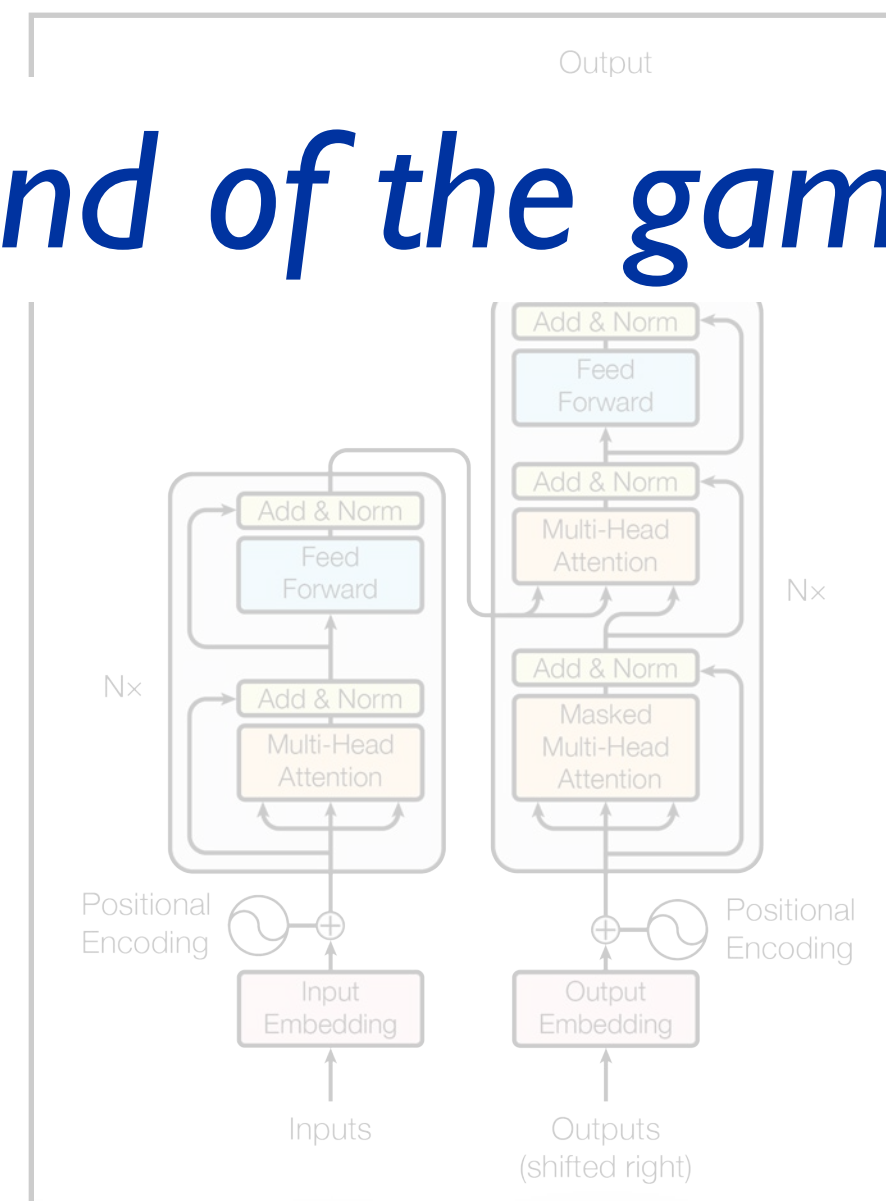
Image/Video Generation



Protein Structure Prediction  
(e.g. *Nature* 596 (2021) 583)



*End of the game?*



Scattering Amplitude Computation  
(e.g. *MLST* 5 (2024) 035073)



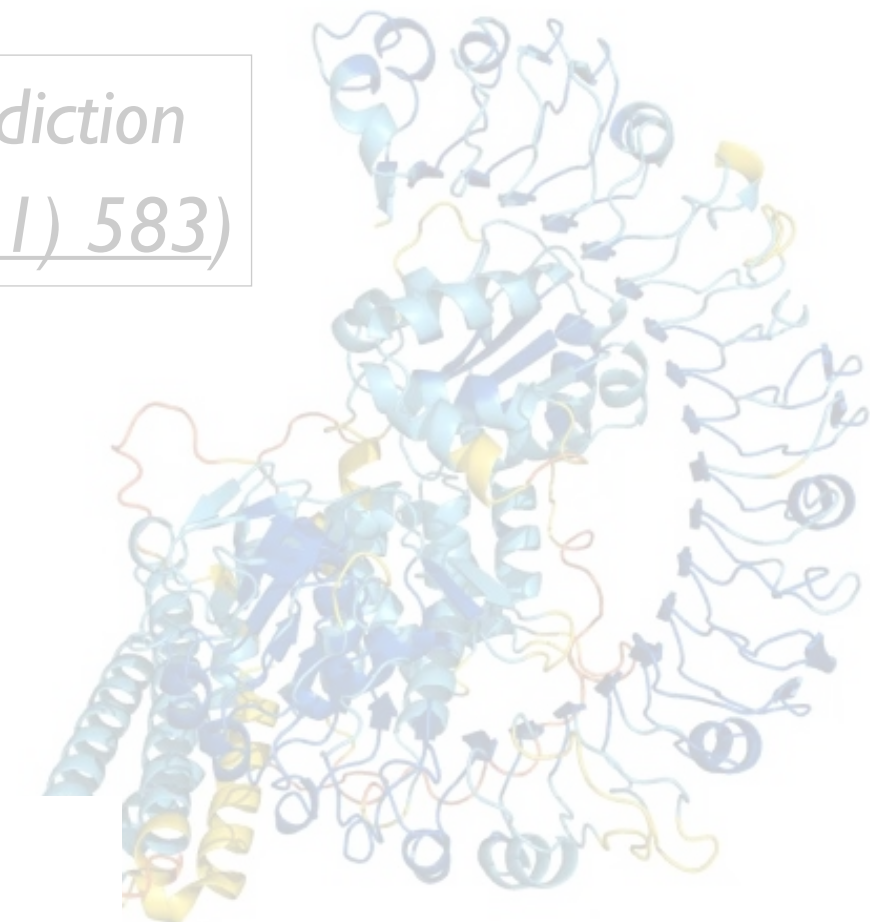
# ATTENTION IS ALL YOU NEED?

- Transformers are taking over every task

Large Language Models



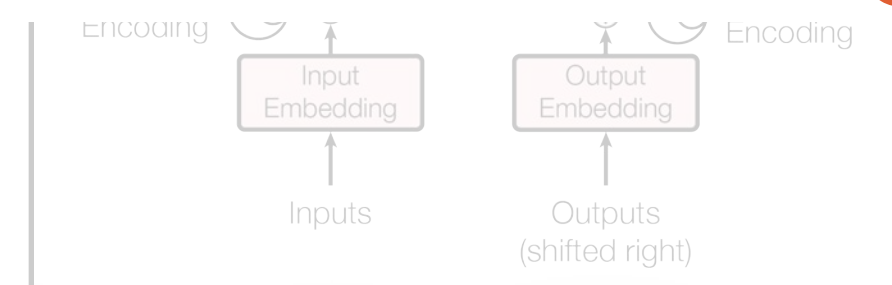
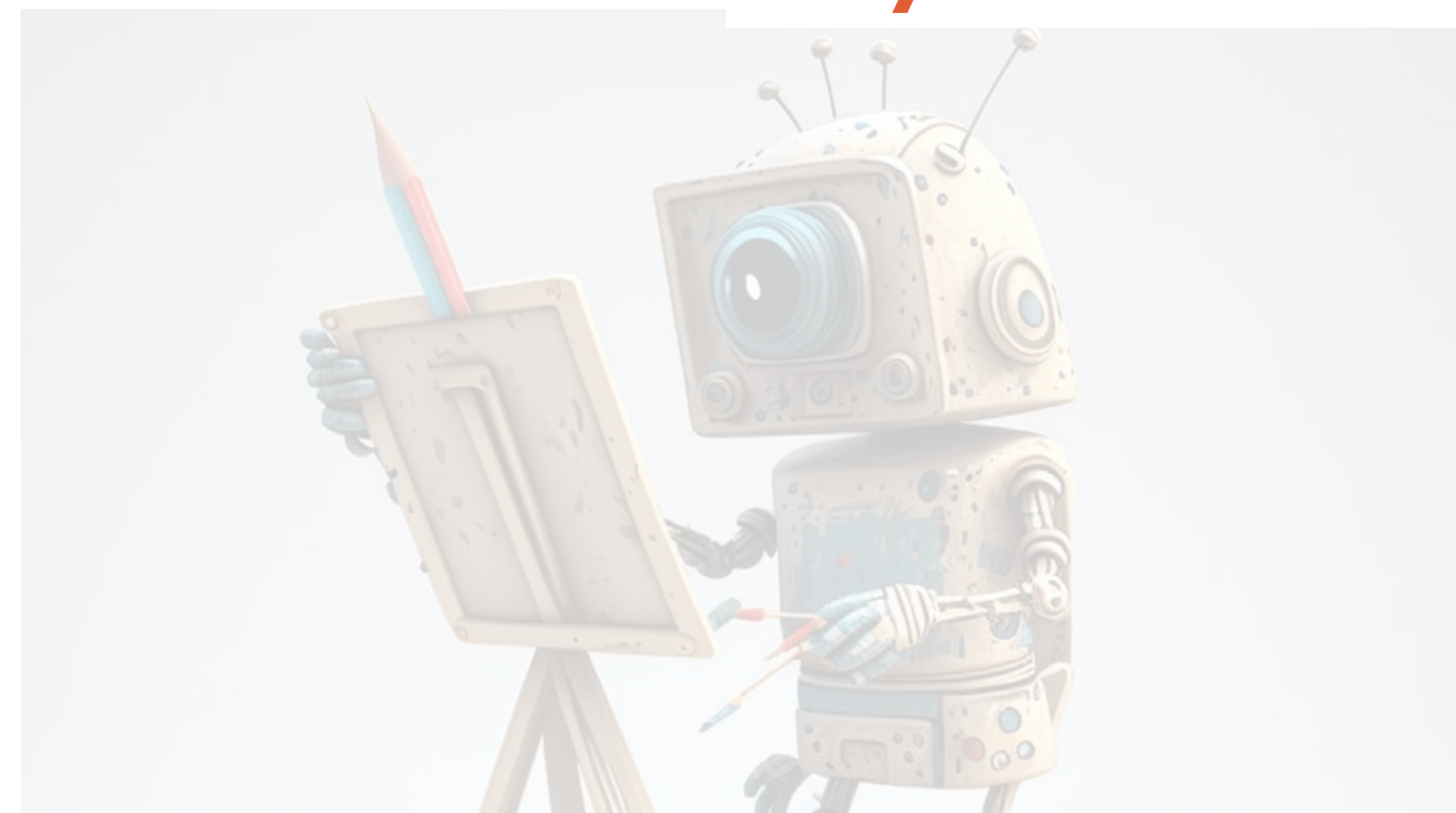
Protein Structure Prediction  
(e.g. *Nature* 596 (2021) 583)



*End of the game?*

*Not really...*

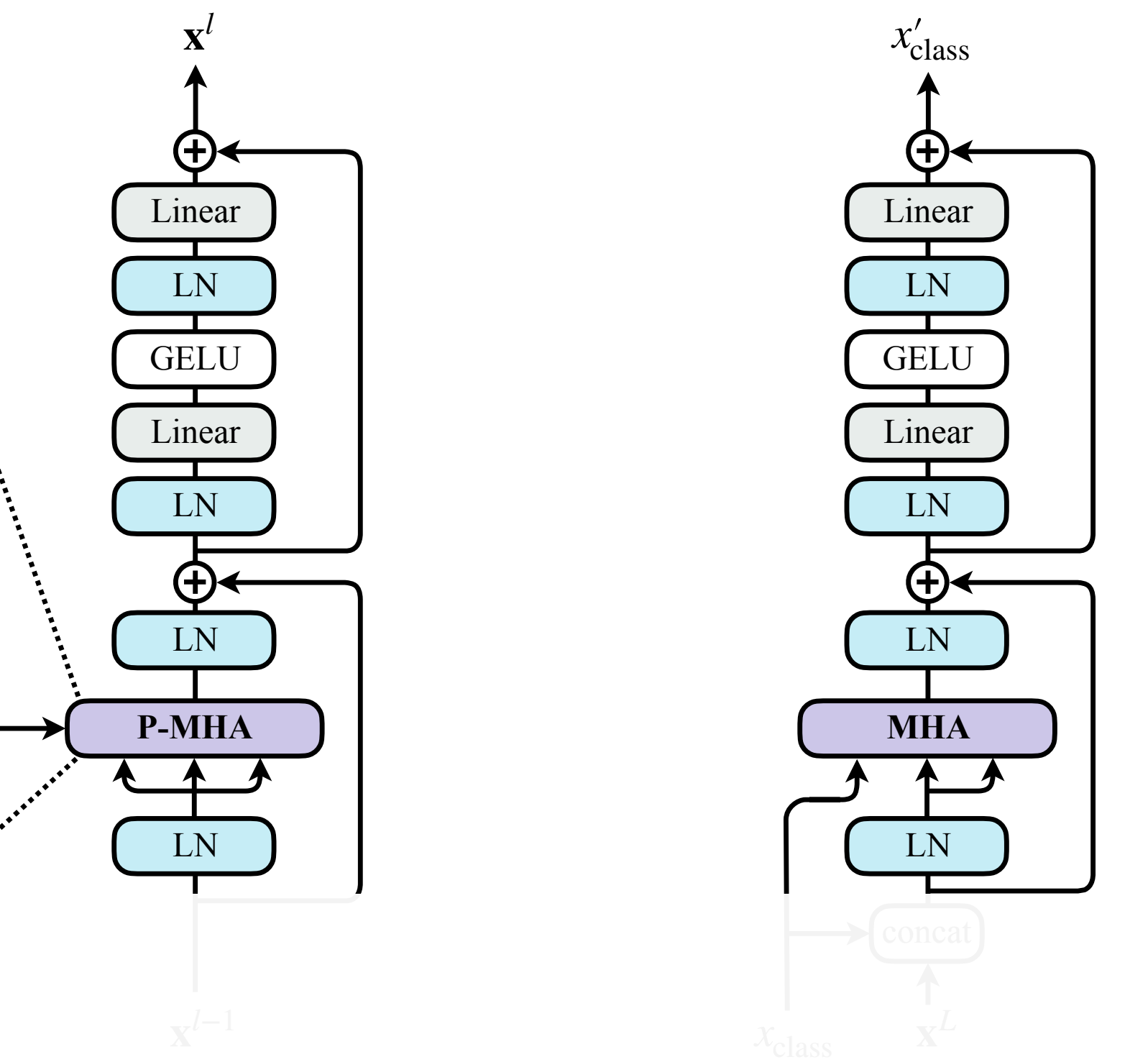
*Physics + network “co-design” is the key.*



Scattering Amplitude Computation  
(e.g. *MLST* 5 (2024) 035073)



(a) Particle Transformer



Block

(c) Class Attention Block

Physics-driven enhancement

$$\Delta = \sqrt{(y_a - y_b)^2 + (\phi_a - \phi_b)^2},$$
$$k_T = \min(p_{T,a}, p_{T,b}) \Delta,$$
$$z = \min(p_{T,a}, p_{T,b}) / (p_{T,a} + p_{T,b}),$$
$$m^2 = (E_a + E_b)^2 - \|\mathbf{p}_a + \mathbf{p}_b\|^2,$$

U

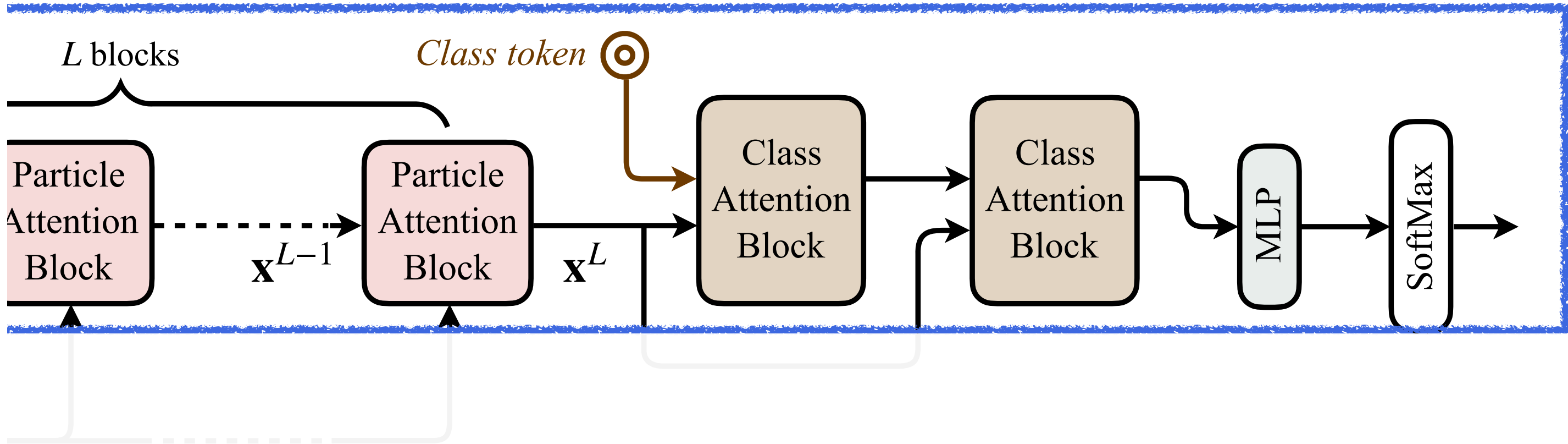
+

Mask

# RMER

mer model tailored for particle physics

HQ, C. Li, S. Qian,  
ICML 2022



x^l

x\_{class}^l

$$\text{P-MHA}(Q, K, V) = \text{SoftMax}(QK^T / \sqrt{d_k} + \text{U})V,$$

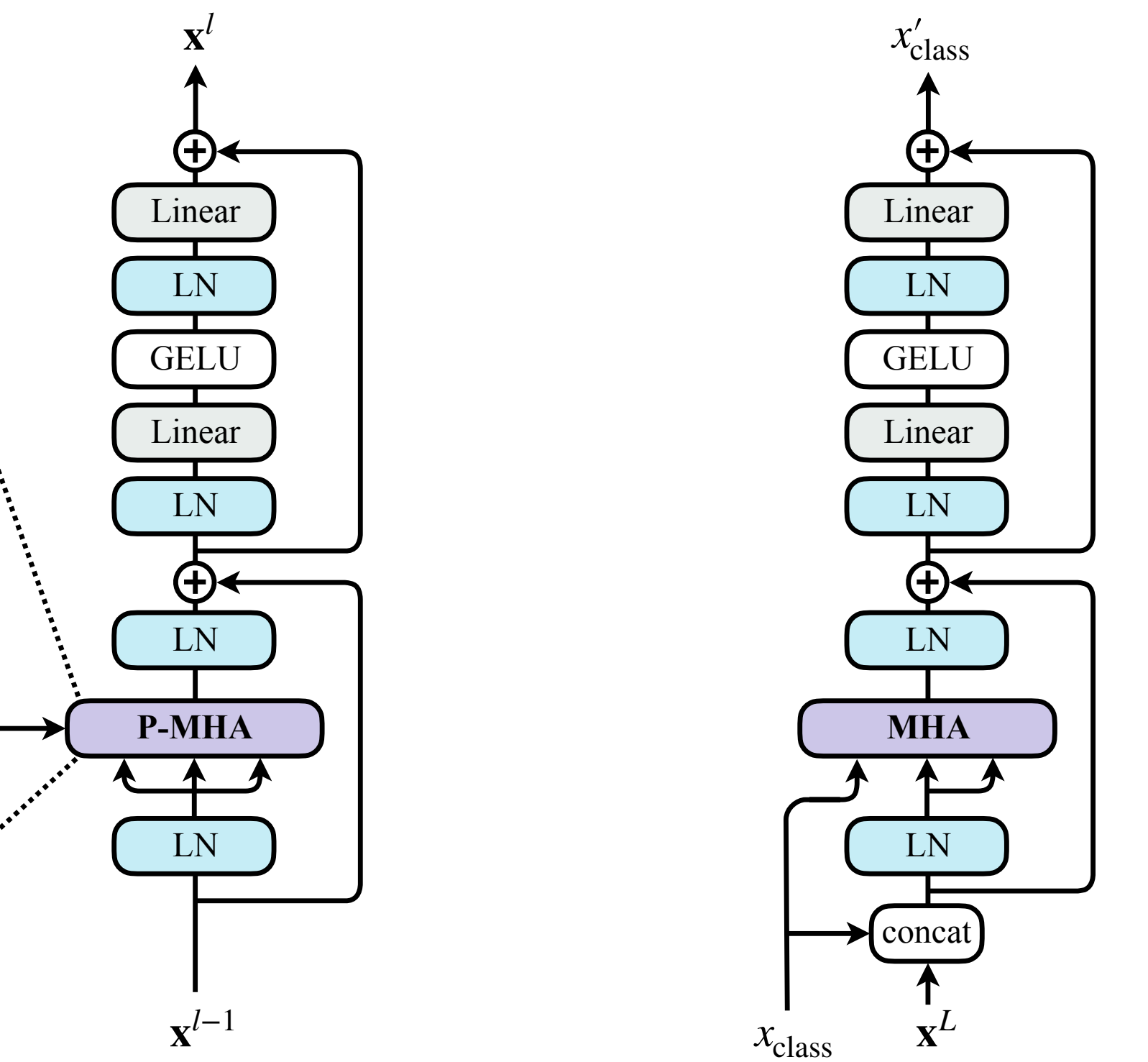
Linear  
LN  
GELU  
Linear

Injection of (physics-inspired) pairwise features to “bias” the dot-product self-attention

Linear  
LN  
GELU  
Linear



(a) Particle Transformer



Block (c) Class Attention Block

Physics-driven enhancement

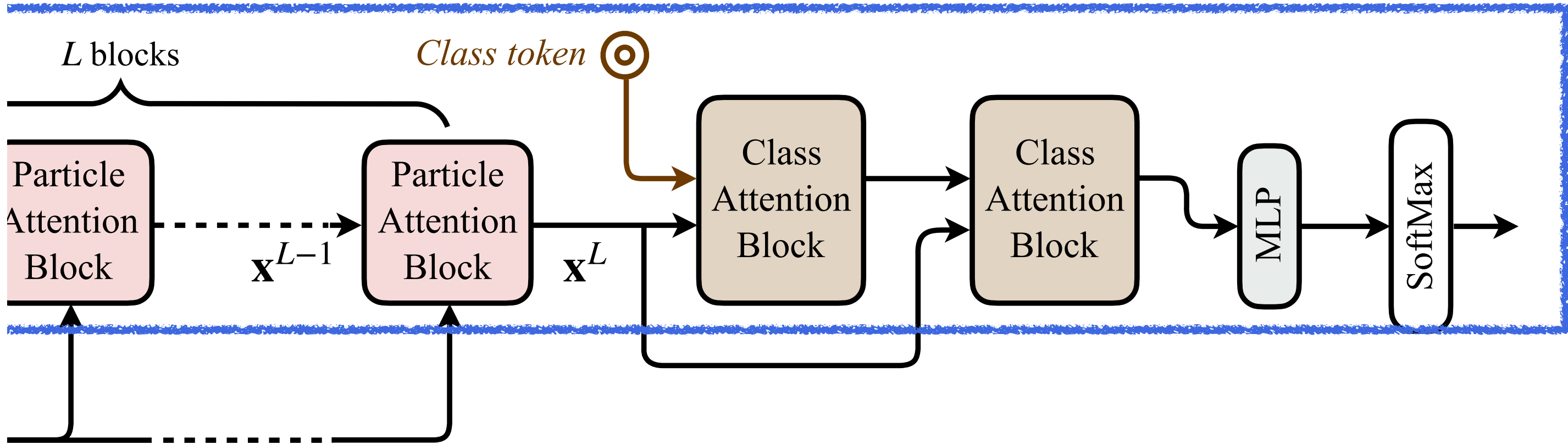
$$\Delta = \sqrt{(y_a - y_b)^2 + (\phi_a - \phi_b)^2},$$
$$k_T = \min(p_{T,a}, p_{T,b}) \Delta,$$
$$z = \min(p_{T,a}, p_{T,b}) / (p_{T,a} + p_{T,b}),$$
$$m^2 = (E_a + E_b)^2 - \|\mathbf{p}_a + \mathbf{p}_b\|^2,$$

U ⊕ Mask

# RMER

mer model tailored for particle physics

HQ, C. Li, S. Qian,  
ICML 2022



$$\text{P-MHA}(Q, K, V) = \text{SoftMax}(QK^T / \sqrt{d_k} + \mathbf{U})V,$$

Linear LN GELU Linear LN

Injection of (physics-inspired) pairwise features to “bias” the dot-product self-attention



