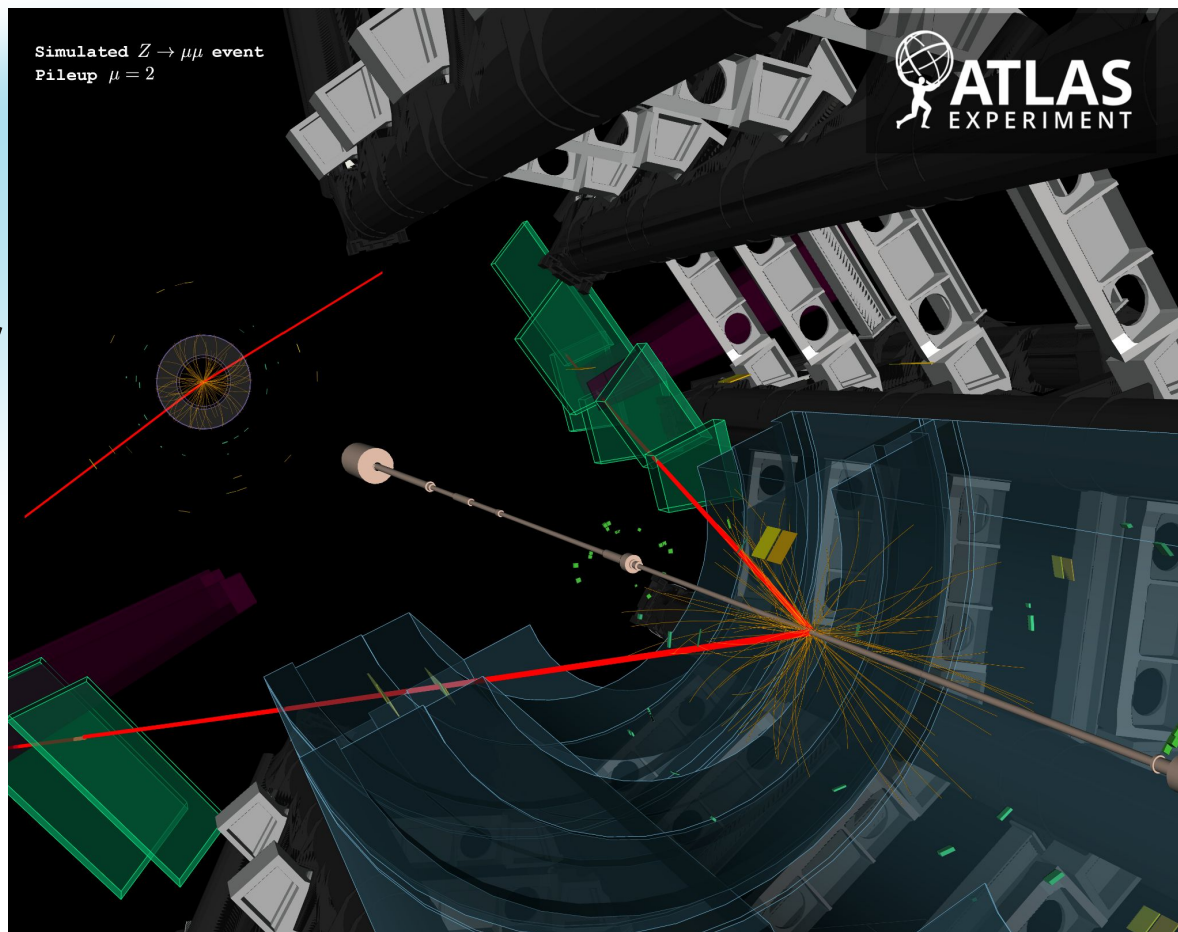




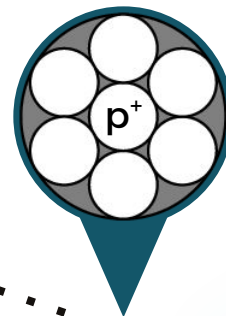
Nathan Suri, Vinicius Mikuni, Ben Nachman
ML4Jets 2026