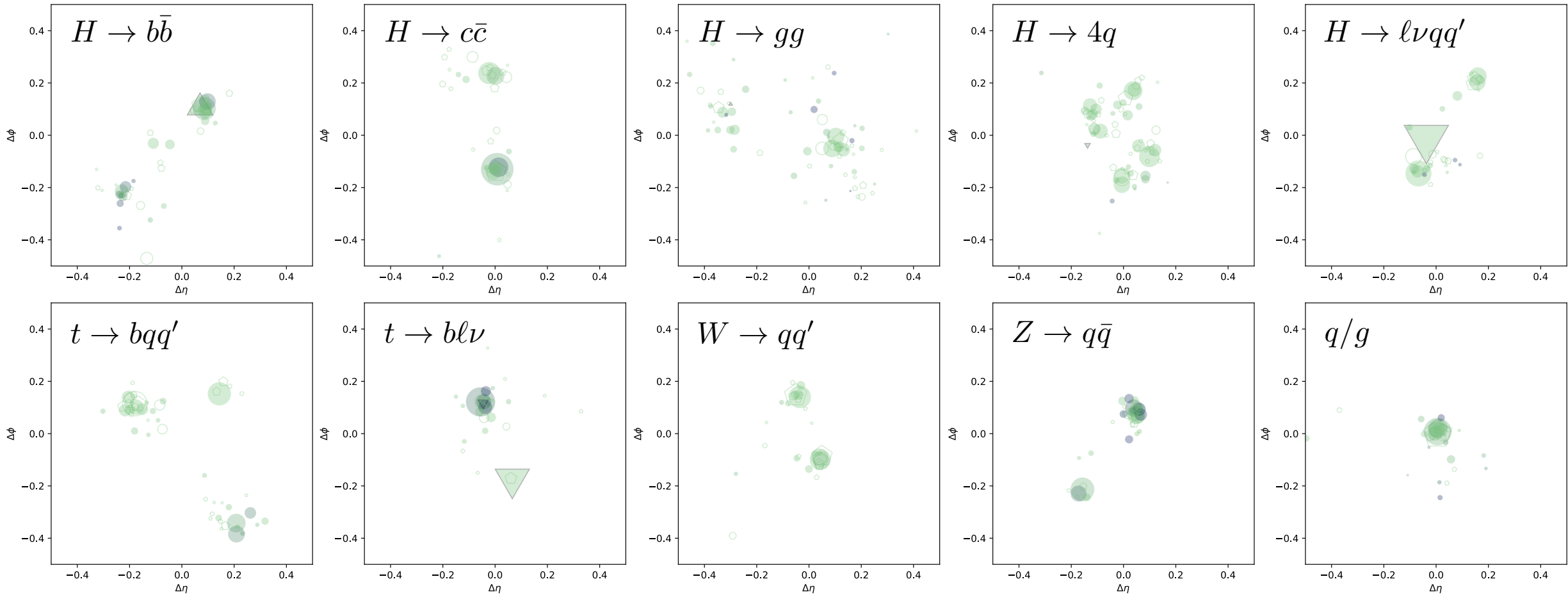
# PARTICLE TRANSFORMER: PERFORMANCE

HQ, C. Li, S. Qian,  
ICML 2022

|             | All classes  |               | $H \rightarrow b\bar{b}$ | $H \rightarrow c\bar{c}$ | $H \rightarrow gg$ | $H \rightarrow 4q$ | $H \rightarrow \ell\nu qq'$ | $t \rightarrow bqq'$ | $t \rightarrow b\ell\nu$ | $W \rightarrow qq'$ | $Z \rightarrow q\bar{q}$ |
|-------------|--------------|---------------|--------------------------|--------------------------|--------------------|--------------------|-----------------------------|----------------------|--------------------------|---------------------|--------------------------|
|             | Accuracy     | AUC           | Rej <sub>50%</sub>       | Rej <sub>50%</sub>       | Rej <sub>50%</sub> | Rej <sub>50%</sub> | Rej <sub>99%</sub>          | Rej <sub>50%</sub>   | Rej <sub>99.5%</sub>     | Rej <sub>50%</sub>  | Rej <sub>50%</sub>       |
| PFN         | 0.772        | 0.9714        | 2924                     | 841                      | 75                 | 198                | 265                         | 797                  | 721                      | 189                 | 159                      |
| P-CNN       | 0.809        | 0.9789        | 4890                     | 1276                     | 88                 | 474                | 947                         | 2907                 | 2304                     | 241                 | 204                      |
| ParticleNet | 0.844        | 0.9849        | 7634                     | 2475                     | 104                | 954                | 3339                        | 10526                | 11173                    | 347                 | 283                      |
| <b>ParT</b> | <b>0.861</b> | <b>0.9877</b> | <b>10638</b>             | <b>4149</b>              | <b>123</b>         | <b>1864</b>        | <b>5479</b>                 | <b>32787</b>         | <b>15873</b>             | <b>543</b>          | <b>402</b>               |

- Particle Transformer (ParT):
  - significant gains compared to ParticleNet
    - 1.7% increase in accuracy
    - up to 3x increase in background rejection (Rej<sub>x%</sub>)

JETCLASS dataset (100M jets)





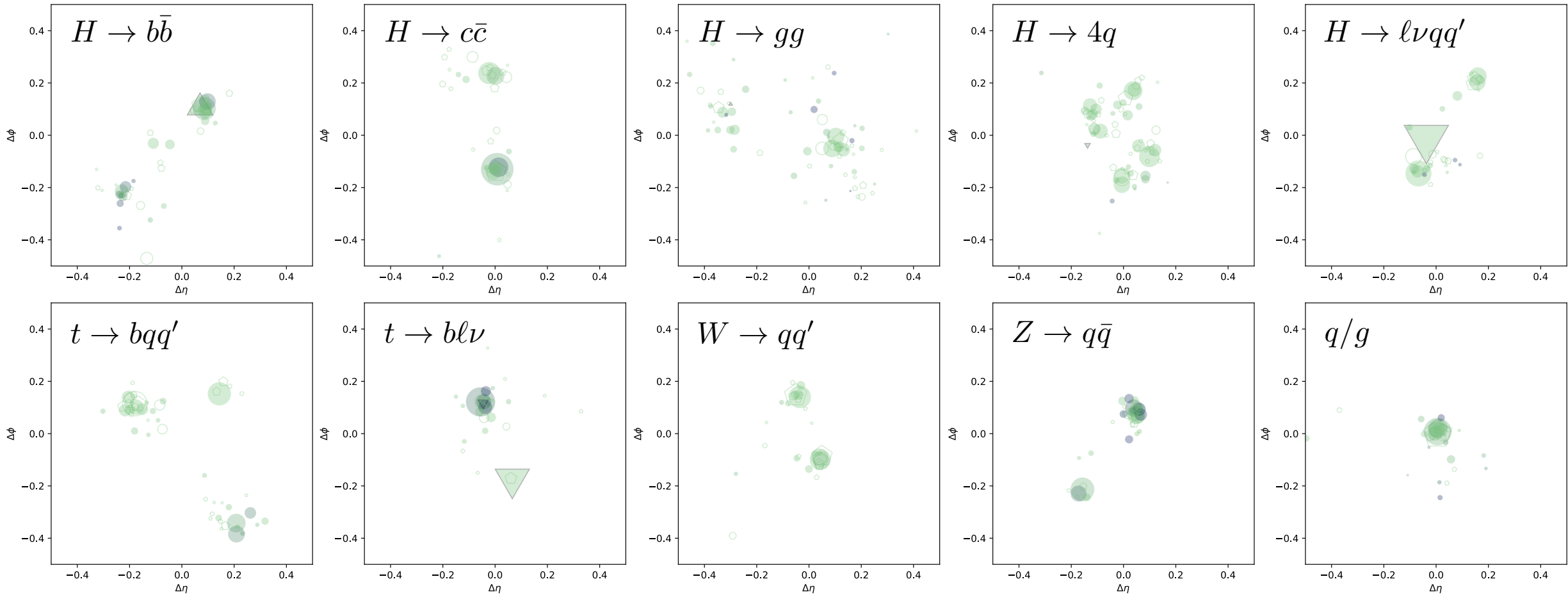
# PARTICLE TRANSFORMER: PERFORMANCE

HQ, C. Li, S. Qian,  
ICML 2022

|              | All classes  |               | $H \rightarrow b\bar{b}$ | $H \rightarrow c\bar{c}$ | $H \rightarrow gg$ | $H \rightarrow 4q$ | $H \rightarrow \ell\nu qq'$ | $t \rightarrow bqq'$ | $t \rightarrow b\ell\nu$ | $W \rightarrow qq'$ | $Z \rightarrow q\bar{q}$ |
|--------------|--------------|---------------|--------------------------|--------------------------|--------------------|--------------------|-----------------------------|----------------------|--------------------------|---------------------|--------------------------|
|              | Accuracy     | AUC           | Rej <sub>50%</sub>       | Rej <sub>50%</sub>       | Rej <sub>50%</sub> | Rej <sub>50%</sub> | Rej <sub>99%</sub>          | Rej <sub>50%</sub>   | Rej <sub>99.5%</sub>     | Rej <sub>50%</sub>  | Rej <sub>50%</sub>       |
| PFN          | 0.772        | 0.9714        | 2924                     | 841                      | 75                 | 198                | 265                         | 797                  | 721                      | 189                 | 159                      |
| P-CNN        | 0.809        | 0.9789        | 4890                     | 1276                     | 88                 | 474                | 947                         | 2907                 | 2304                     | 241                 | 204                      |
| ParticleNet  | 0.844        | 0.9849        | 7634                     | 2475                     | 104                | 954                | 3339                        | 10526                | 11173                    | 347                 | 283                      |
| <b>ParT</b>  | <b>0.861</b> | <b>0.9877</b> | <b>10638</b>             | <b>4149</b>              | <b>123</b>         | <b>1864</b>        | <b>5479</b>                 | <b>32787</b>         | <b>15873</b>             | <b>543</b>          | <b>402</b>               |
| ParT (plain) | 0.849        | 0.9859        | 9569                     | 2911                     | 112                | 1185               | 3868                        | 17699                | 12987                    | 384                 | 311                      |

- Particle Transformer (ParT):
  - significant gains compared to ParticleNet
    - 1.7% increase in accuracy
    - up to 3x increase in background rejection (Rej<sub>x%</sub>)
- ParT (plain): Transformer w/o pairwise features
  - 1.2% drop in accuracy compared to full ParT
  - Physics-driven modification plays a key role!

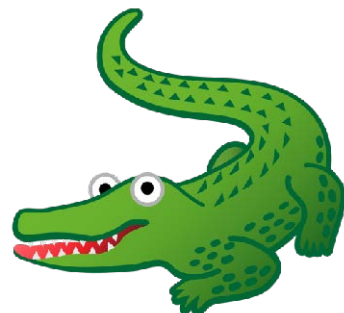
*JETCLASS* dataset (100M jets)



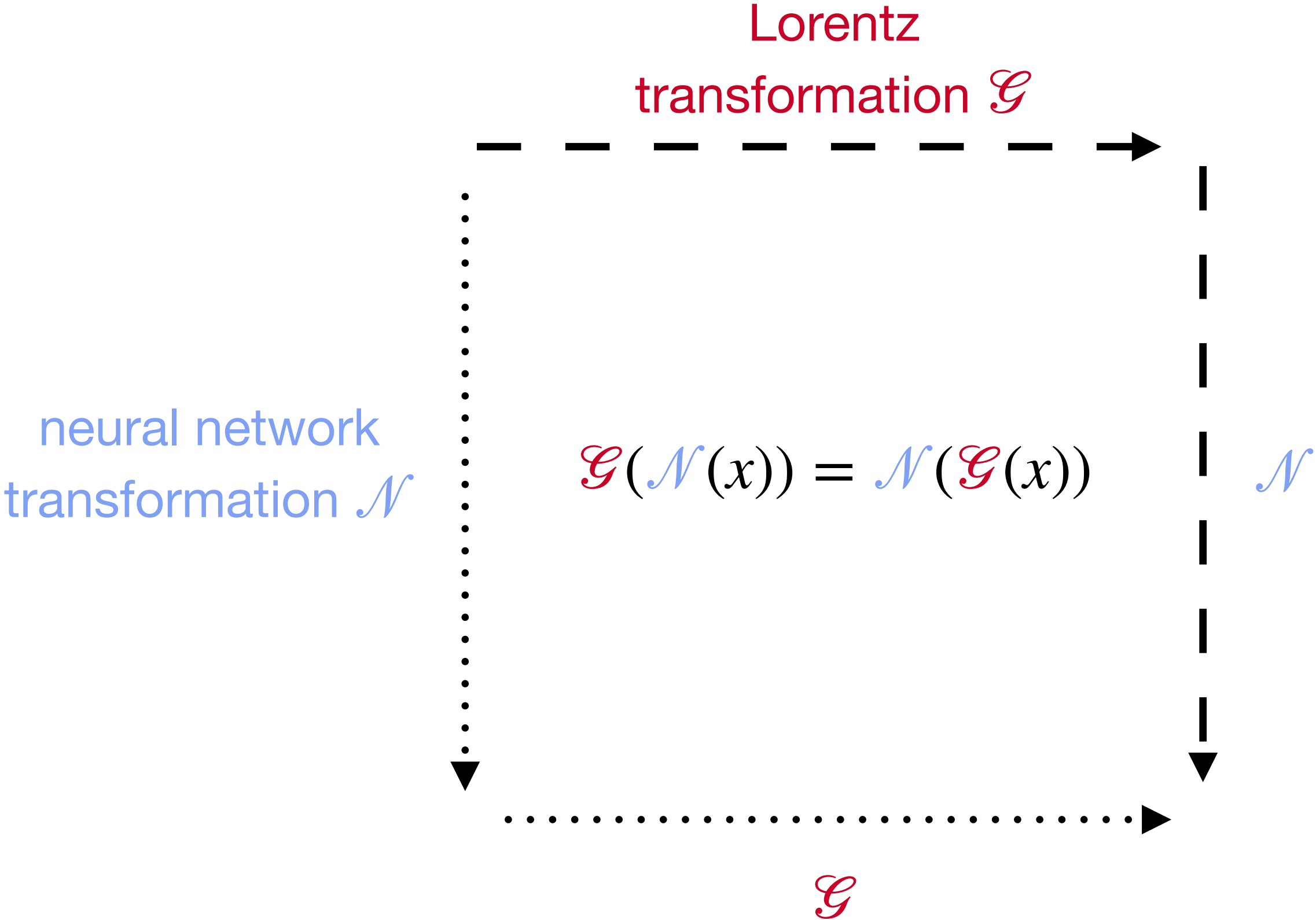


# SYMMETRY-AWARE ARCHITECTURES

- One step further: symmetry-aware architectures
  - one important symmetry for particle physics: Lorentz symmetry
- Lorentz-equivariant architectures
  - equivariance (covariance) – preserve the symmetry by construction



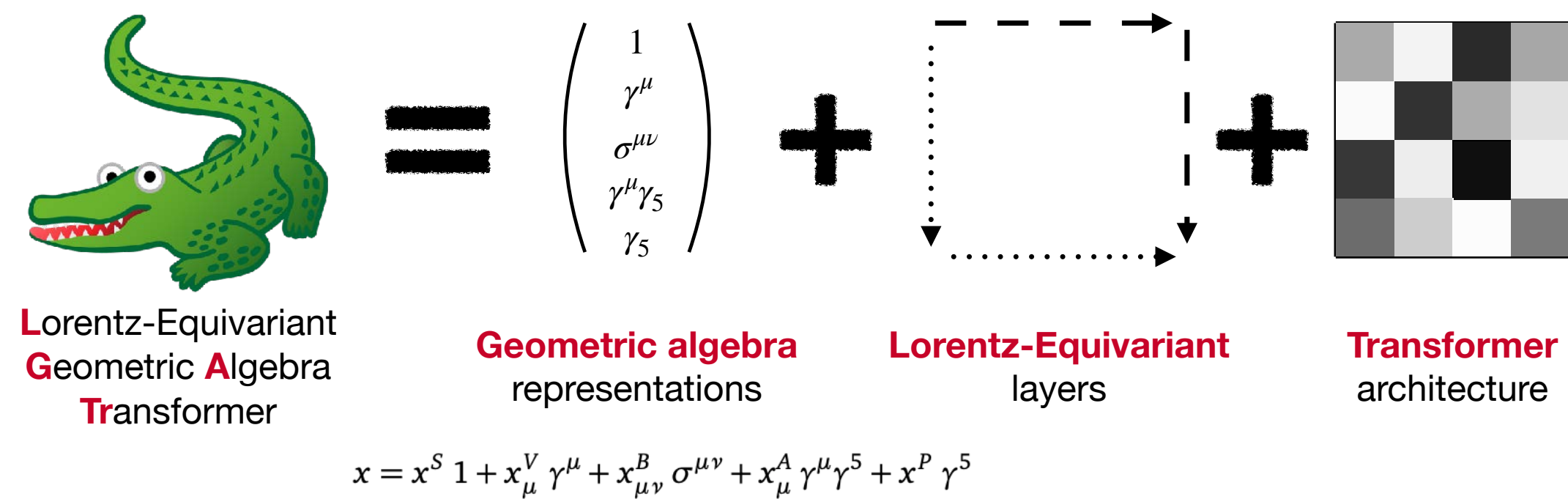
$$= \begin{pmatrix} 1 \\ \gamma^\mu \\ \sigma^{\mu\nu} \\ \gamma^\mu \gamma_5 \\ \gamma_5 \end{pmatrix} + \begin{array}{|c|} \hline \begin{array}{c} \text{---} \rightarrow \\ \vdots \\ \downarrow \text{---} \\ \vdots \\ \leftarrow \text{---} \\ \vdots \\ \vdots \end{array} \\ \hline \end{array} + \begin{array}{|c|} \hline \begin{array}{c} \text{---} \rightarrow \\ \vdots \\ \downarrow \text{---} \\ \vdots \\ \leftarrow \text{---} \\ \vdots \\ \vdots \end{array} \\ \hline \end{array}$$





# L-GATr & LLoCa

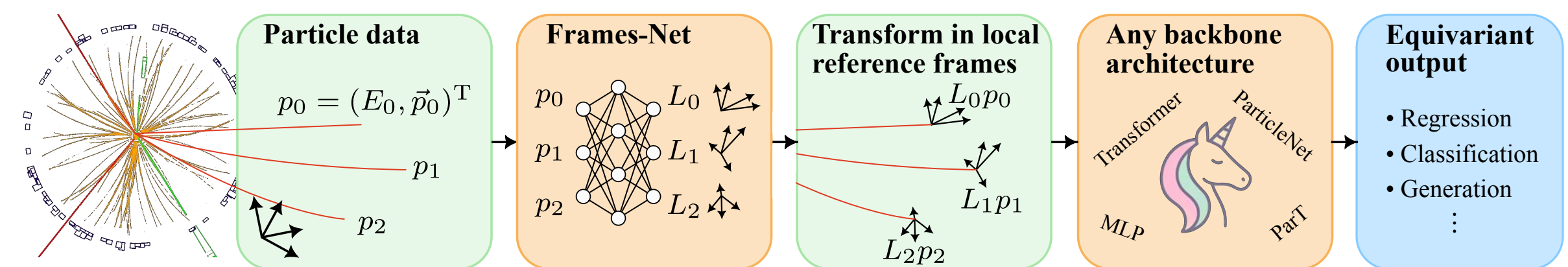
## L-GATr: Lorentz-Equivariant Geometric Algebra Transformer



*Specialized Lorentz-equivariant linear/attention/norm layers based on geometric algebra representations.*

J. Brehmer, V. Bresó, P. Haan, T. Plehn, HQ, J. Spinner and J. Thaler,  
[arXiv: 2411.00446](https://arxiv.org/abs/2411.00446)

## LLoCa: Lorentz Local Canonicalization



*Learned local reference frames to transform the inputs into a (learned) canonical frame and then the outputs back. Can be applied to any architecture (GNN, transformer).*

L. Favaro, G. Gerhartz, F.A. Hamprecht, P. Lippmann, S. Pitz, T. Plehn, HQ and J. Spinner, [arXiv: 2508.14898](https://arxiv.org/abs/2508.14898)



# L-GATr & LLoCA PERFORMANCE

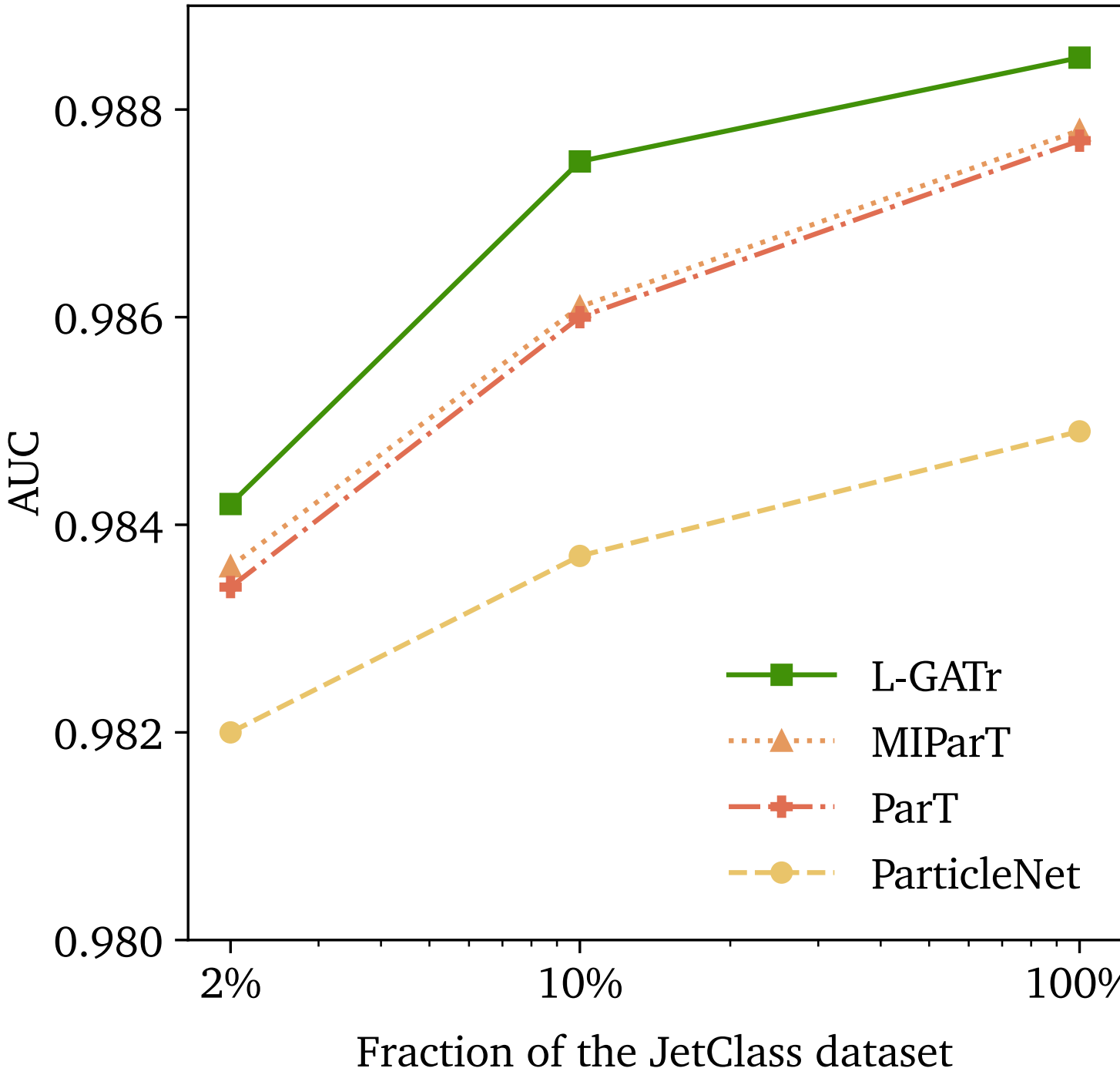
Performance on JetClass (100M)

| Network                 | Accuracy | AUC    | Time | FLOPs | Memory | Parameters |
|-------------------------|----------|--------|------|-------|--------|------------|
| MIParT-L [26]           | 0.861    | 0.9878 | 43h  | 225M  | 53.6G  | 2380k      |
| LorentzNet* [11]        | 0.847    | 0.9856 | 64h  | 676M  | 20.5G  | 223k       |
| PELICAN-lite* [17]      | 0.851    | 0.9862 | 97h  | 1370M | 27.4G  | 244k       |
| ParticleNet [9]         | 0.844    | 0.9849 | 25h  | 413M  | 16.5G  | 366k       |
| LLoCa-ParticleNet* [17] | 0.848    | 0.9857 | 43h  | 517M  | 23.5G  | 385k       |
| ParT [18]               | 0.861    | 0.9877 | 33h  | 211M  | 13.3G  | 2141k      |
| LLoCa-ParT* [17]        | 0.864    | 0.9882 | 66h  | 315M  | 19.9G  | 2160k      |
| Transformer [17]        | 0.855    | 0.9867 | 15h  | 210M  | 2.3G   | 1979k      |
| LLoCa-Transformer* [17] | 0.864    | 0.9882 | 28h  | 219M  | 4.1G   | 1980k      |
| L-GATr* [16]            | 0.866    | 0.9885 | 166h | 2060M | 19.0G  | 1079k      |
| L-GATr-slim*            | 0.866    | 0.9885 | 27h  | 329M  | 8.1G   | 2031k      |

\* *L-GATr-slim*: a more efficient version of L-GATr implemented using only Lorentz scalar and vectors.

A. Petitjean, T. Plehn, J. Spinner and U. Köthe, [arXiv: 2512.17011](#)

Scaling w/ training set size

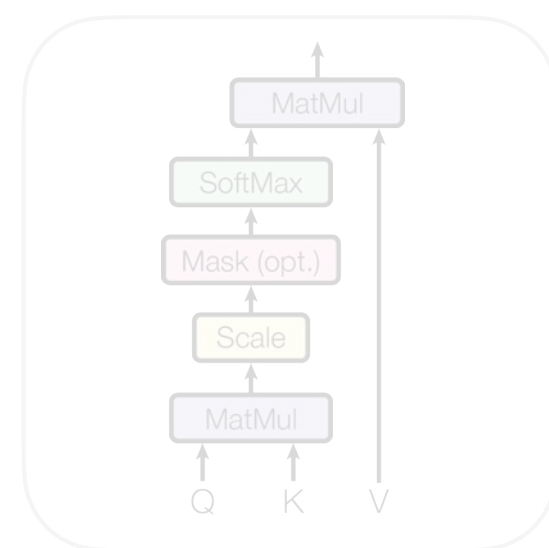






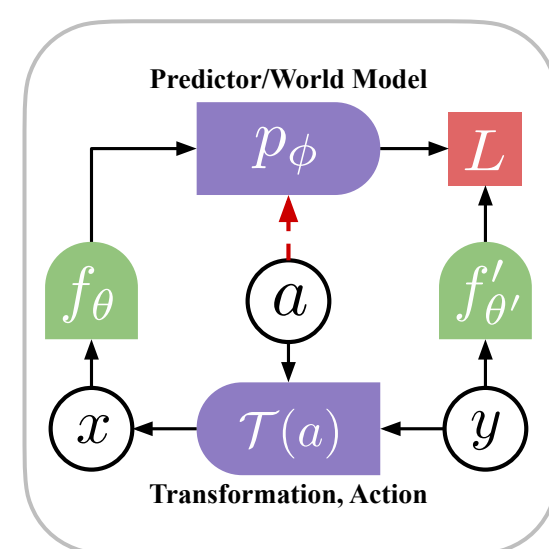
## Evolution of Jet Representations

*image, sequence, point cloud, ...*



## Attention Is All You Need?

*or how physics can help...*



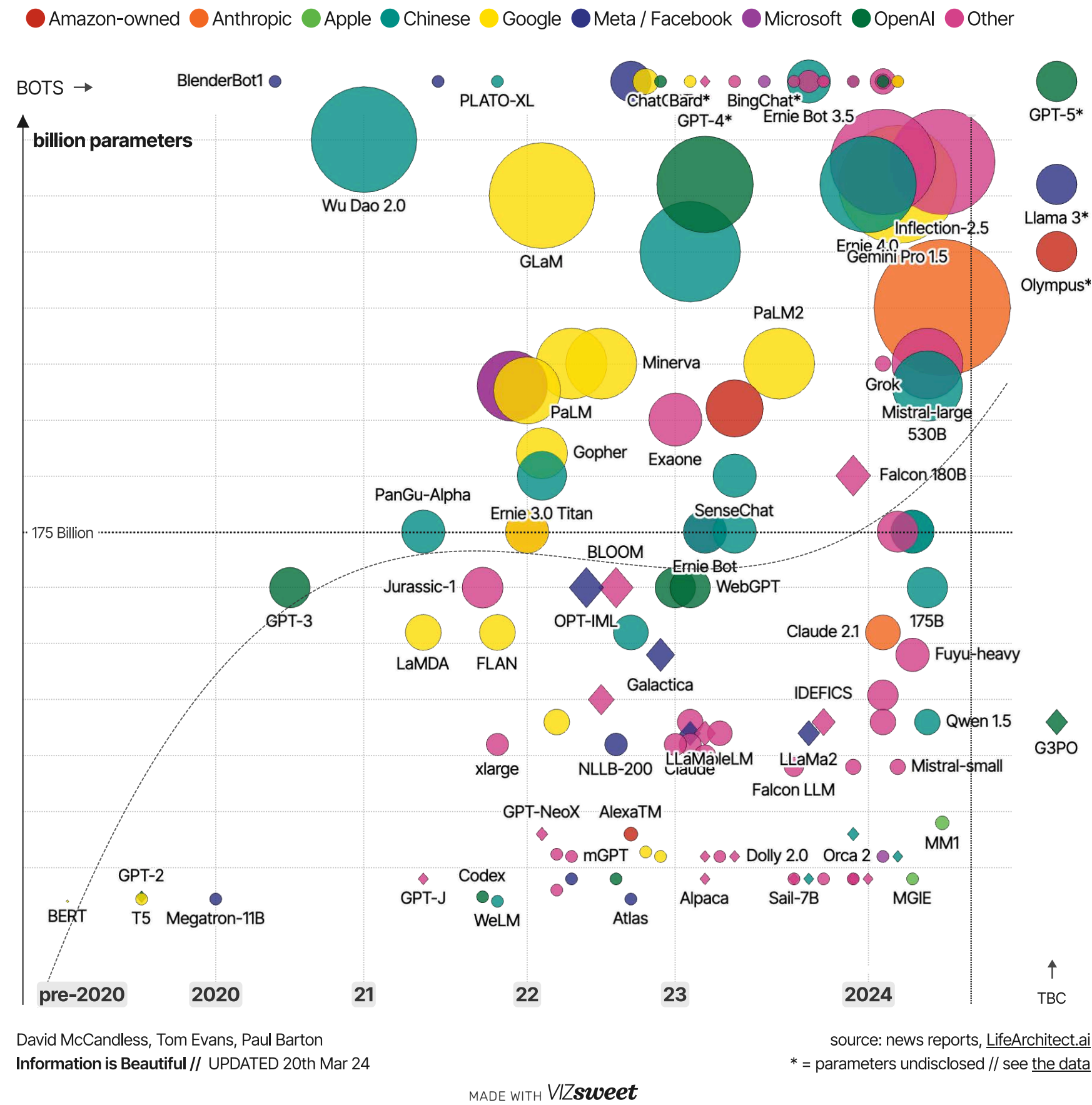
## Towards a Foundation Model

*scaling, self-supervision, and more...*



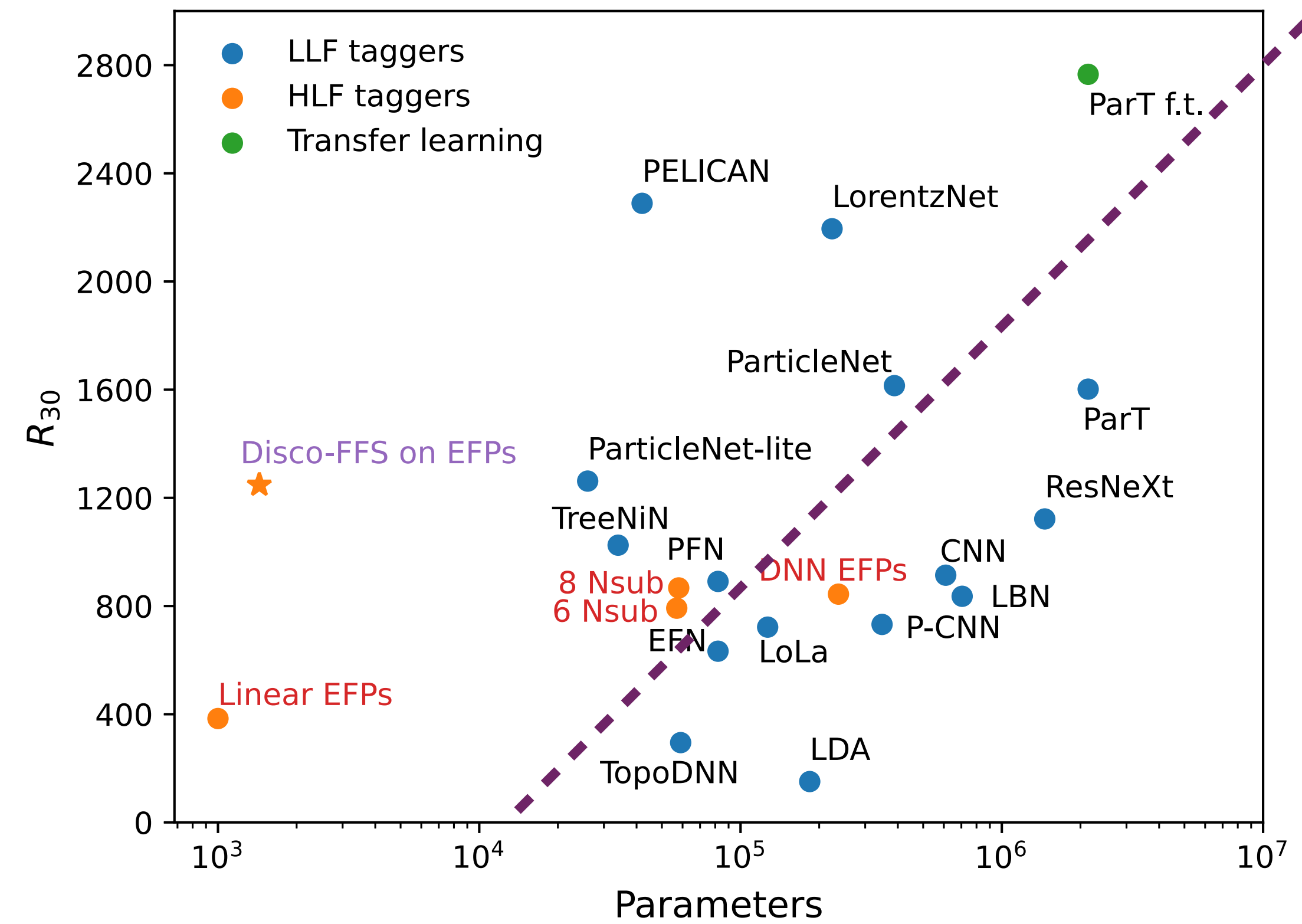
 information is beautiful

## Natural language models



Source: [informationisbeautiful.net](http://informationisbeautiful.net)

*HEP models (e.g., jet tagging)*

R. Das, G. Kasieczka and D. Shih, *PRD* 109, 054009



# SCALING LAWS

- Neural scaling law [arXiv: 2001.08361, 2203.15556]

- **question:** for a fixed compute budget (C), how to determine the optimal dataset size (D) and model parameters (N), such that the loss is minimized?

- **empirical power law:**

$$\hat{L}(N, D) \triangleq E + \frac{A}{N^\alpha} + \frac{B}{D^\beta}.$$

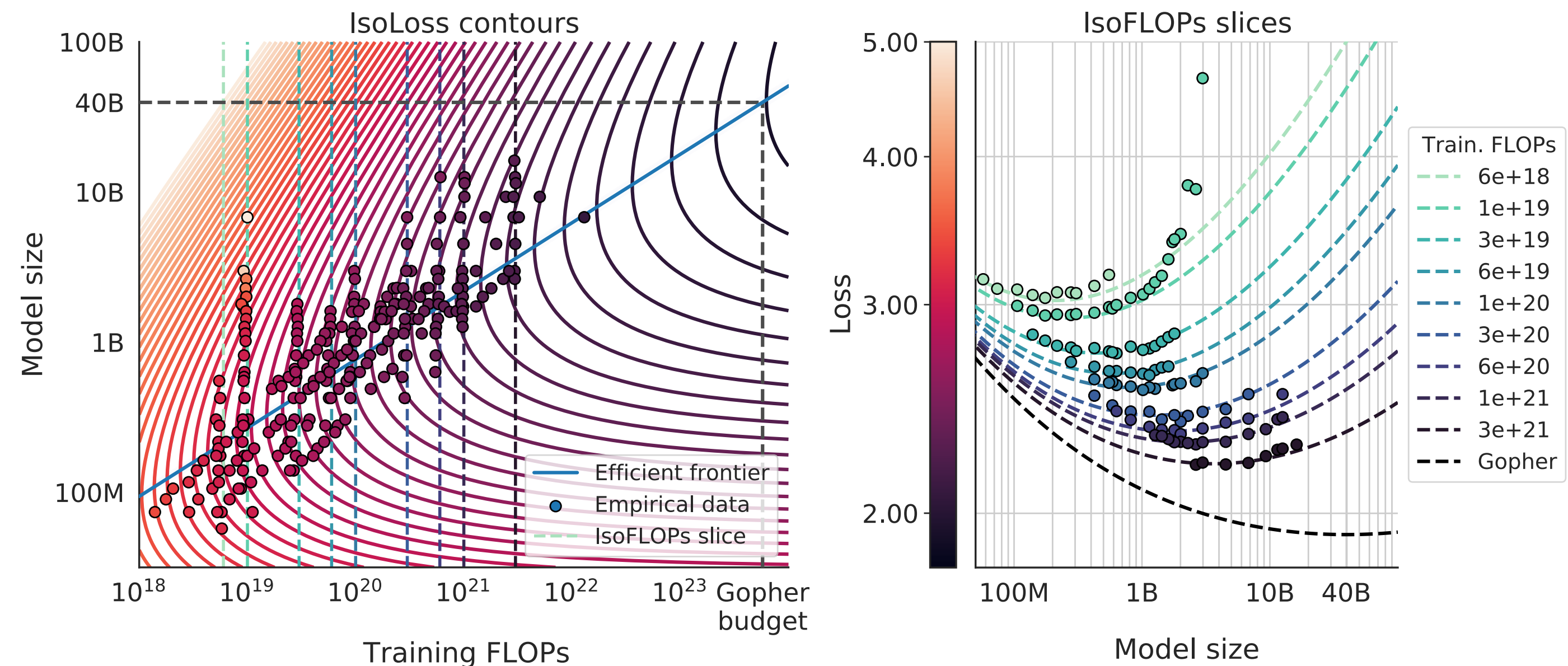
- for a standard transformer:

$$\text{FLOPs}(N, D) \approx 6ND$$

- **compute-optimal scaling:**

$$N_{\text{opt}}(C) = G \left( \frac{C}{6} \right)^a, \quad D_{\text{opt}}(C) = G^{-1} \left( \frac{C}{6} \right)^b,$$

where  $G = \left( \frac{\alpha A}{\beta B} \right)^{\frac{1}{\alpha+\beta}}$ ,  $a = \frac{\beta}{\alpha+\beta}$ , and  $b = \frac{\alpha}{\alpha+\beta}$ .



*Interestingly, for LLMs, the optimal scaling is found to be  $a \approx b \approx 0.50$  — same scaling for data and parameters.*



# JET SCALING LAW: STANDARD TRANSFORMER

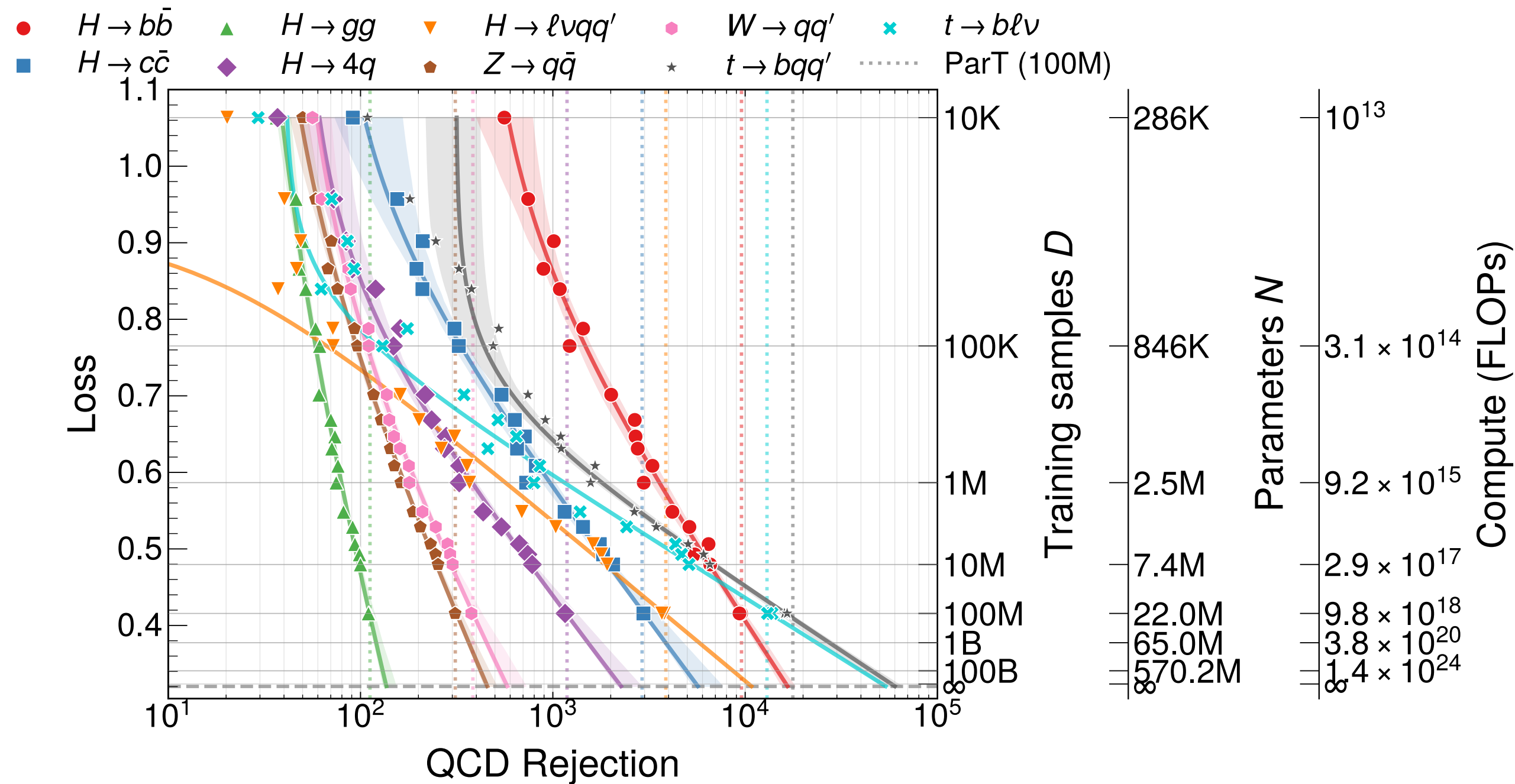
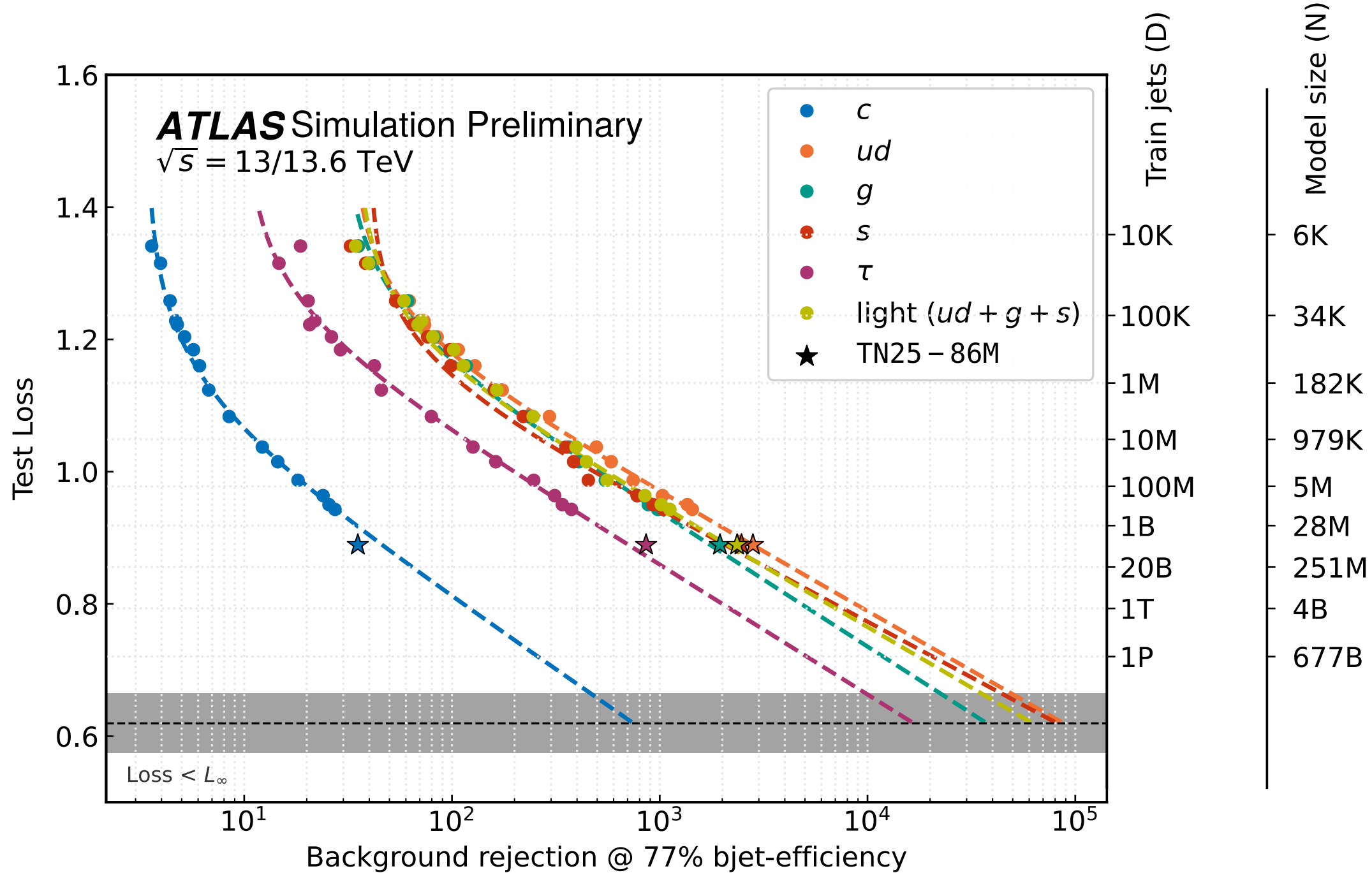
$$L(N, D) = L_{\infty} + \frac{A}{N^{\alpha}} + \frac{B}{D^{\beta}},$$

Jet flavor tagging (Small-R)

Boosted jet tagging (Large-R)

ATL-SOFT-PUB-2026-002

M.Vigl, N. Hartman, M. Kagan and L. Heinrich, [arXiv:2602.15781](#)



| Quantity                   | Symbol       | Description                                             | Value                |
|----------------------------|--------------|---------------------------------------------------------|----------------------|
| Irreducible loss           | $L_{\infty}$ | $L(N \rightarrow \text{inf}, D \rightarrow \text{inf})$ | 0.619 (0.539, 0.671) |
| Model scaling exponent     | $\alpha$     | $L_N \propto N^{-\alpha}$                               | 0.677 (0.644, 0.764) |
| Data scaling exponent      | $\beta$      | $L_D \propto D^{-\beta}$                                | 0.077 (0.065, 0.086) |
| $C = 6ND$ optimal exponent | $\gamma$     | $L \propto C^{-\gamma}$                                 | 0.07 (0.064, 0.074)  |

| Parameter    | Value (68% CI)      |
|--------------|---------------------|
| $L_{\infty}$ | 0.32 (0.29, 0.34)   |
| $A$          | 11.27 (8.01, 15.14) |
| $\alpha$     | 0.44 (0.40, 0.48)   |
| $B$          | 7.22 (6.66, 7.89)   |
| $\beta$      | 0.22 (0.21, 0.23)   |
| $\gamma$     | 0.15 (0.13, 0.16)   |



# JET SCALING LAW: STANDARD TRANSFORMER

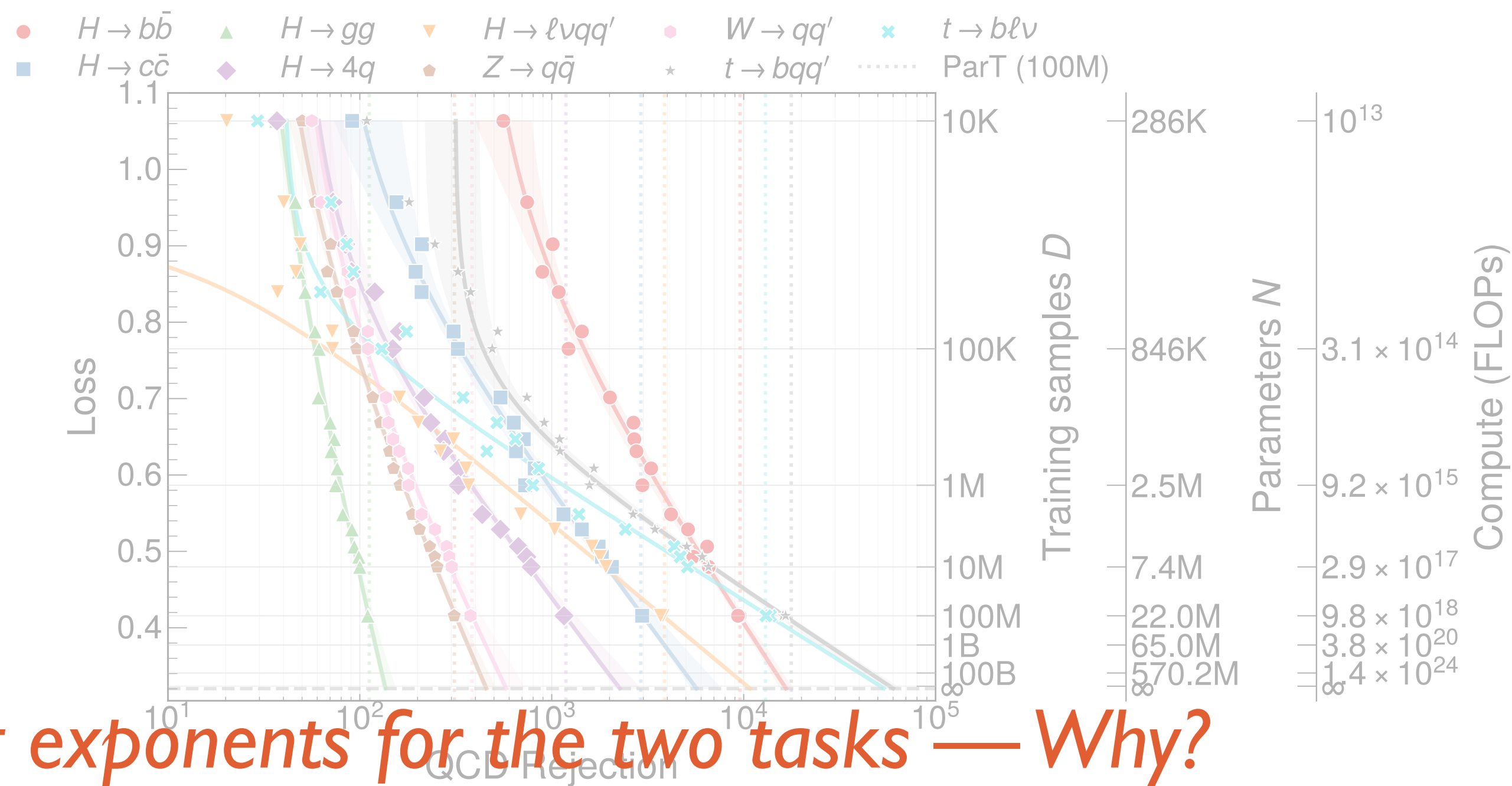
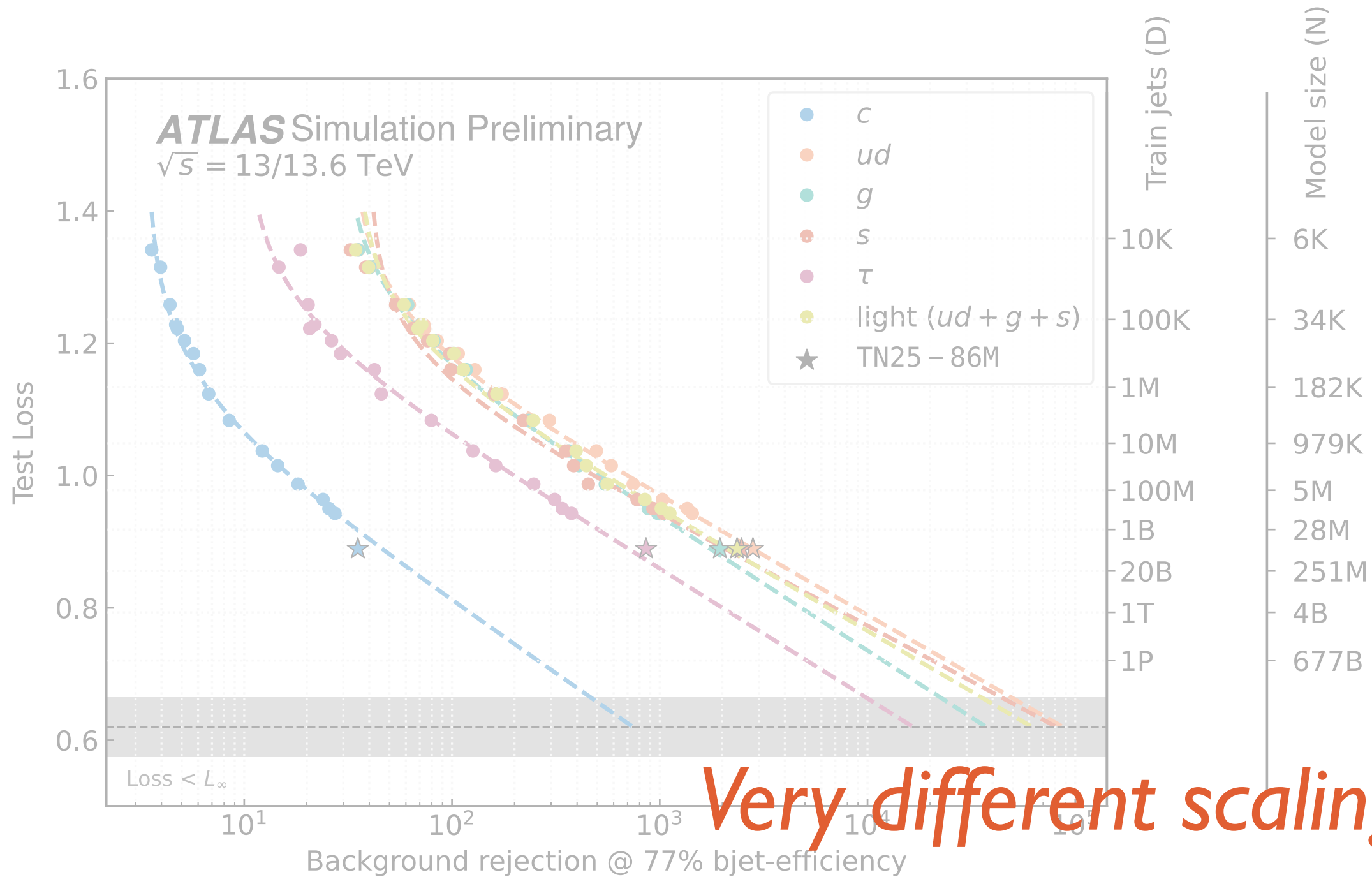
$$L(N, D) = L_{\infty} + \frac{A}{N^{\alpha}} + \frac{B}{D^{\beta}},$$

Jet flavor tagging (Small-R)

Boosted jet tagging (Large-R)

ATL-SOFT-PUB-2026-002

M.Vigl, N. Hartman, M. Kagan and L. Heinrich, [arXiv:2602.15781](#)



Very different scaling exponents for the two tasks — Why?

| Quantity                   | Symbol       | Description                                             | Value                |
|----------------------------|--------------|---------------------------------------------------------|----------------------|
| Irreducible loss           | $L_{\infty}$ | $L(N \rightarrow \text{inf}, D \rightarrow \text{inf})$ | 0.619 (0.539, 0.671) |
| Model scaling exponent     | $\alpha$     | $L_N \propto N^{-\alpha}$                               | 0.677 (0.644, 0.764) |
| Data scaling exponent      | $\beta$      | $L_D \propto D^{-\beta}$                                | 0.077 (0.065, 0.086) |
| $C = 6ND$ optimal exponent | $\gamma$     | $L \propto C^{-\gamma}$                                 | 0.07 (0.064, 0.074)  |

| Parameter    | Value (68% CI)      |
|--------------|---------------------|
| $L_{\infty}$ | 0.32 (0.29, 0.34)   |
| $A$          | 11.27 (8.01, 15.14) |
| $\alpha$     | 0.44 (0.40, 0.48)   |
| $B$          | 7.22 (6.66, 7.89)   |
| $\beta$      | 0.22 (0.21, 0.23)   |
| $\gamma$     | 0.15 (0.13, 0.16)   |



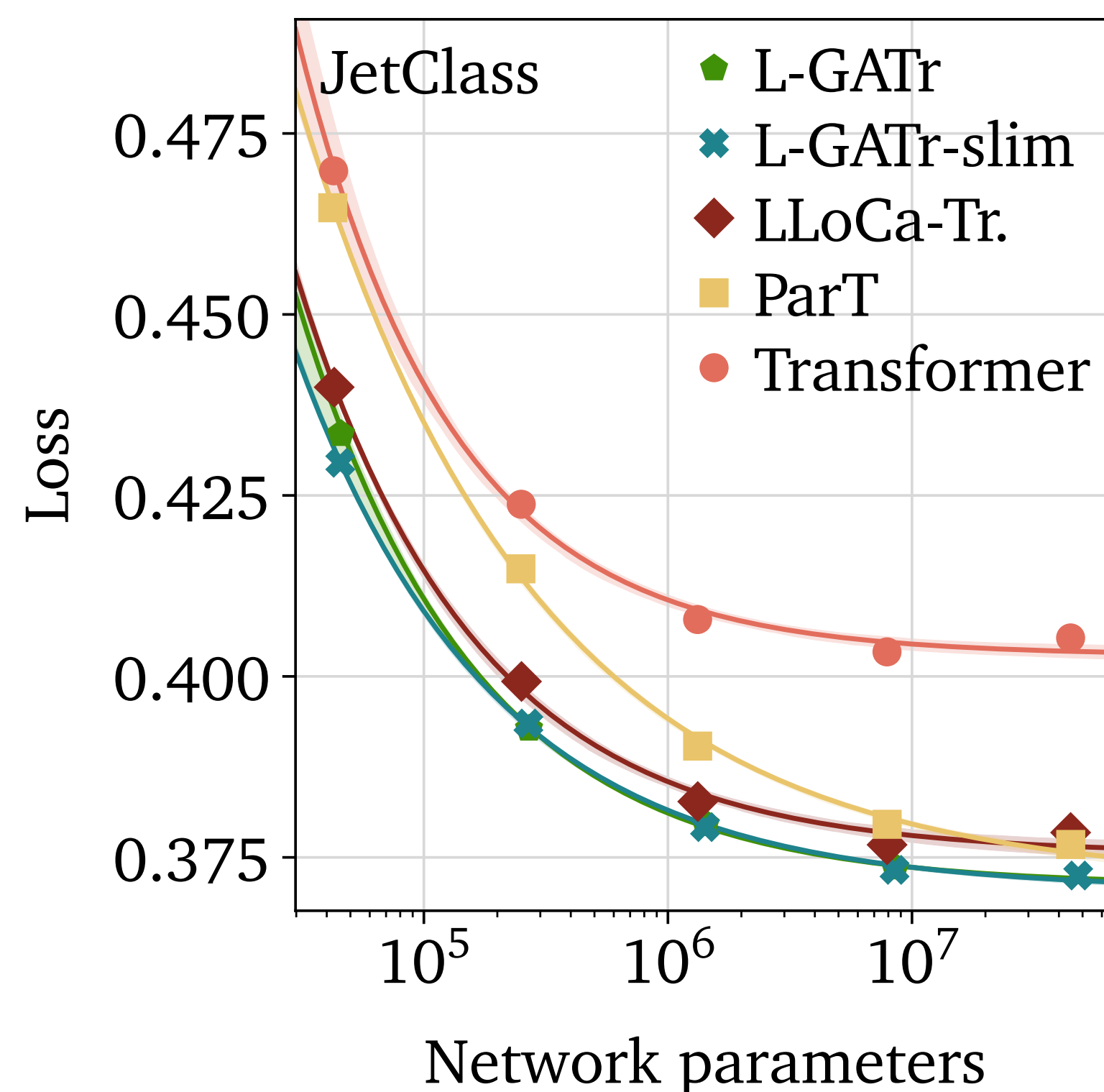
# JET SCALING LAW: EQUIVARIANT TRANSFORMERS



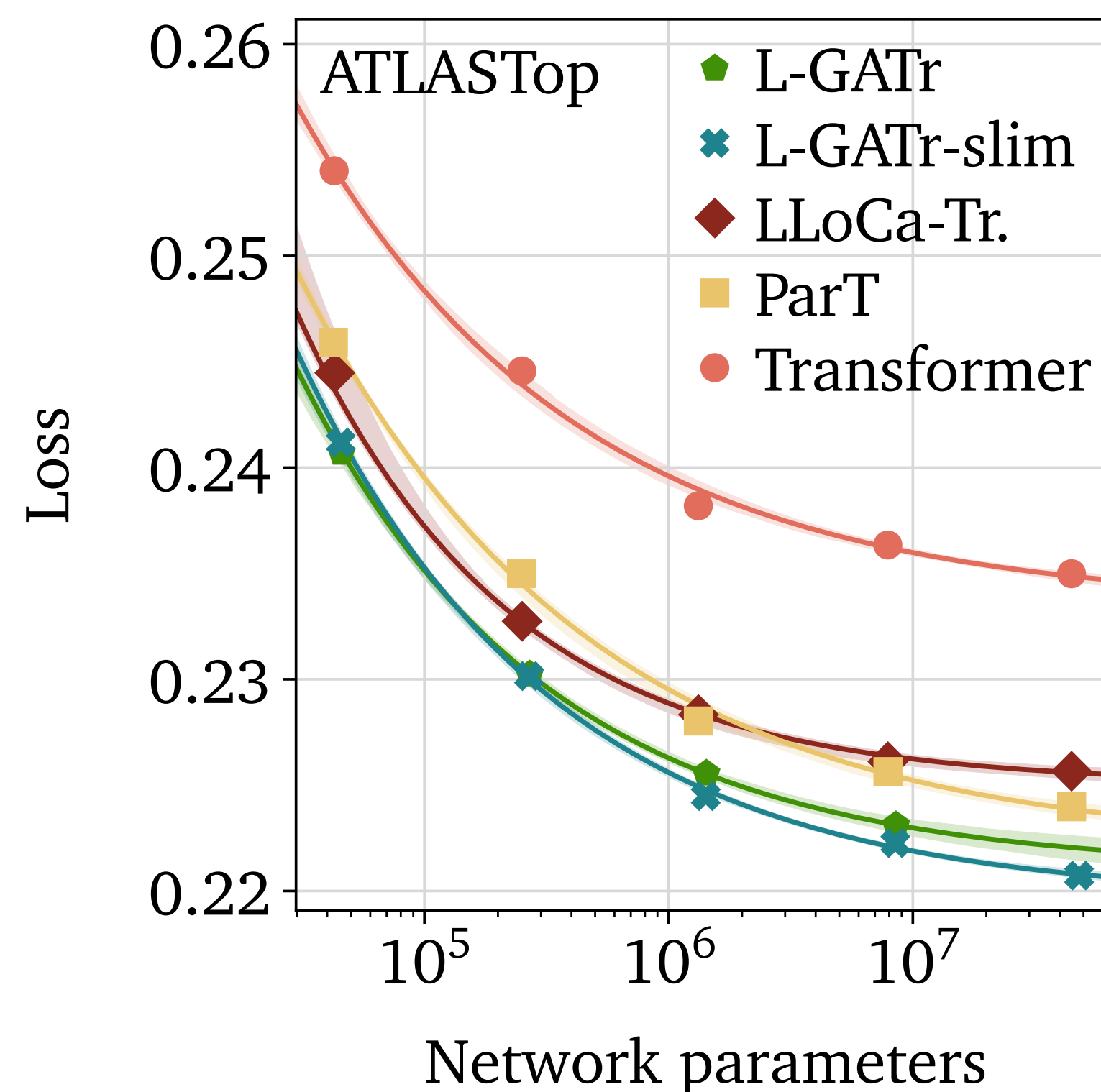
L. Favaro, T. Plehn, HQ, J. Spinner  
[arXiv: 2608.02735](https://arxiv.org/abs/2608.02735)

- How do Lorentz equivariant transformers scale?
  - fixed dataset size (~100M); only varying model sizes

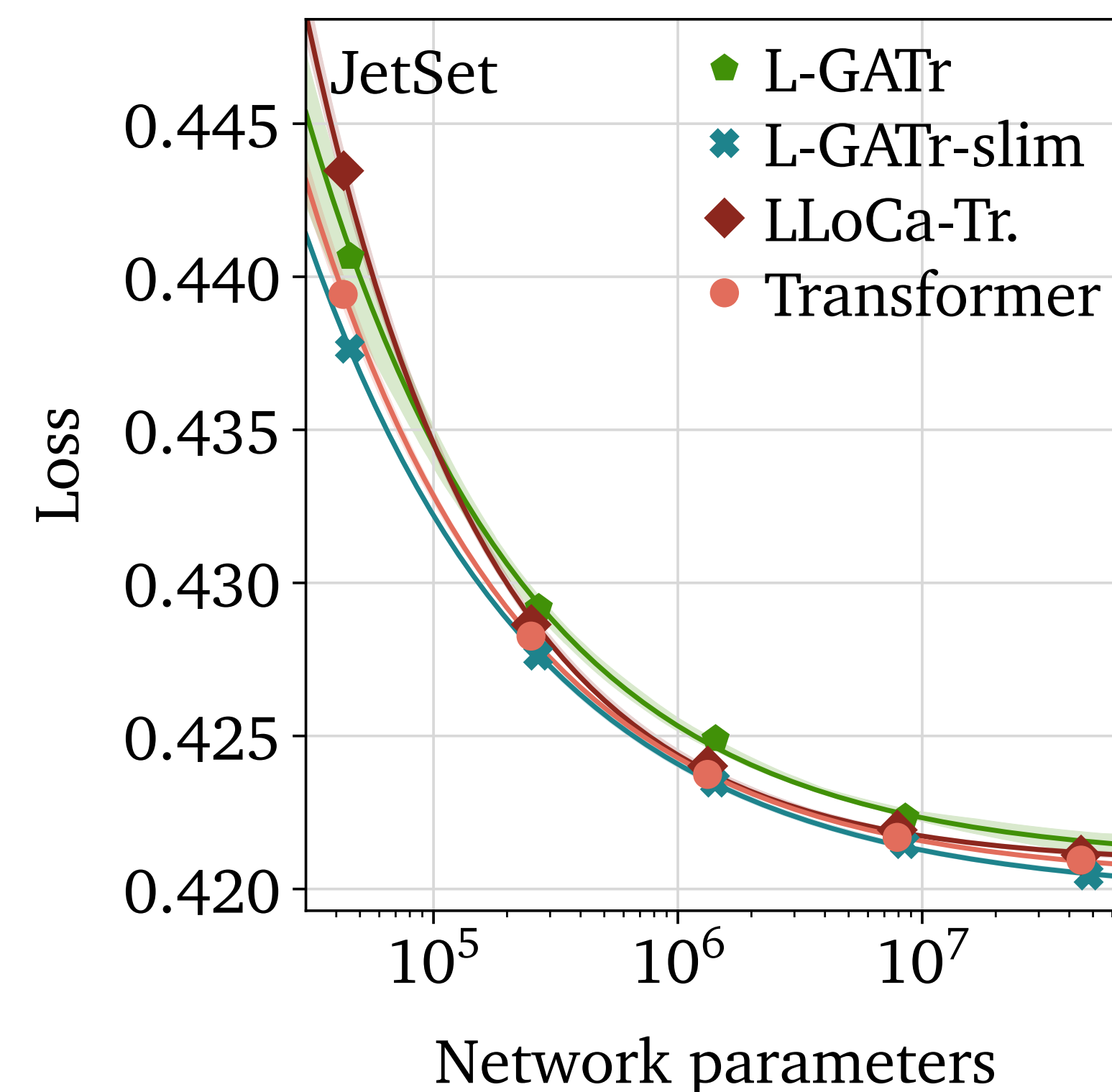
*Boosted jet tagging  
(JetClass, all features)*



*Boosted jet tagging  
(Top vs QCD, 4-vector only)*



*Jet flavor tagging  
(Small-R jets, tracks only)*





# JET SCALING LAW: EQUIVARIANT TRANSFORMERS



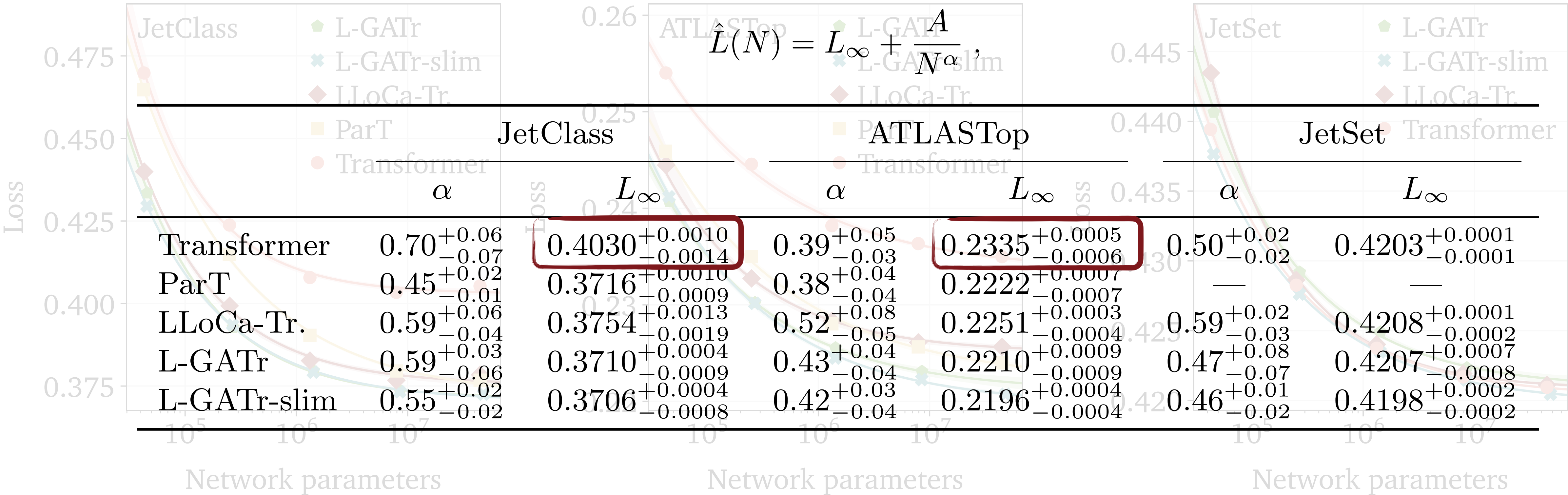
L. Favaro, T. Plehn, HQ, J. Spinner  
[arXiv: 2608.02735](https://arxiv.org/abs/2608.02735)

- How do Lorentz equivariant transformers scale?
  - fixed dataset size (~100M); only varying model sizes

Boosted jet tagging  
(JetClass, all features)

Boosted jet tagging  
(Top vs QCD, 4-vector only)

Jet flavor tagging  
(Small-R jets, tracks only)

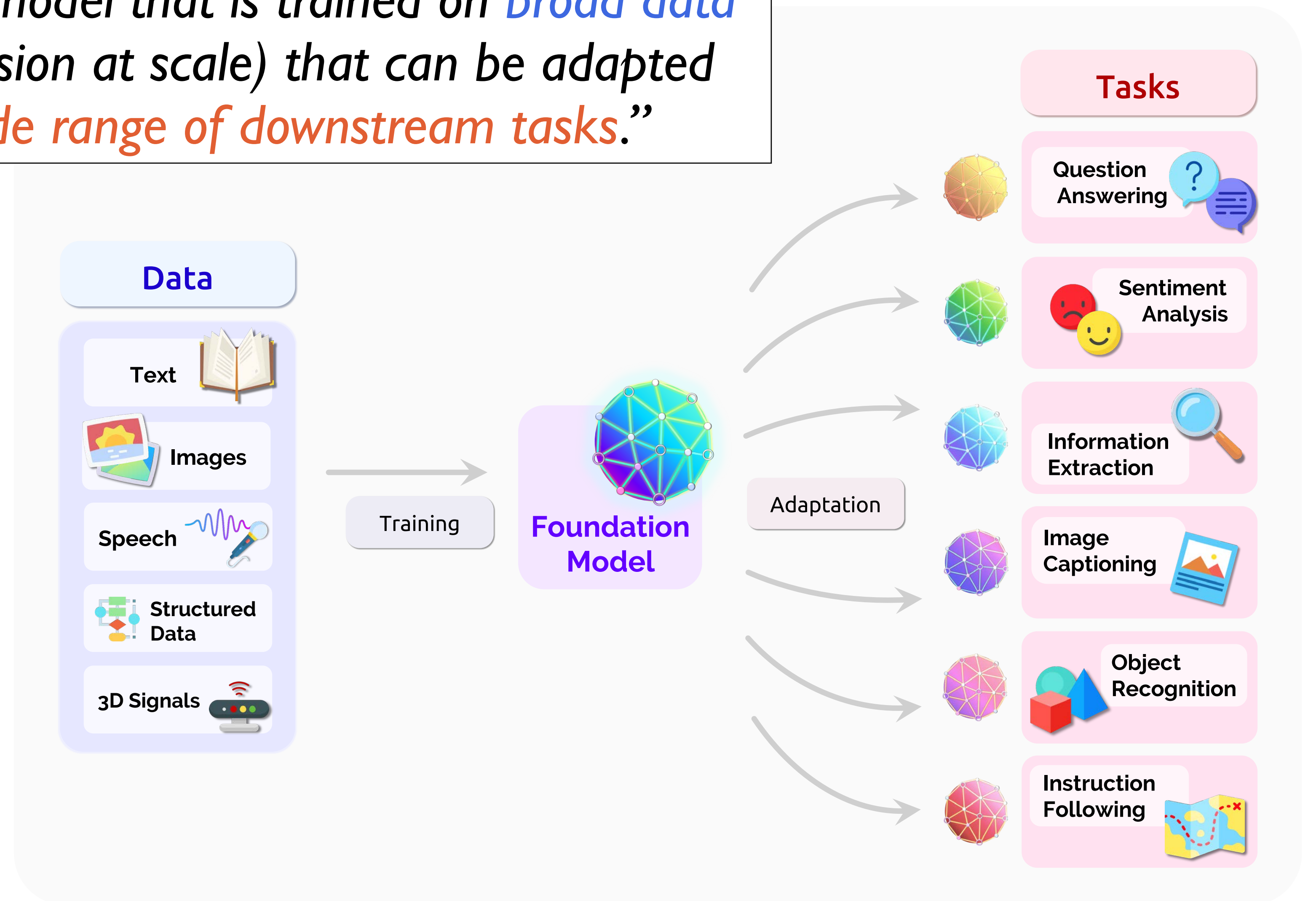




# FOUNDATION MODEL

“A foundation model is any model that is trained on *broad data* (generally using self-supervision at scale) that can be adapted (e.g., fine-tuned) to *a wide range of downstream tasks*.”

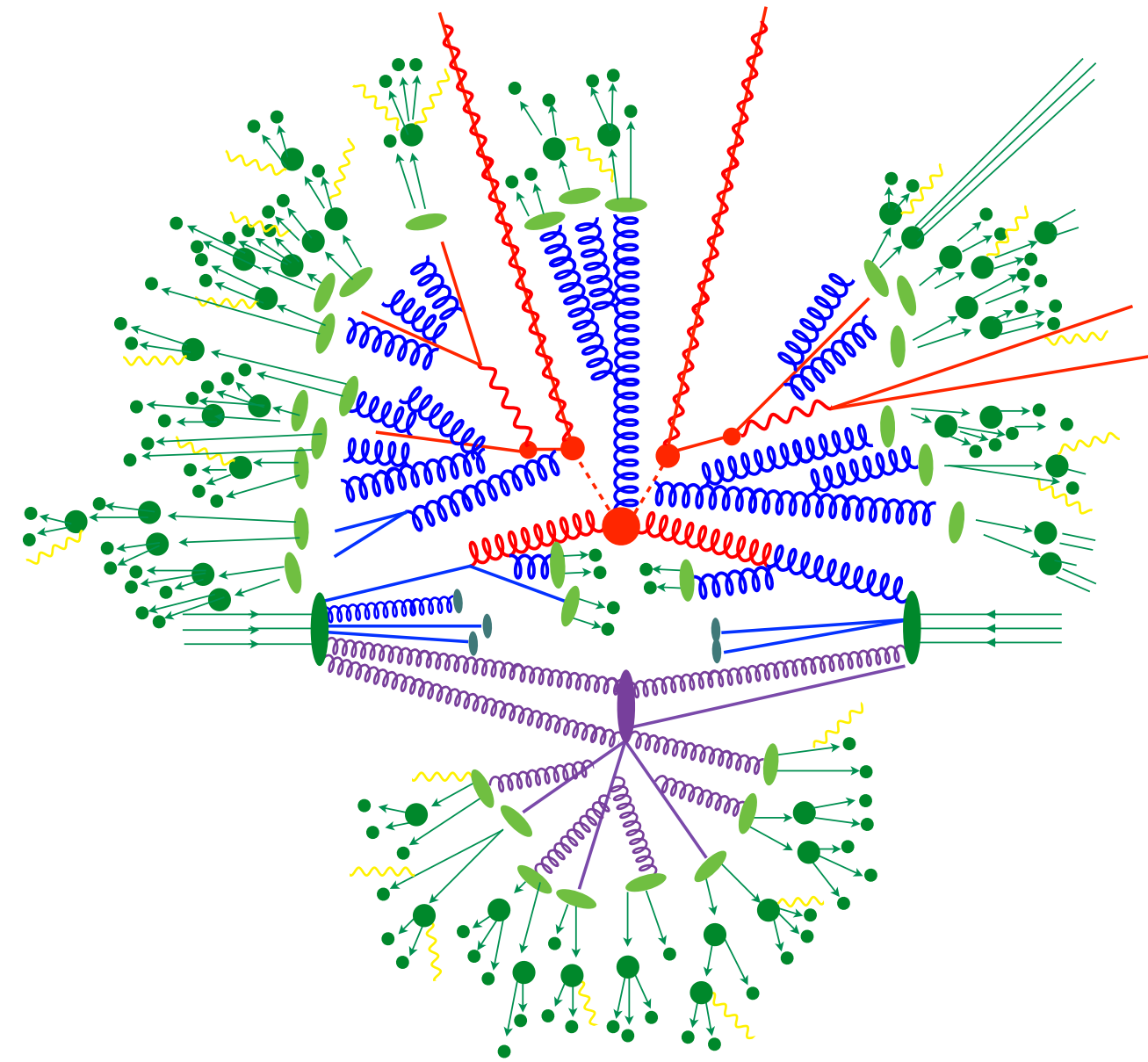
[arXiv: 2108.07258]





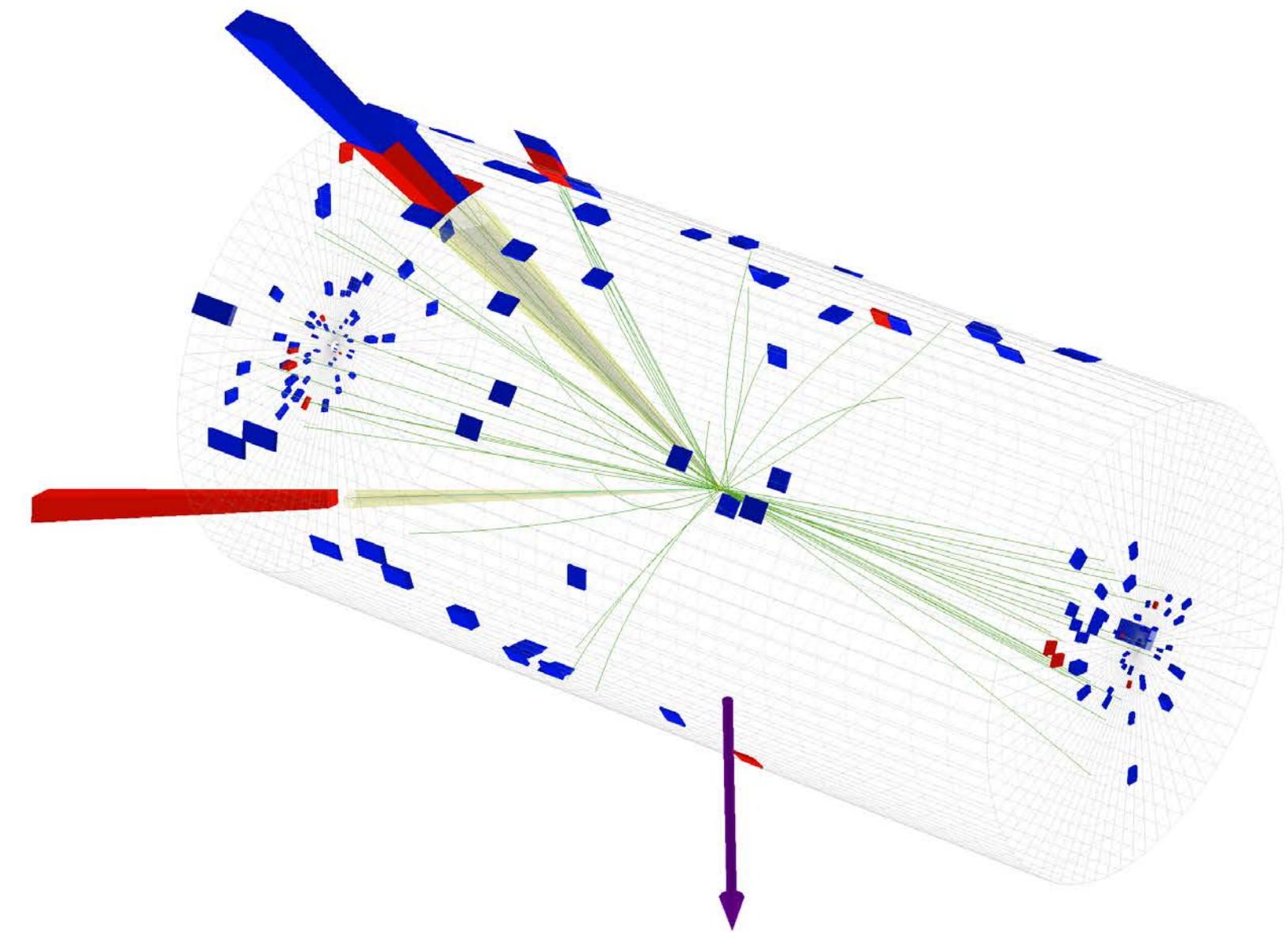
# TOWARDS A FOUNDATION MODEL

- Foundation models need to be (pre-)trained on large-scale dataset



*Train on MC simulations*

*Labels are available—  
Supervised Learning*



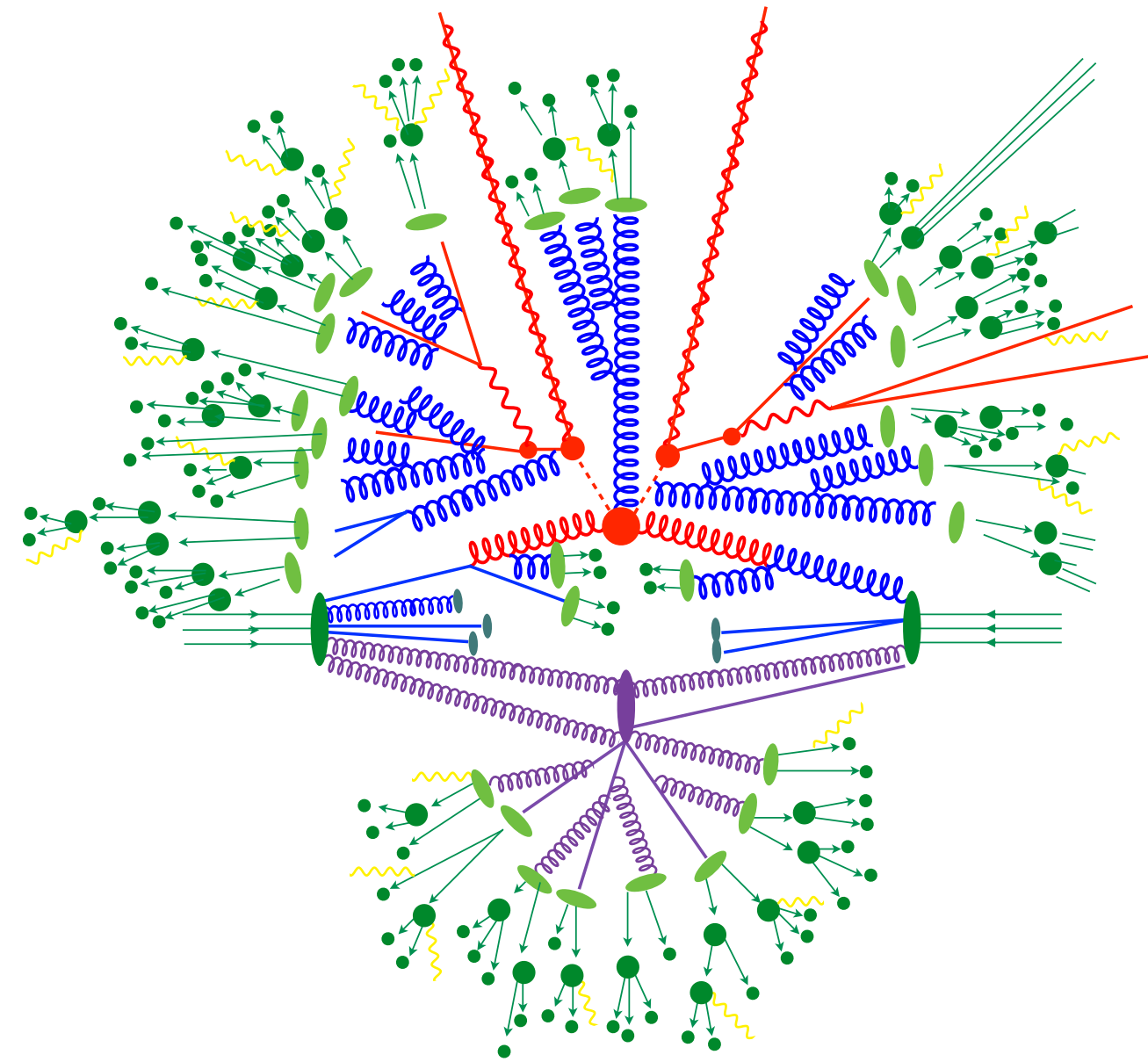
*Train on real data*

*Labels are not available—  
Self-Supervised Learning (SSL)*



# TOWARDS A FOUNDATION MODEL

- Foundation models need to be (pre-)trained on large-scale dataset



*Train on MC simulations*

*Labels are available—  
Supervised Learning*



See e.g., “**OmniLearn**” [V. Mikuni and B. Nachman, [arXiv: 2404.16091](#), [arXiv: 2502.14652](#)],

“**Sophon**” / “**GloParT**” [C. Li, et al., [arXiv: 2405.12972](#)],

“**UParT**” [CMS, [CMS-DP-2024-066](#), [CMS-DP-2025-081](#)],

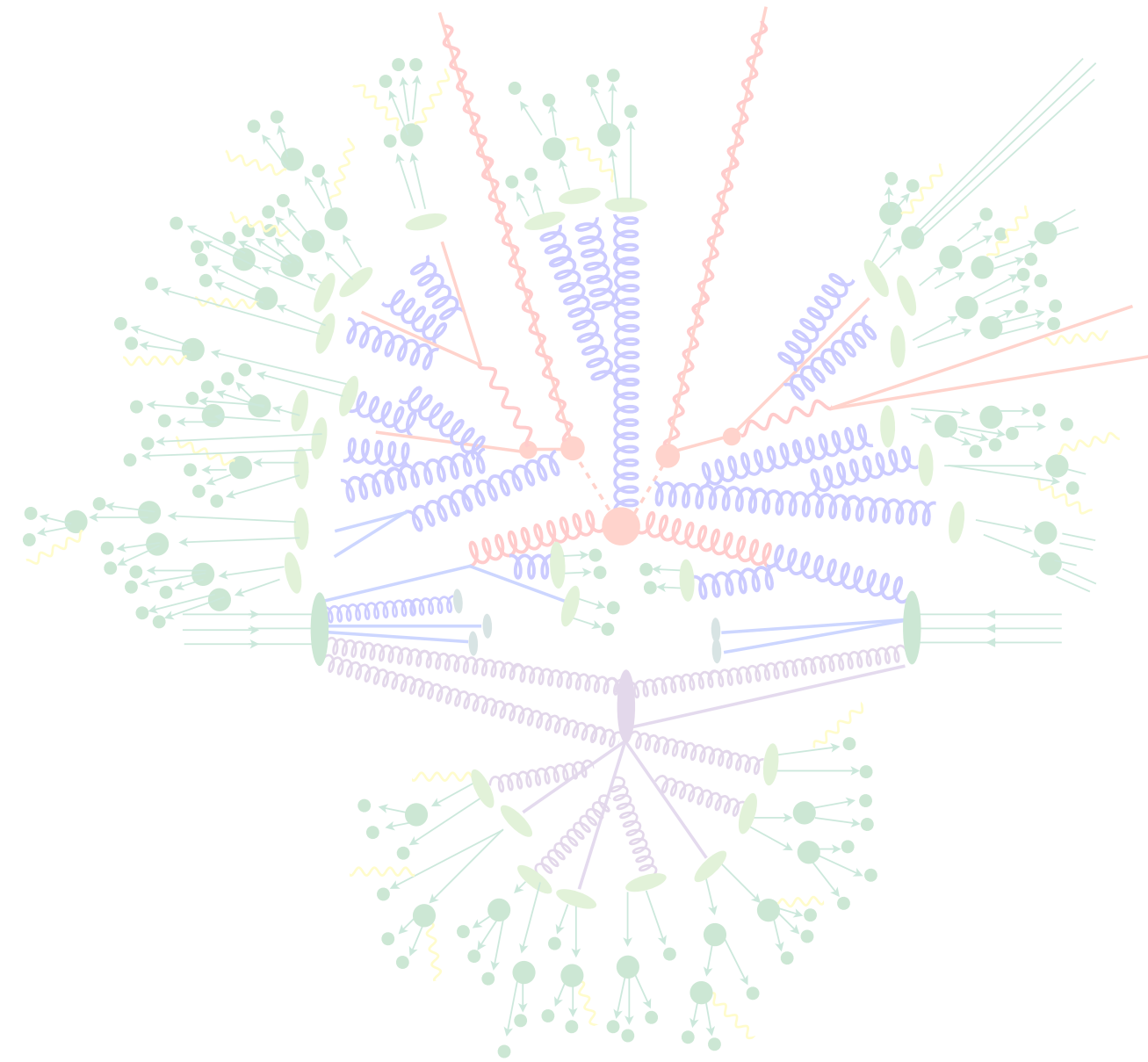
“**OmniLearned**” [W. Bhimji, C. Harris, V. Mikuni and B. Nachman, [arXiv: 2510.24066](#)],

“**EveNet**” [T. H. Hsu, et al., [arXiv: 2601.17126](#)].

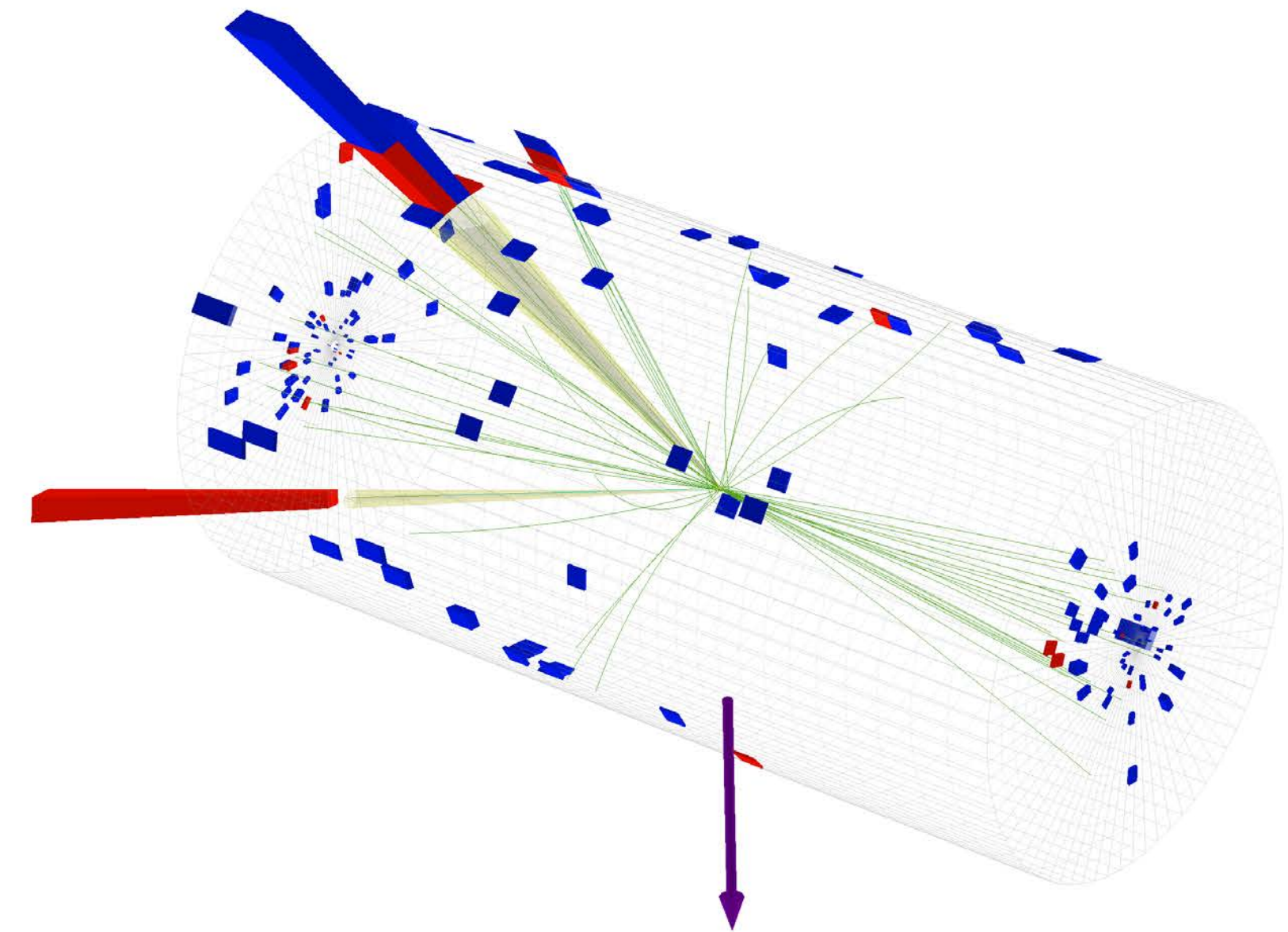


# TOWARDS A FOUNDATION MODEL

- Foundation models need to be (pre-)trained on large-scale dataset



*Train on MC simulations*  
*Labels are available—*  
*Supervised Learning*

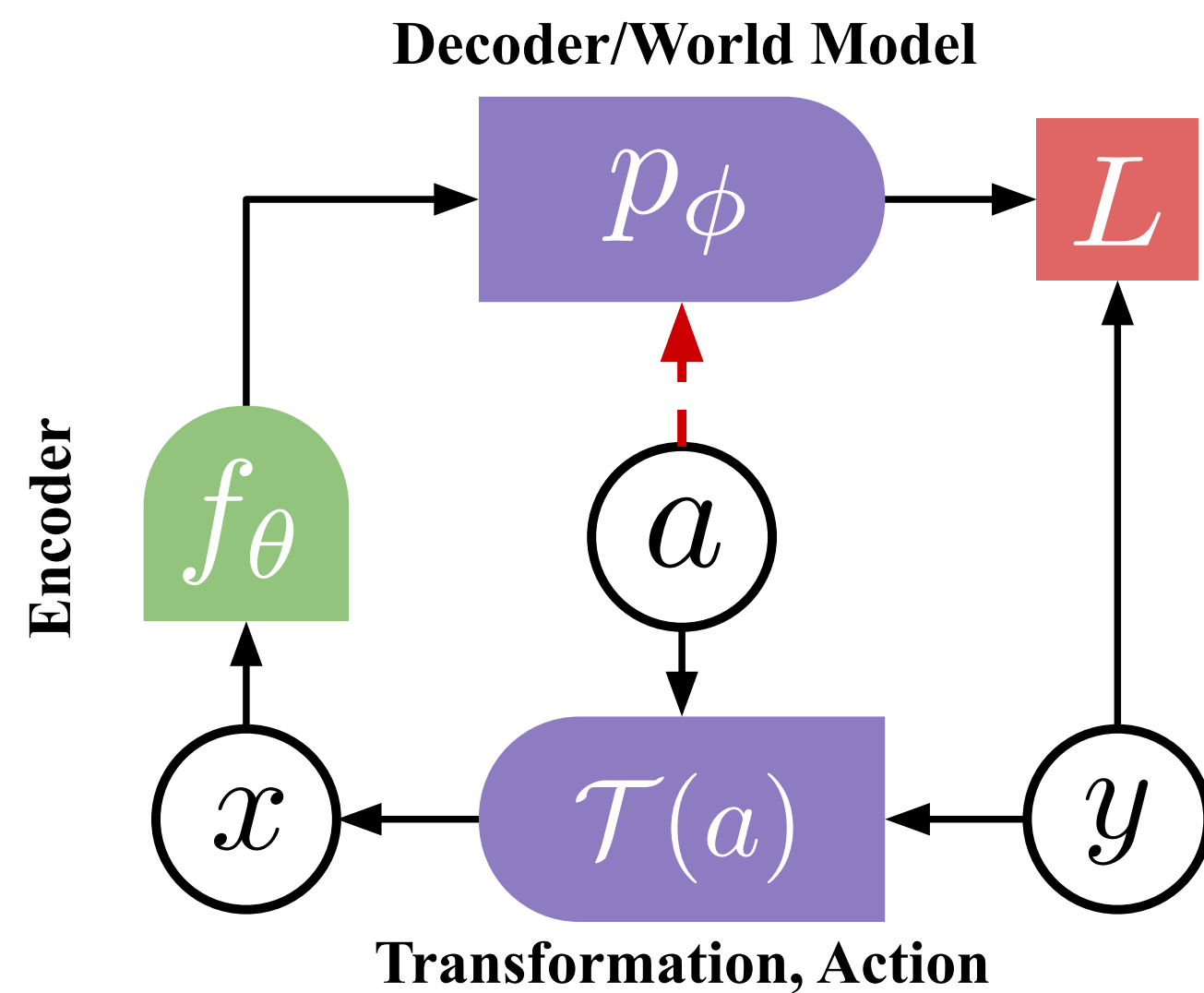


*Train on real data*  
*Labels are not available—*  
*Self-Supervised Learning (SSL)*



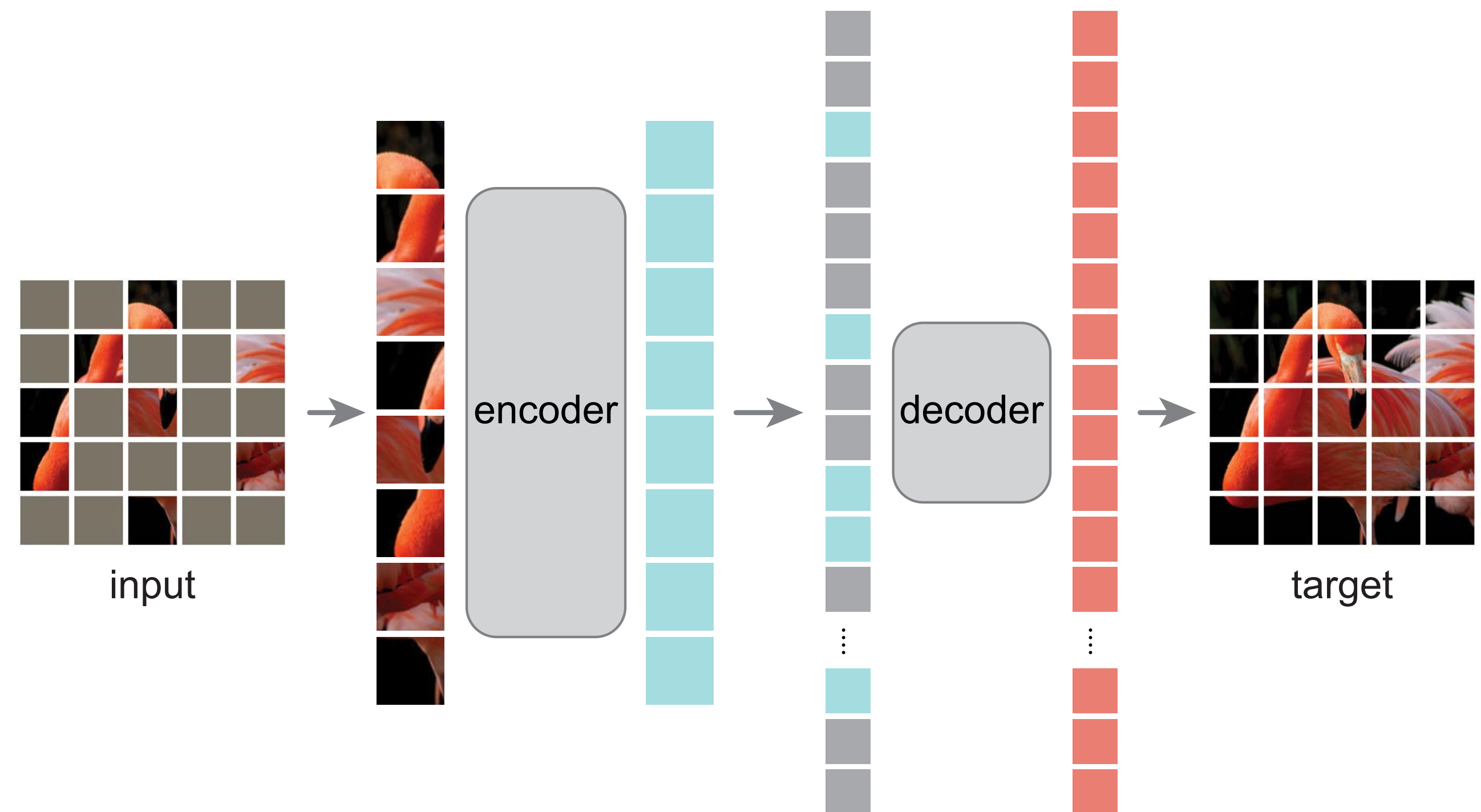
# SELF-SUPERVISION: MASKED MODELING

Generative Architecture



Learns to invert a transformation  
in the **input** space.  
e.g., Masked Autoencoder (MAE)

Masked Autoencoder [arXiv: 2111.06377]

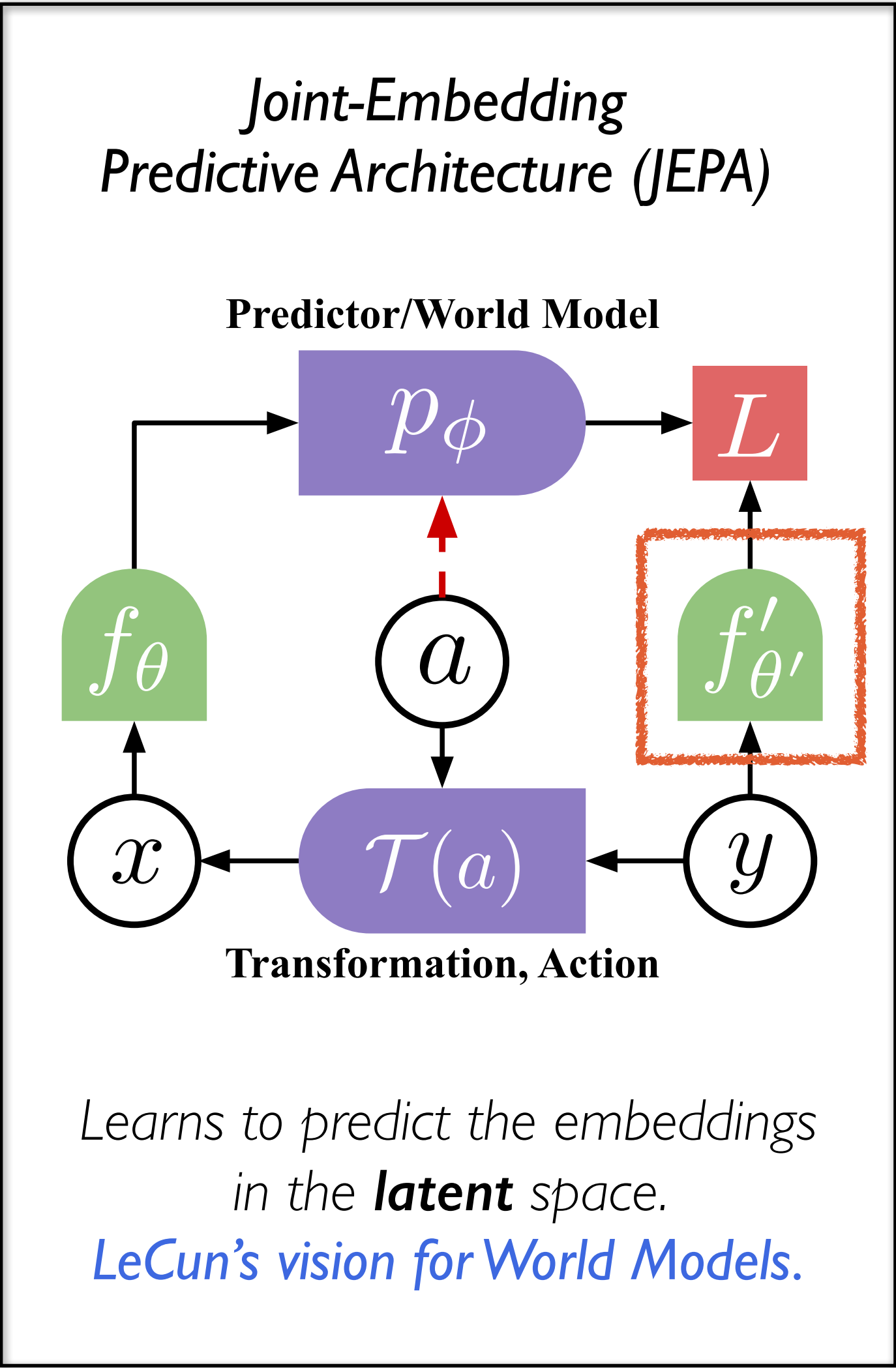


... masks random patches of the input image  
and reconstructs the missing pixels

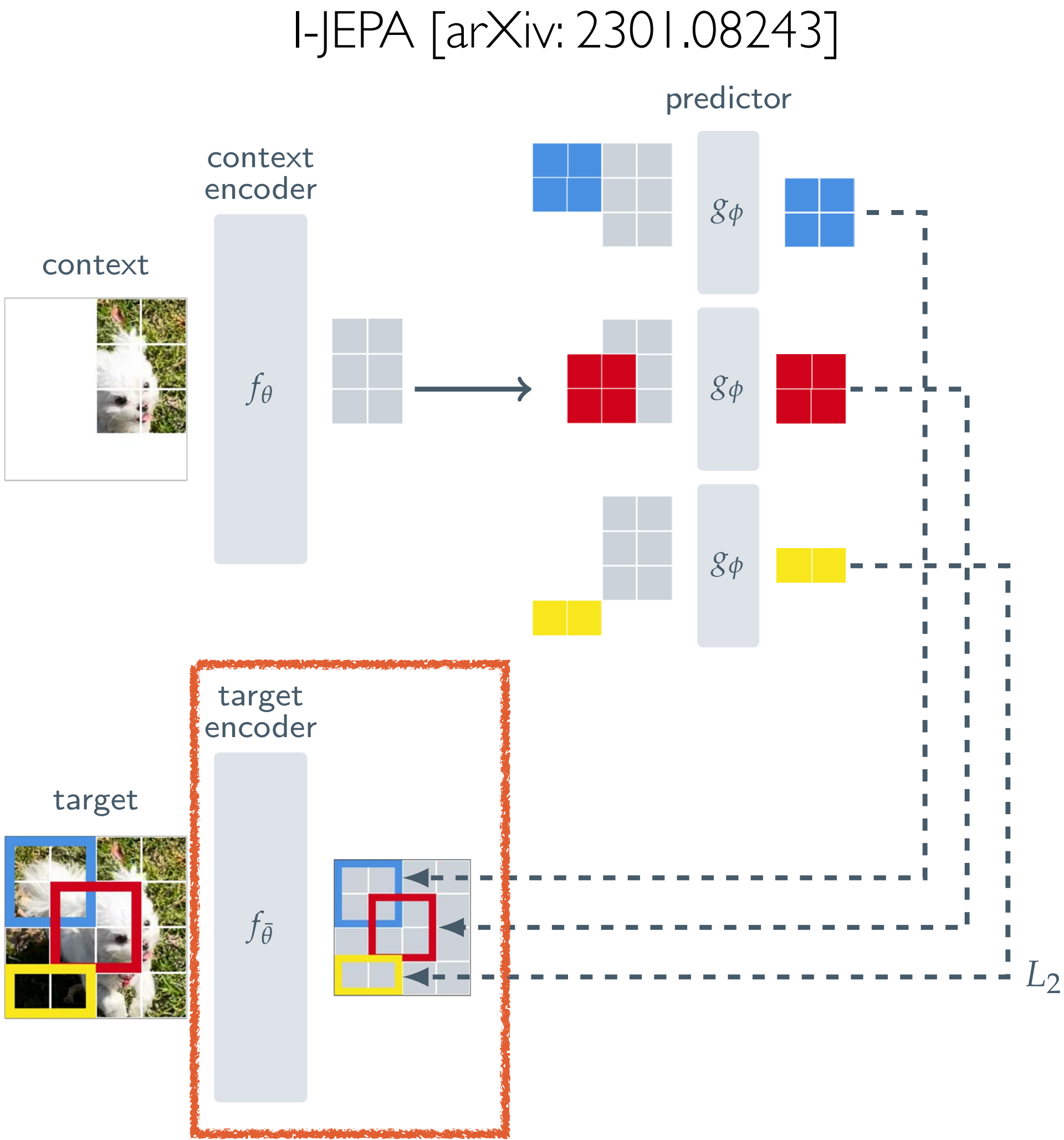
See e.g. "MPM" [T. Gollig, L. Heinrich, M. Kagan, S. Klein, M. Leigh, M. Osadchy and J.A. Raine, [arXiv: 2401.13537](#)],  
"MPMv2" [M. Leigh, S. Klein, F. Charton, T. Gollig, L. Heinrich, M. Kagan, I. Ochoa and M. Osadchy, [arXiv: 2409.12589](#)]



# SELF-SUPERVISION: JEPA



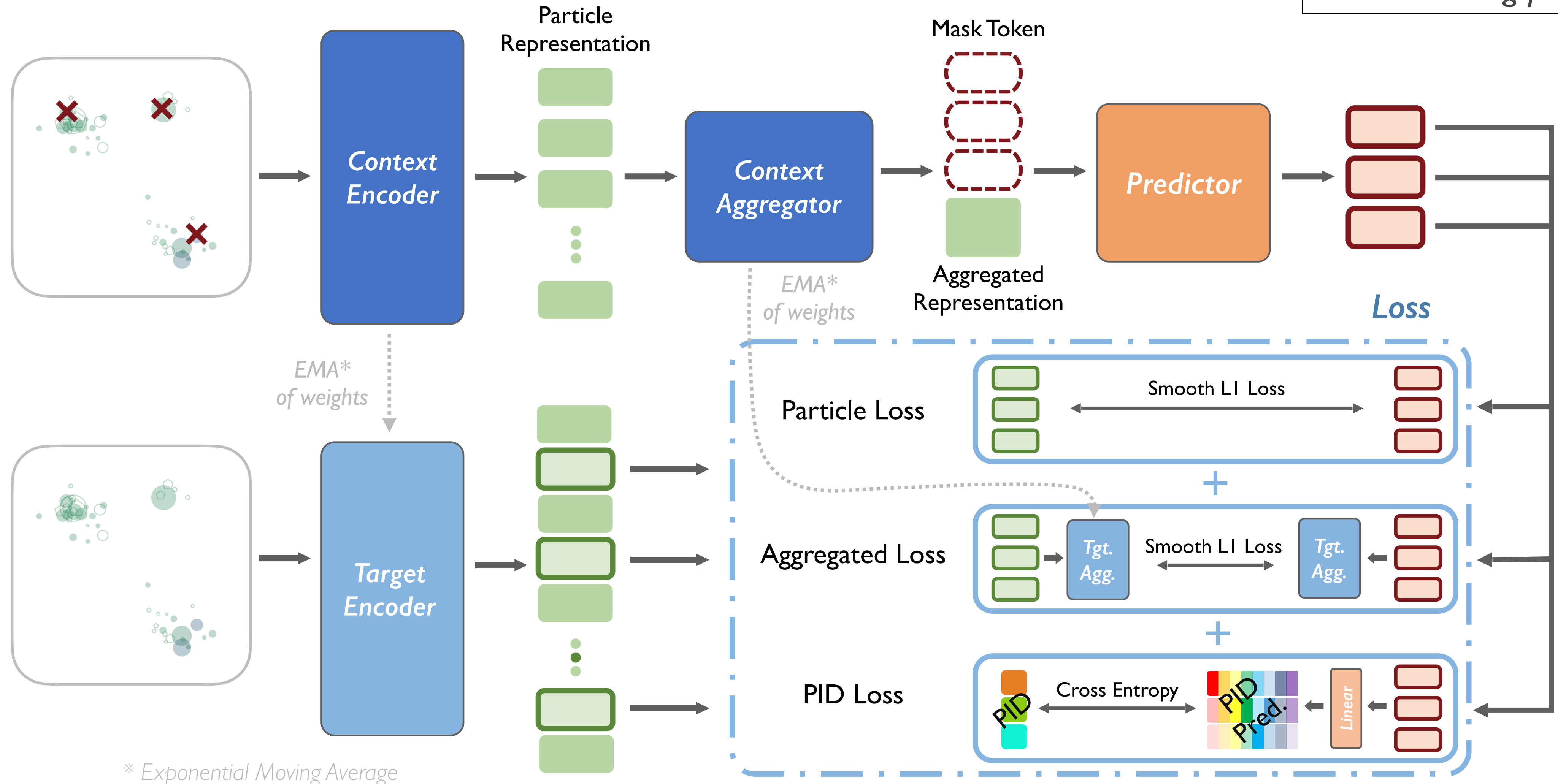
arXiv: 2403.00504



... predicts the embeddings of masked image patches  
in a (learned) latent space.

# P-JEPA

Joint work with  
Qibin Liu, Shudong Wang  
and Congqiao Li

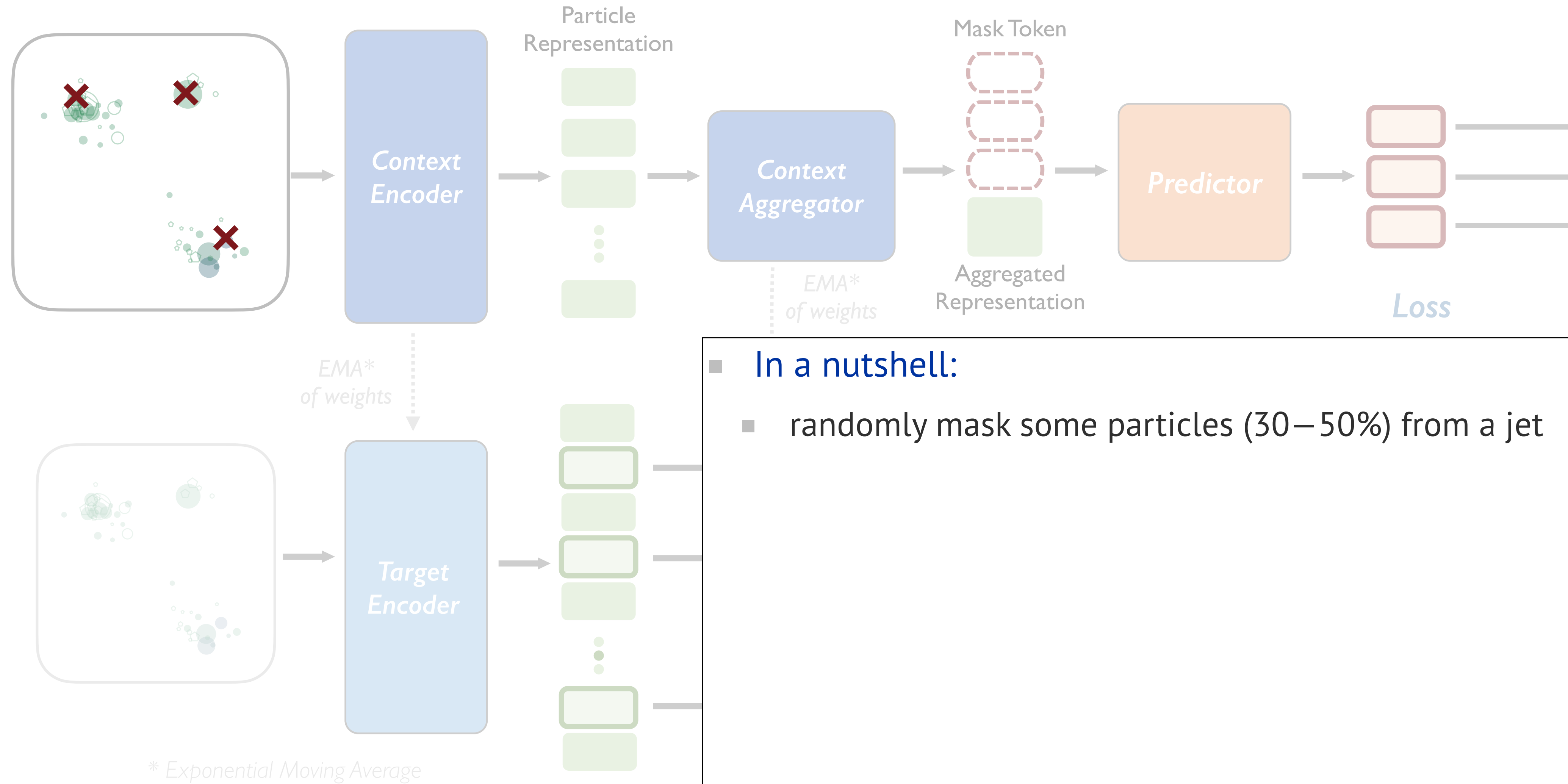


\* Exponential Moving Average

See also: "J-JEPA" [S. Katel, H. Li, Z. Zhao, F. Mokhtar, J. Duarte and R. Kansal, [arXiv: 2412.05333](#)],  
"HEP-JEPA" [J. Bardhan, R. Agrawal, A. Tilak, C. Neeraj and S. Mitra, [arXiv: 2502.03933](#)]

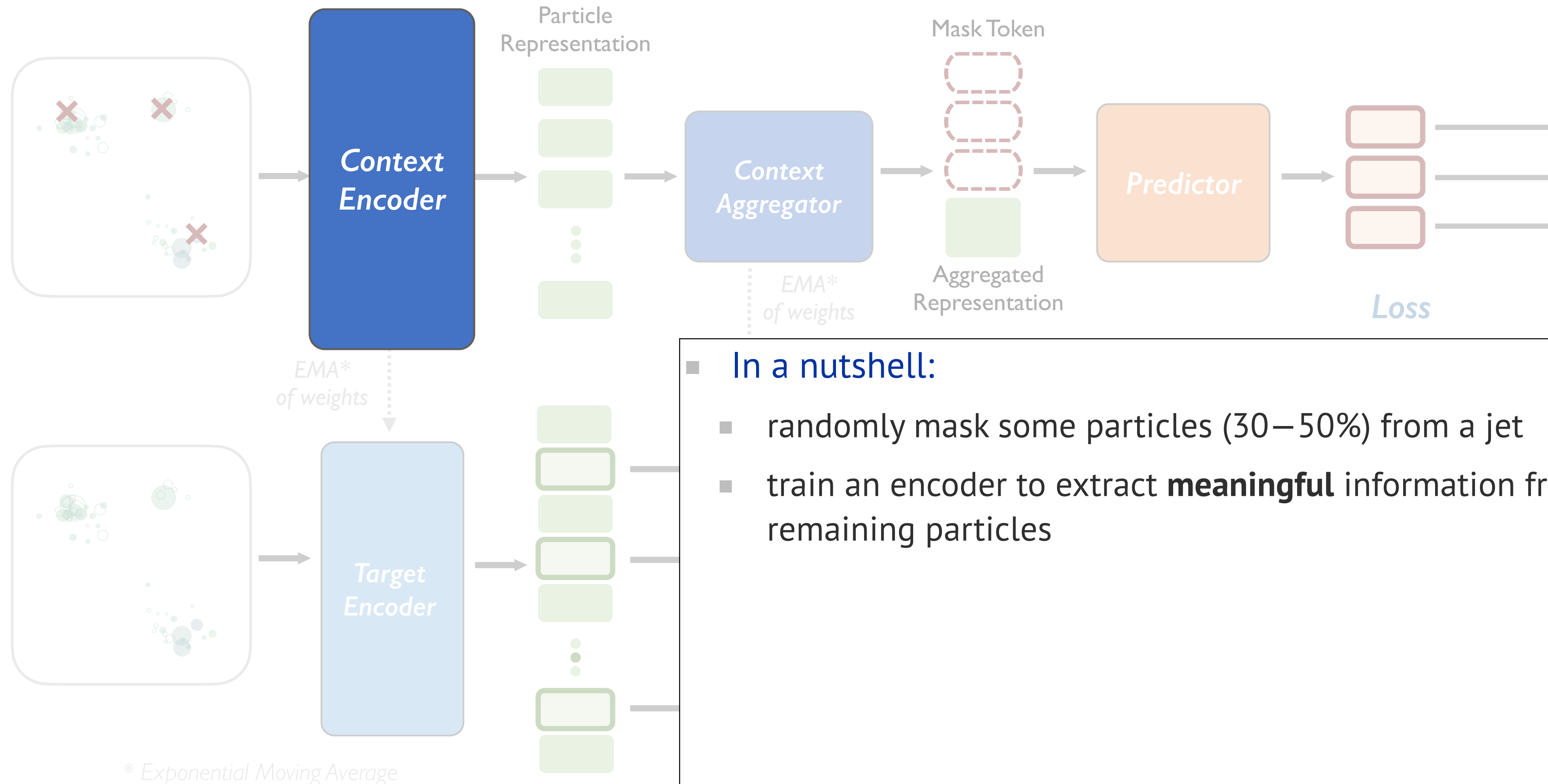


# P-JEPA



- In a nutshell:
  - randomly mask some particles (30–50%) from a jet

# P-JEPA

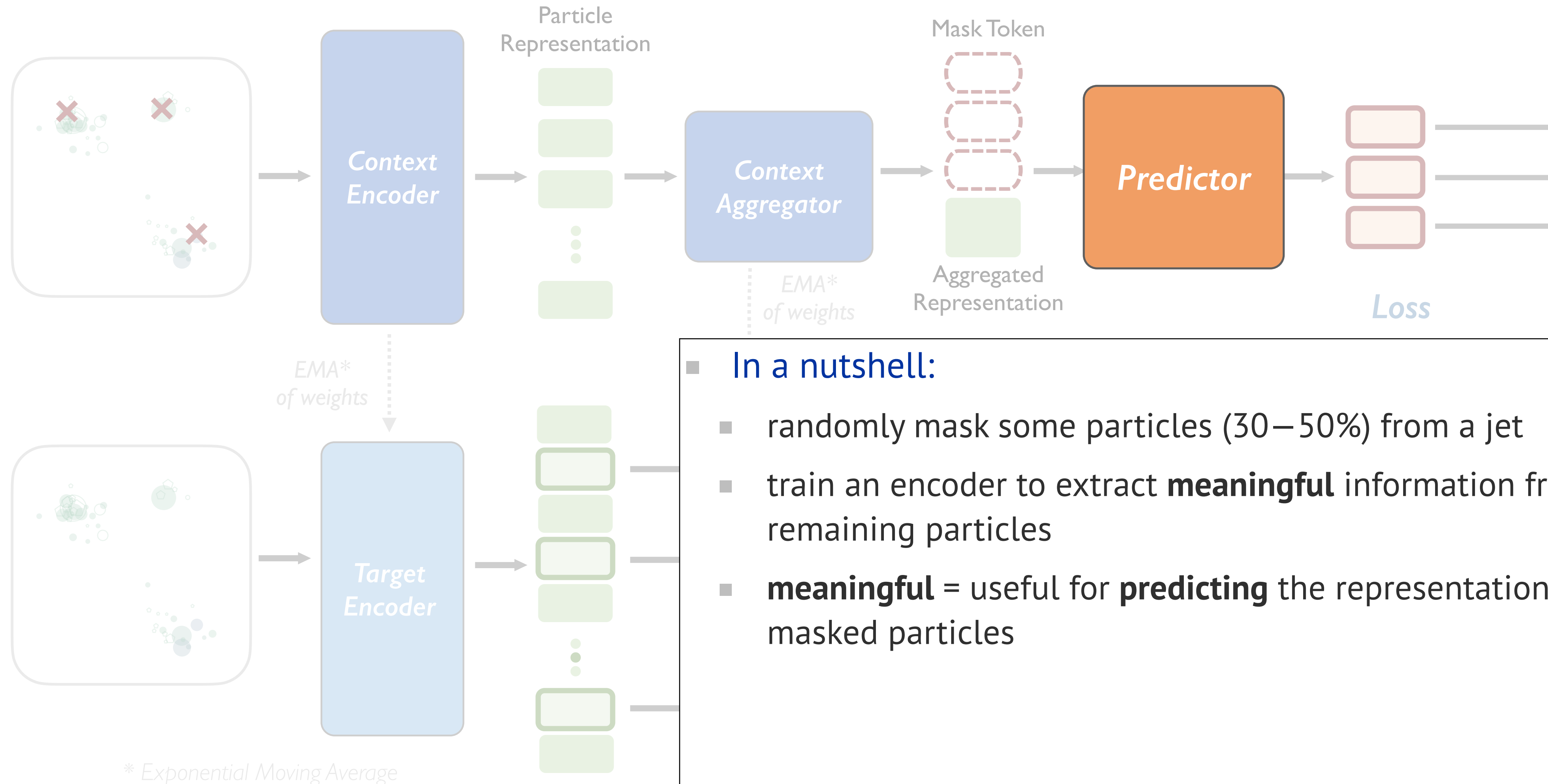


## ■ In a nutshell:

- randomly mask some particles (30–50%) from a jet
- train an encoder to extract **meaningful** information from the remaining particles



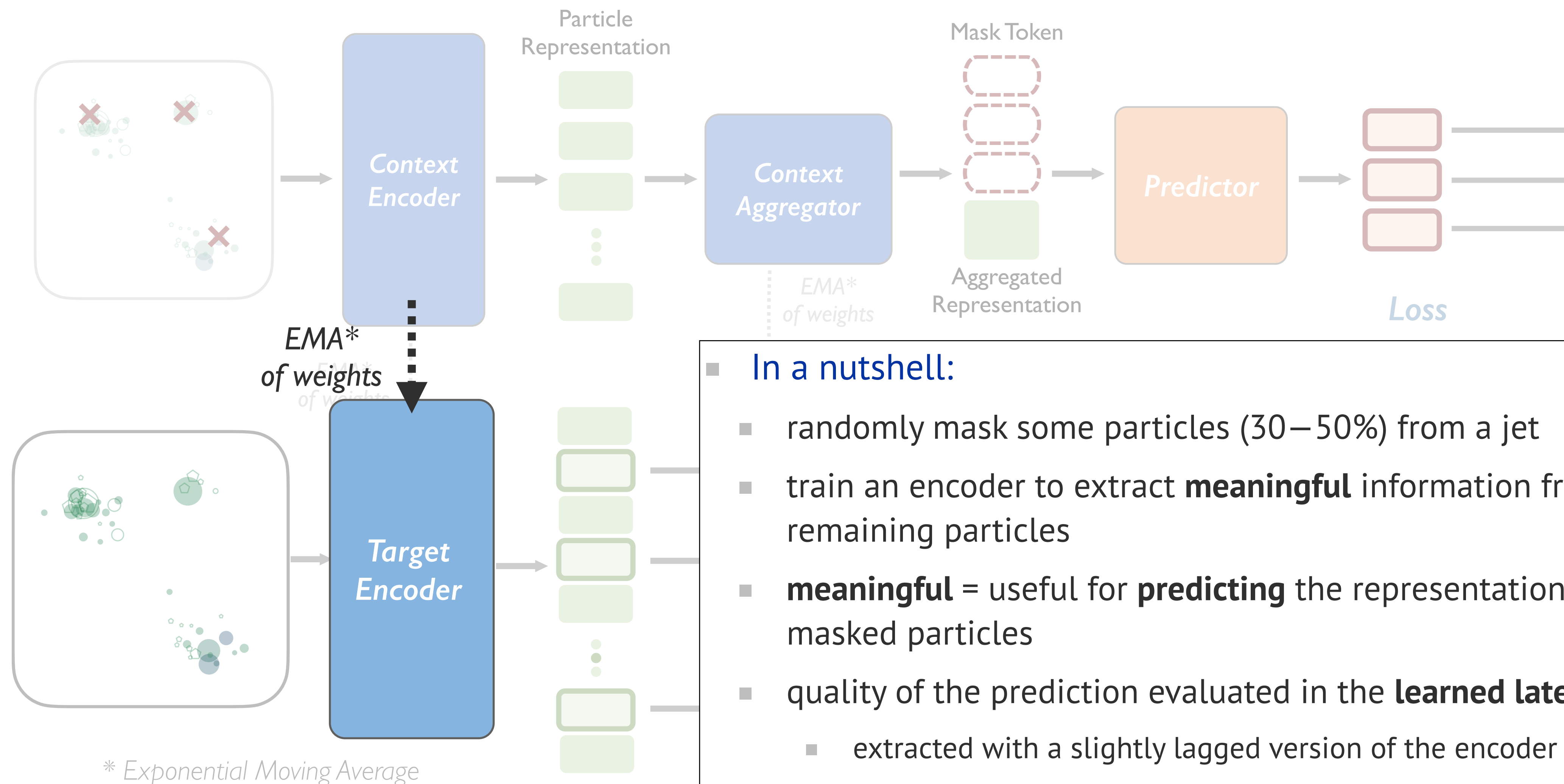
# P-JEPA



## ■ In a nutshell:

- randomly mask some particles (30–50%) from a jet
- train an encoder to extract **meaningful** information from the remaining particles
- **meaningful** = useful for **predicting** the representations of the masked particles

# P-JEPA



## ■ In a nutshell:

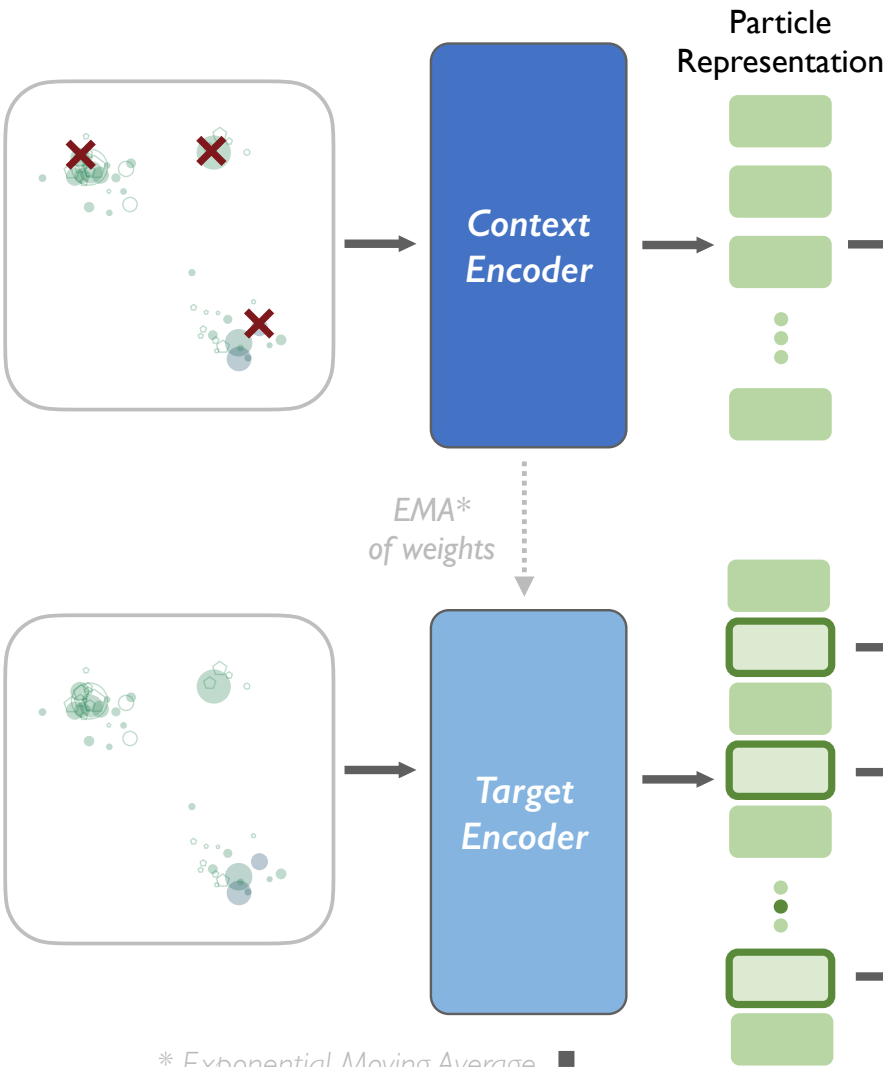
- randomly mask some particles (30–50%) from a jet
- train an encoder to extract **meaningful** information from the remaining particles
- **meaningful** = useful for **predicting** the representations of the masked particles
- quality of the prediction evaluated in the **learned latent space**
  - extracted with a slightly lagged version of the encoder



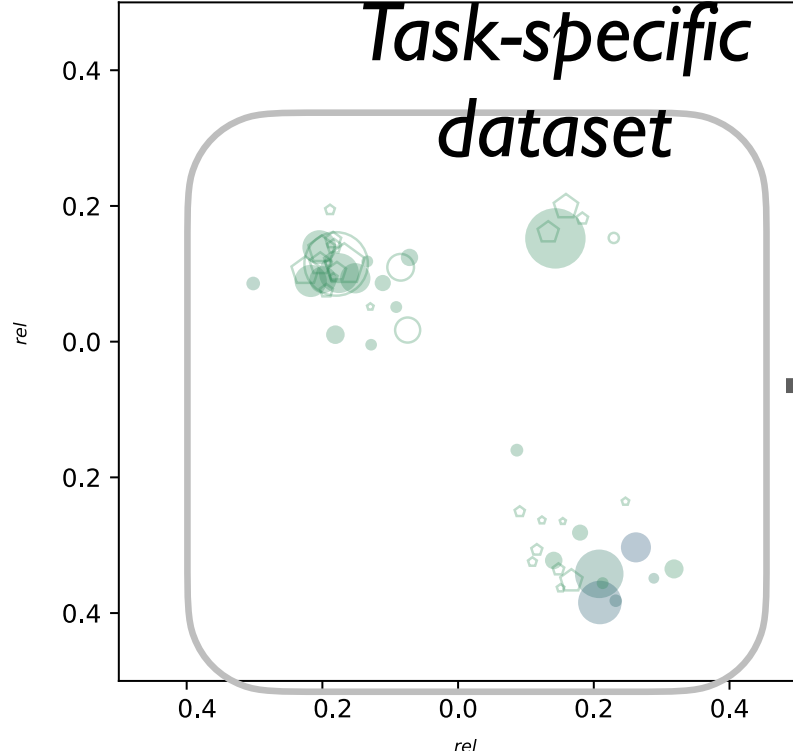
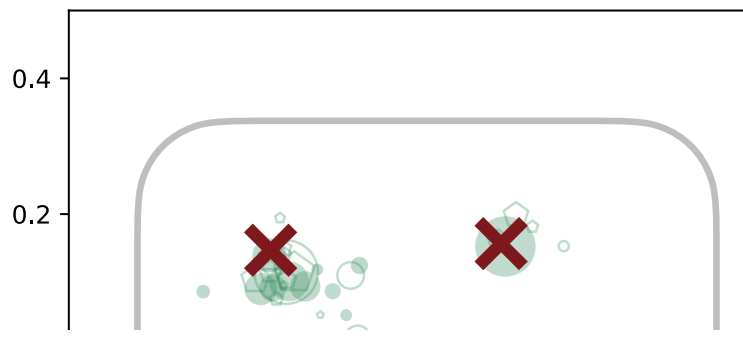
# PRE-TRAINING

- The pre-training of tl
- we demonstrate thi
- Once pre-trained, the
- transfer learning to

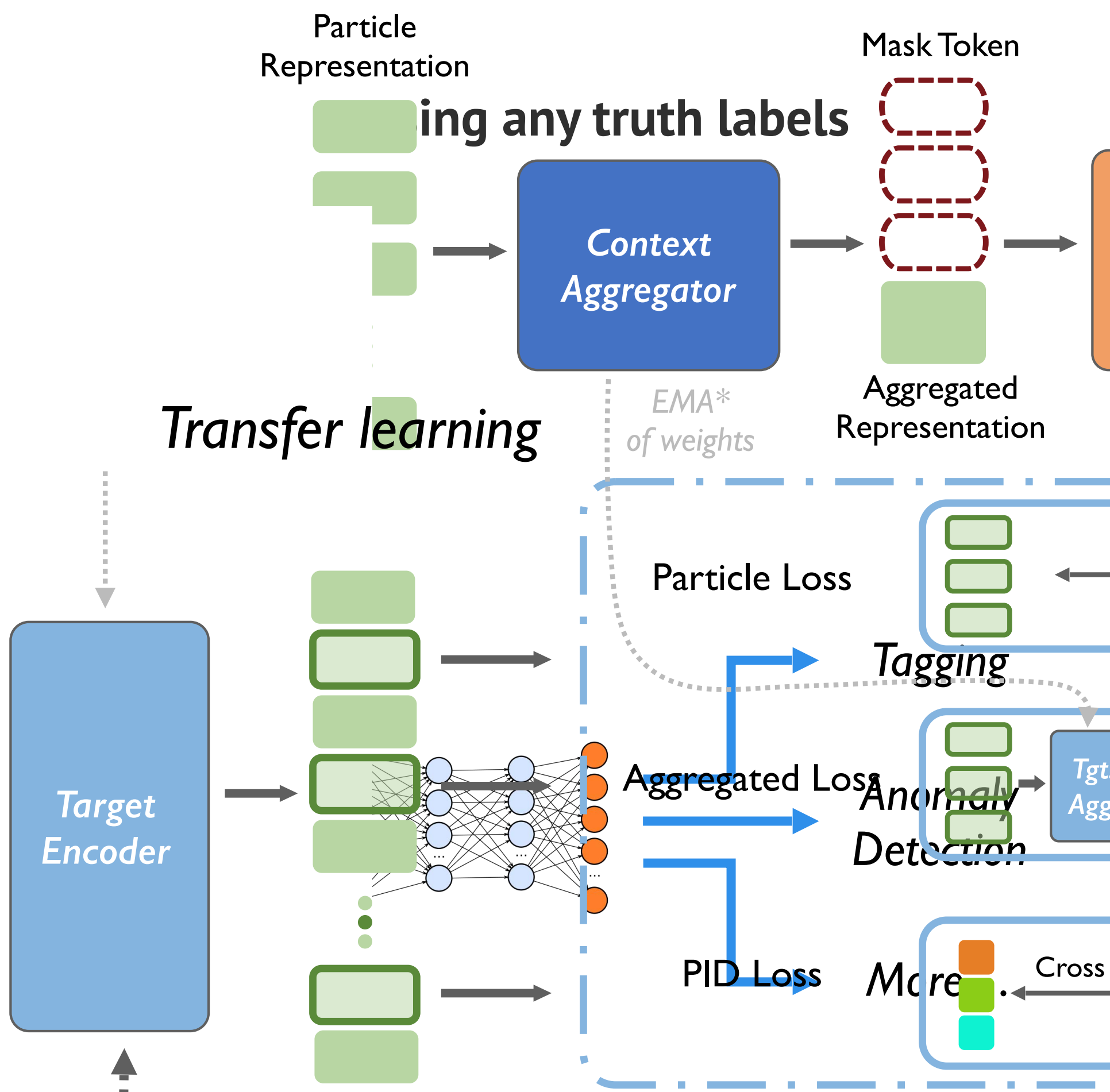
## Pre-training



\* Exponential Moving Average

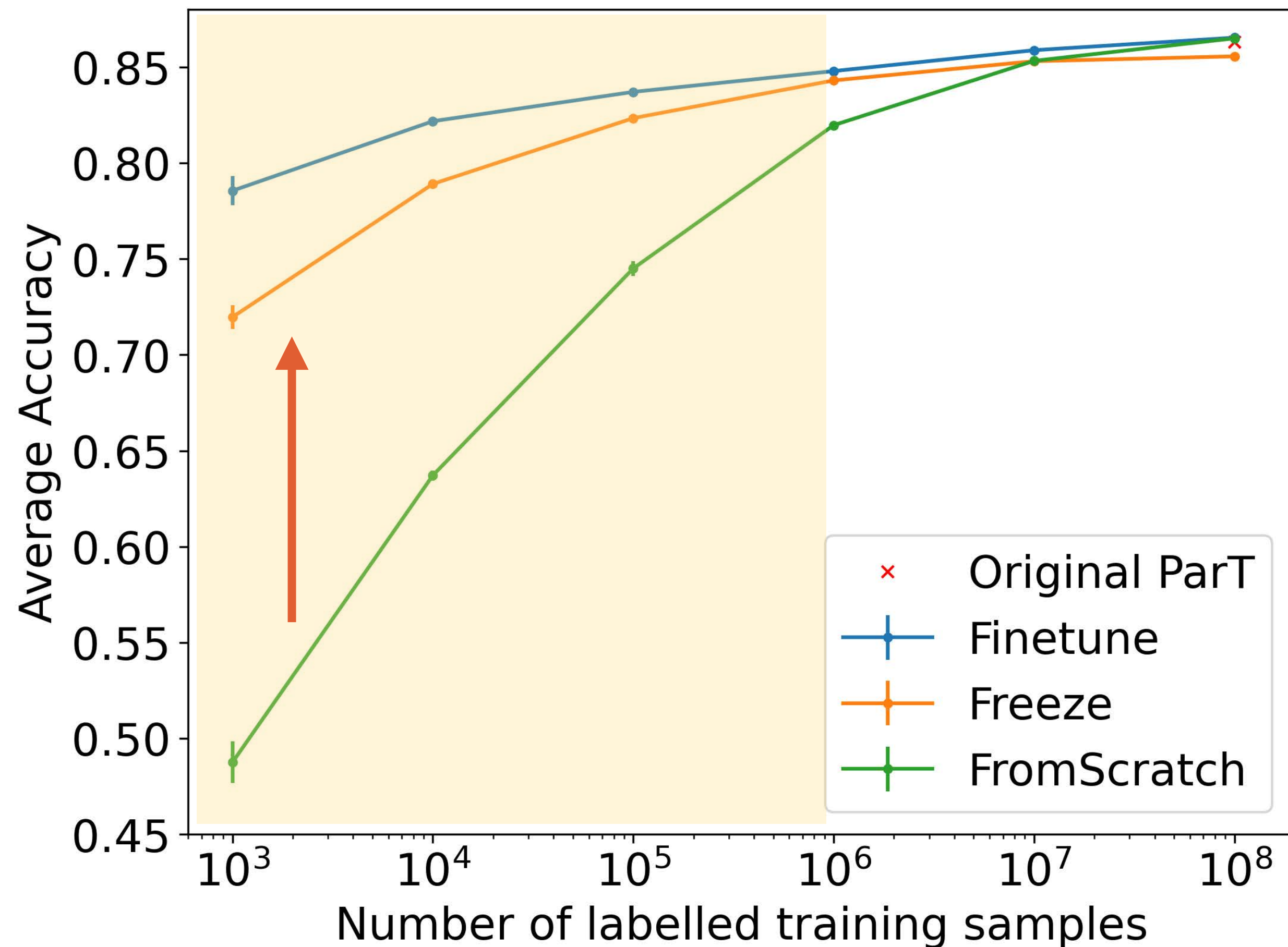


## Transfer learning



# TRANSFER LEARNING: JET TAGGING

- Benchmark: 10-class jet classification on JetClass



## FineTune:

Encoder allowed to be slightly updated when trained with labelled jets for tagging

## Freeze:

Encoder fixed when trained with labelled jets for tagging

## FromScratch:

Same network architecture, but trained with labelled jets starting from randomly initialized weights

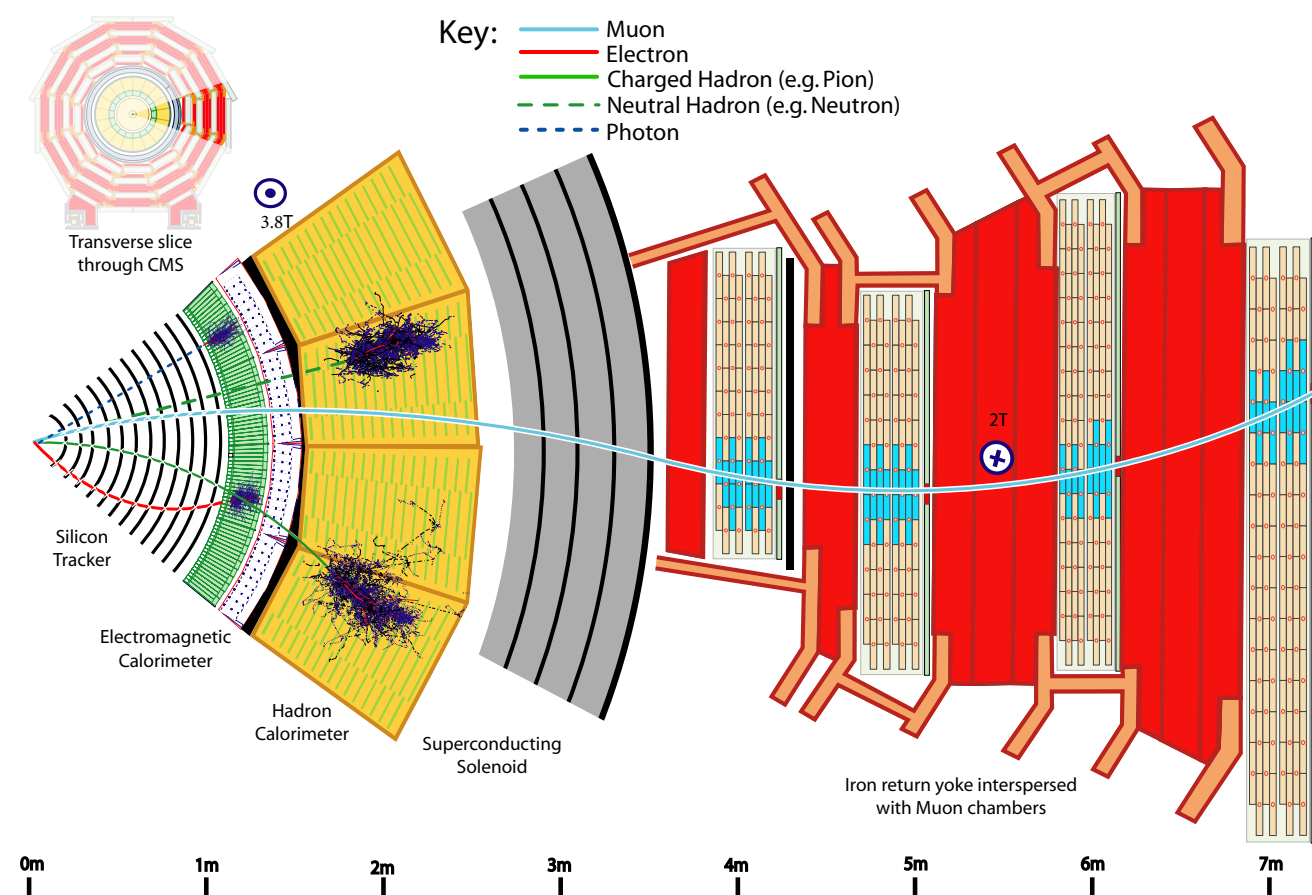
**Pre-training + transfer learning** shows a significant performance boost when **labelled samples are limited**.



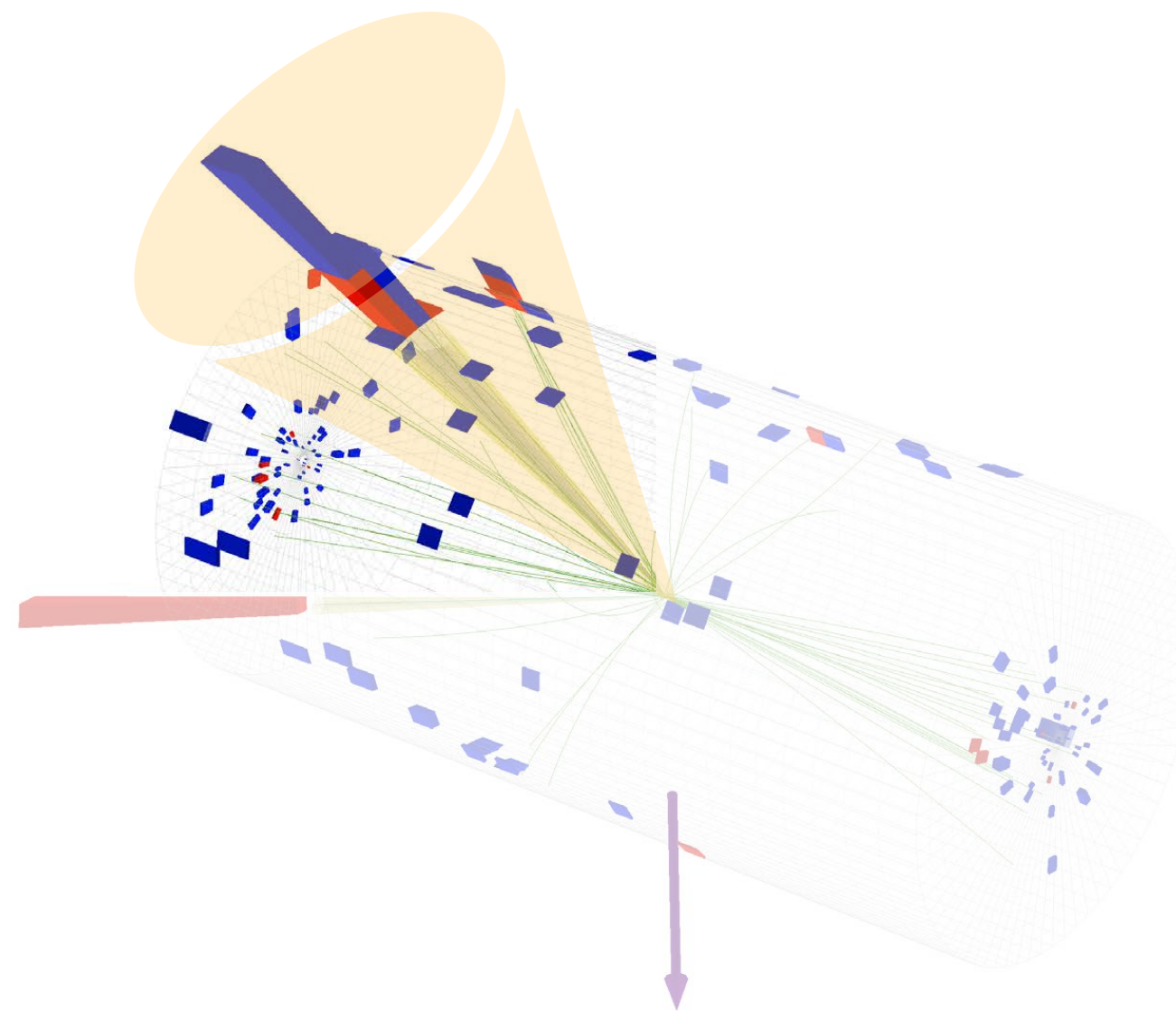
# OUTLOOK: SCALING UP

- **At colliders:** event reconstruction usually follows a hierarchical way
  - due to heterogeneous sub-detectors and high particle multiplicity
- **Going beyond: a foundation model that aligns data representations at various levels and cover all the stages?**
  - requires a multi-modal approach and efficient architectures

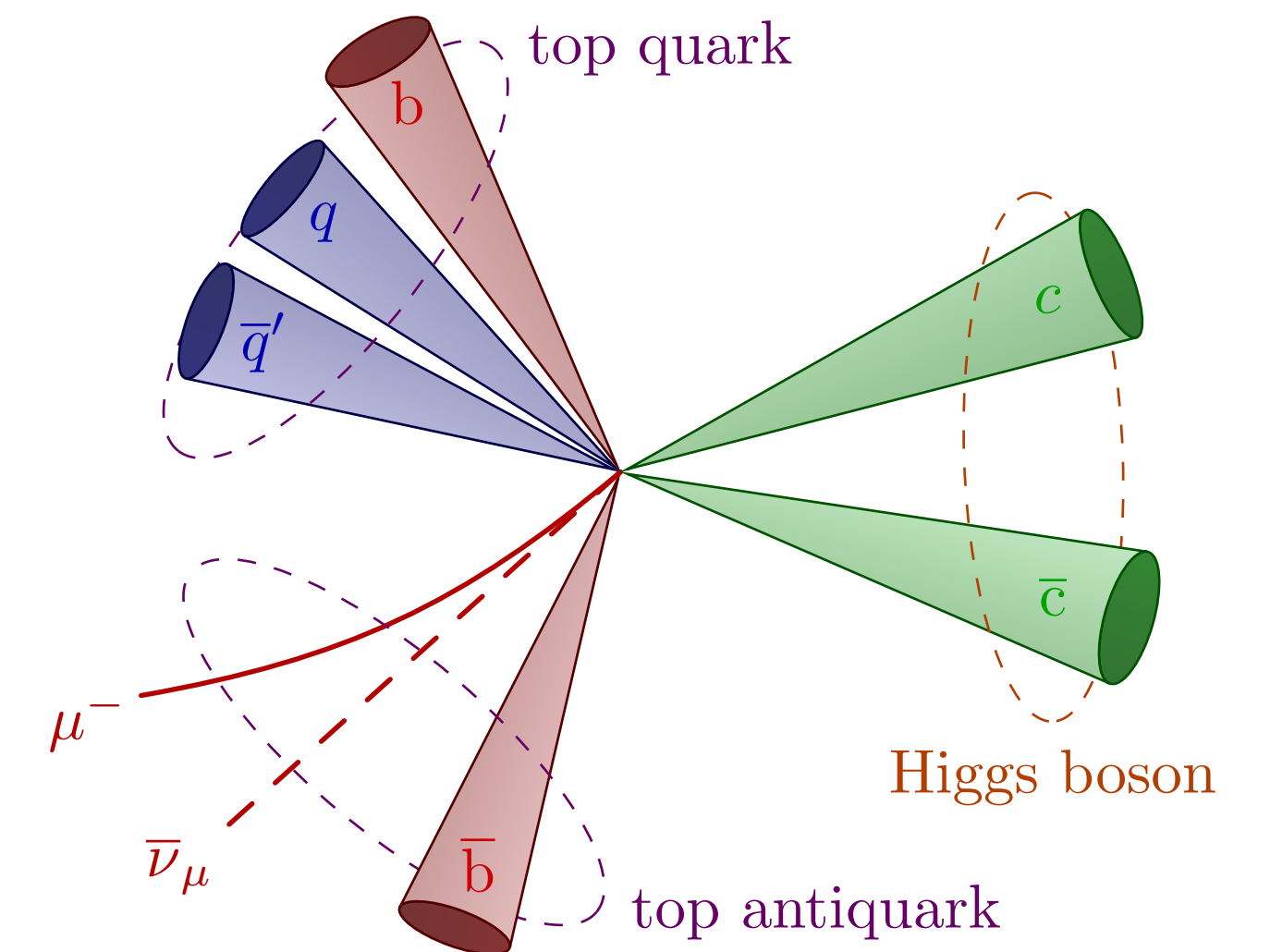
## Detector/Particle-flow reconstruction



## Jet/Object reconstruction



## Event classification/interpretation



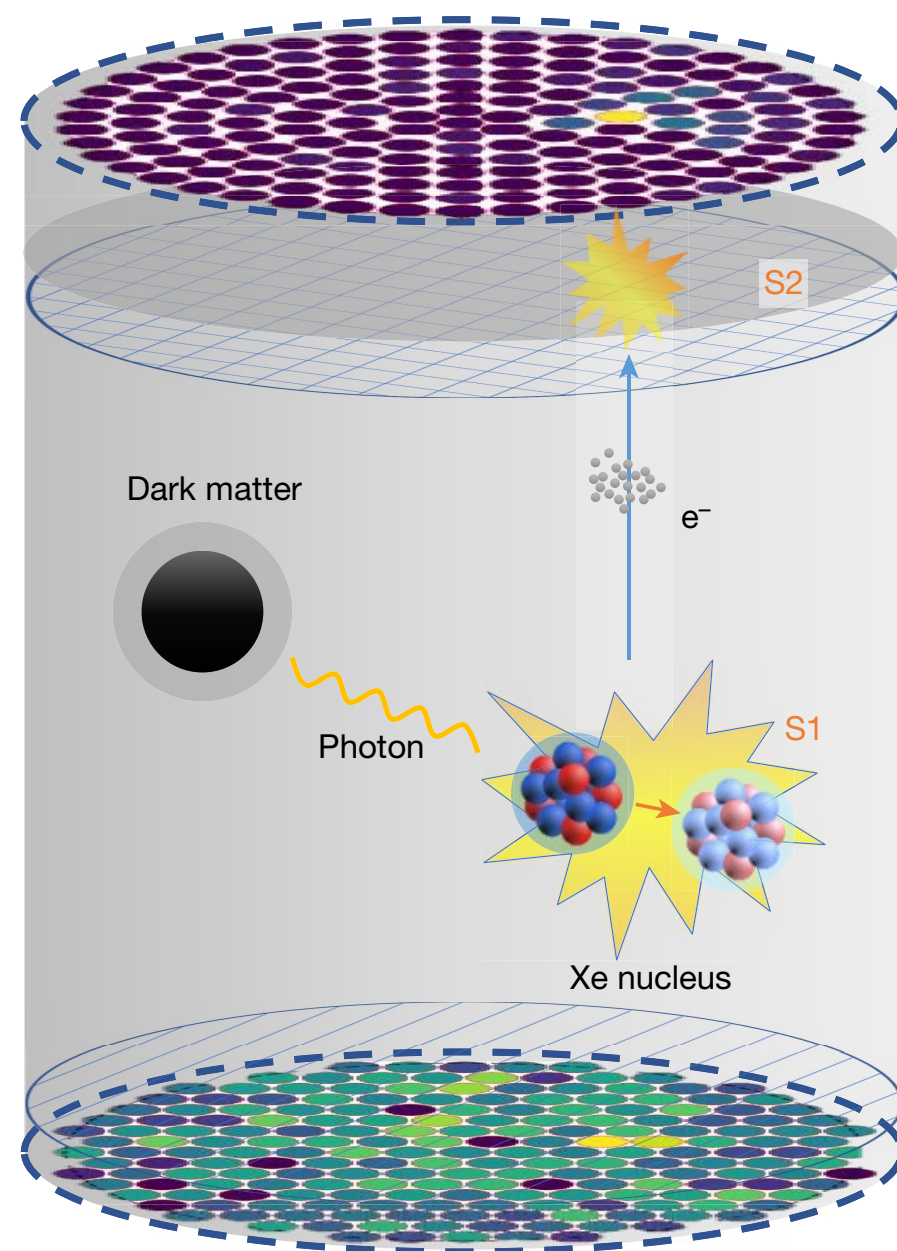
Raw detector hits => Tracks, Showers, Particles => Jets, leptons, taus => Signal vs background processes



# OUTLOOK: SCALING OUT

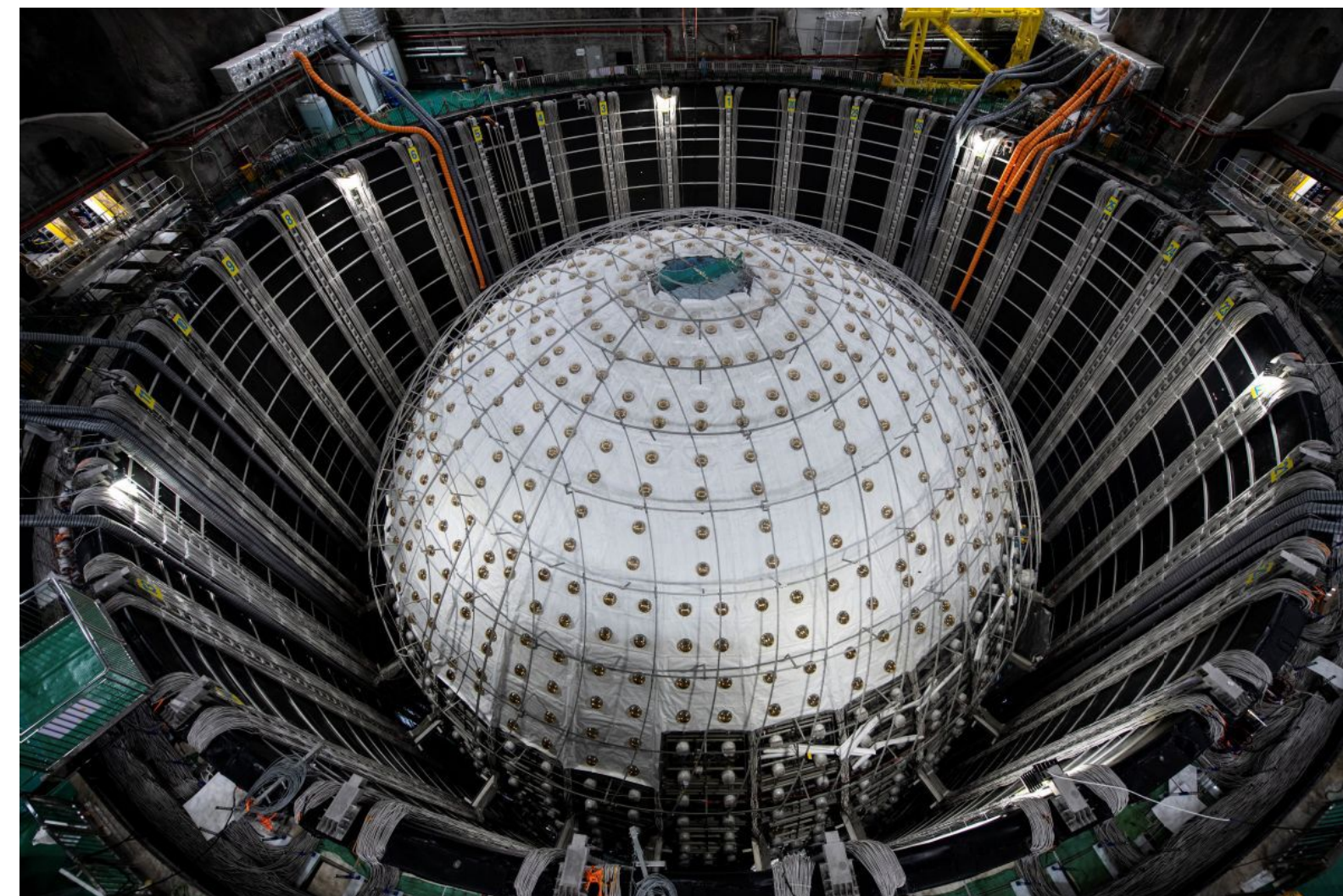
- **Dark matter/neutrino/cosmic ray experiments**
  - typically homogeneous detector — end-to-end ML reconstruction more promising
- **Going beyond: a foundation model that generalizes across experiments and covers all energy scales?**

PandaX



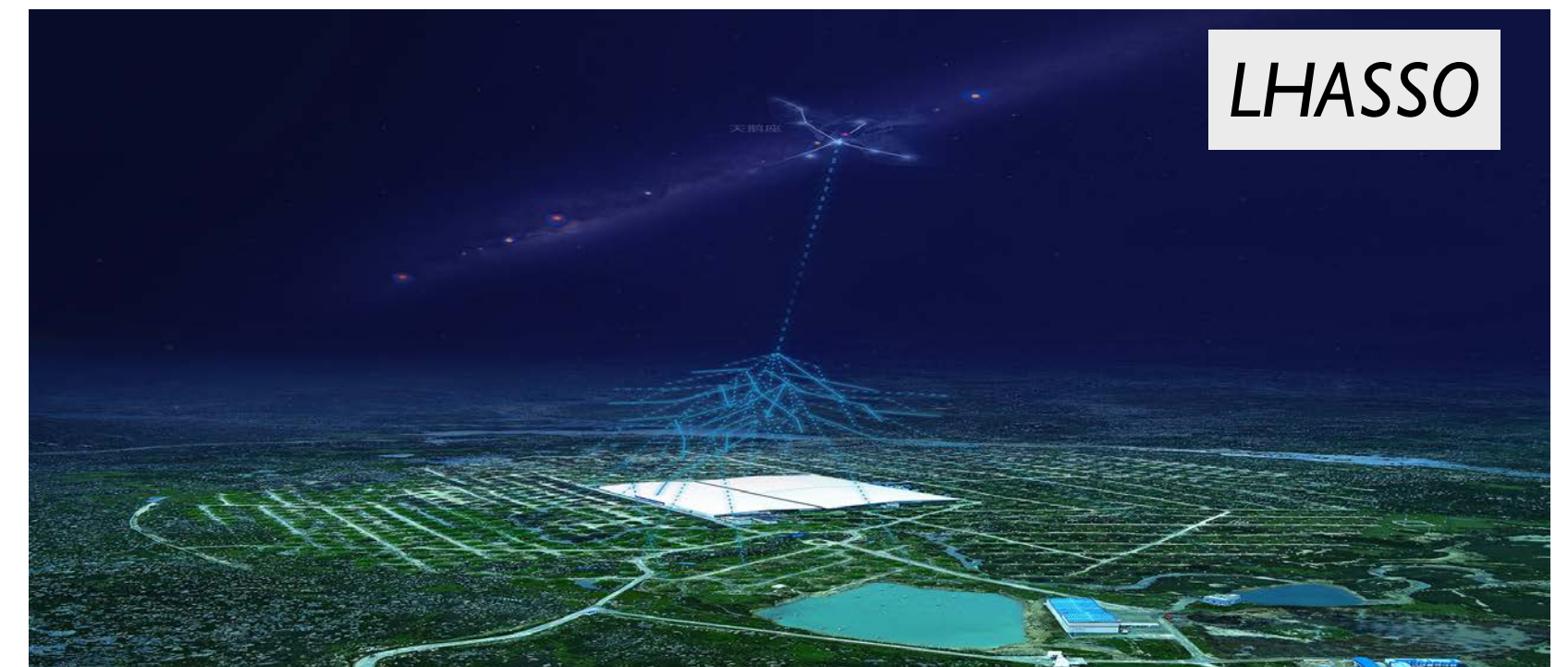
Dark Matter/Neutrino  
(keV~MeV)

JUNO

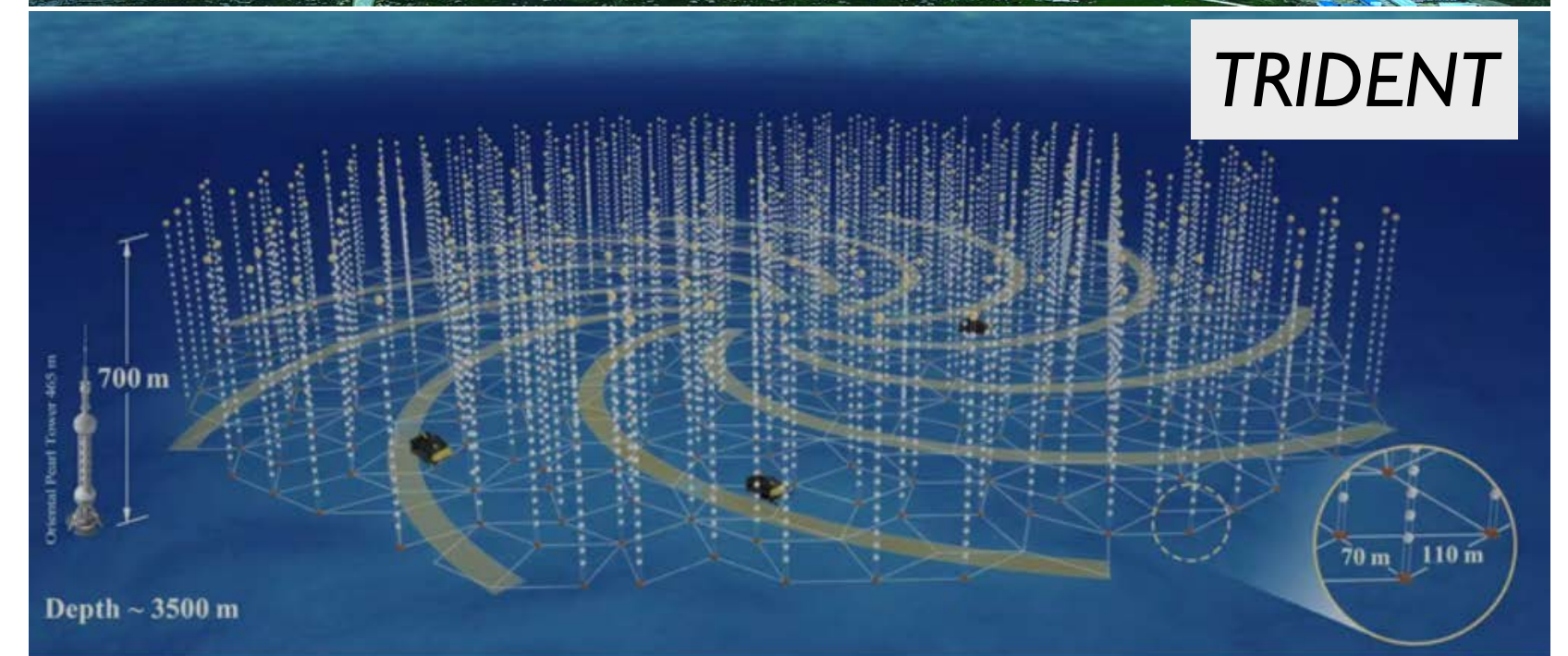


Neutrino  
(MeV~GeV)

LHASSO



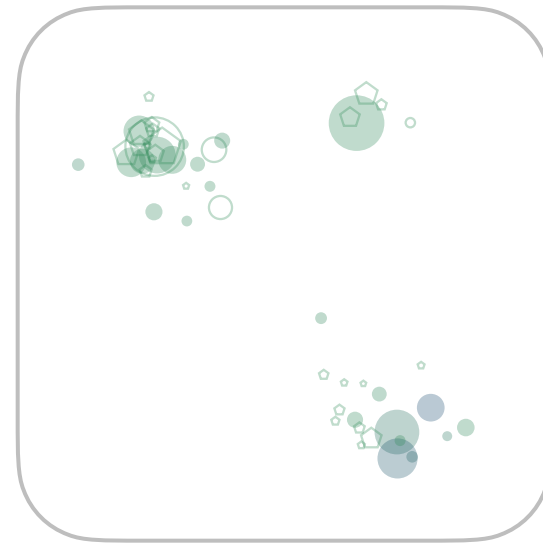
TRIDENT



Cosmic Rays/Neutrinos  
(TeV~EeV)

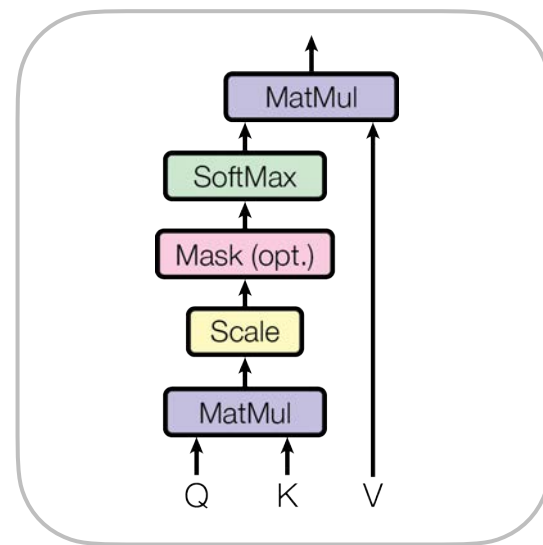


# SUMMARY



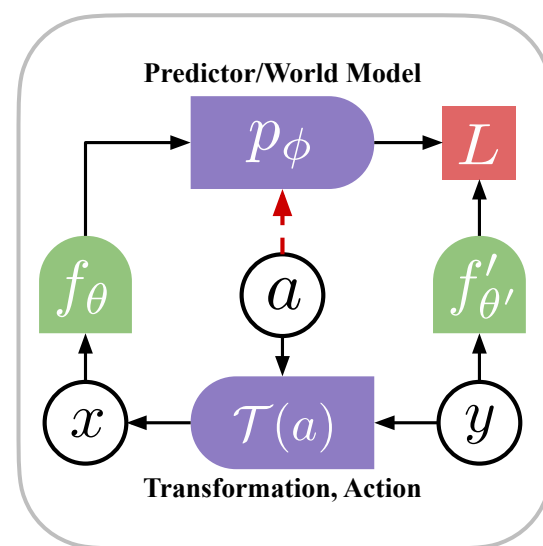
## Evolution of Jet Representations

*... leads to significant performance gains.*



## Attention Isn't All You Need

*... more physics = fast learning & better generalization.*



## Towards a Foundation Model

*... efficient and scalable representation learning is the key!*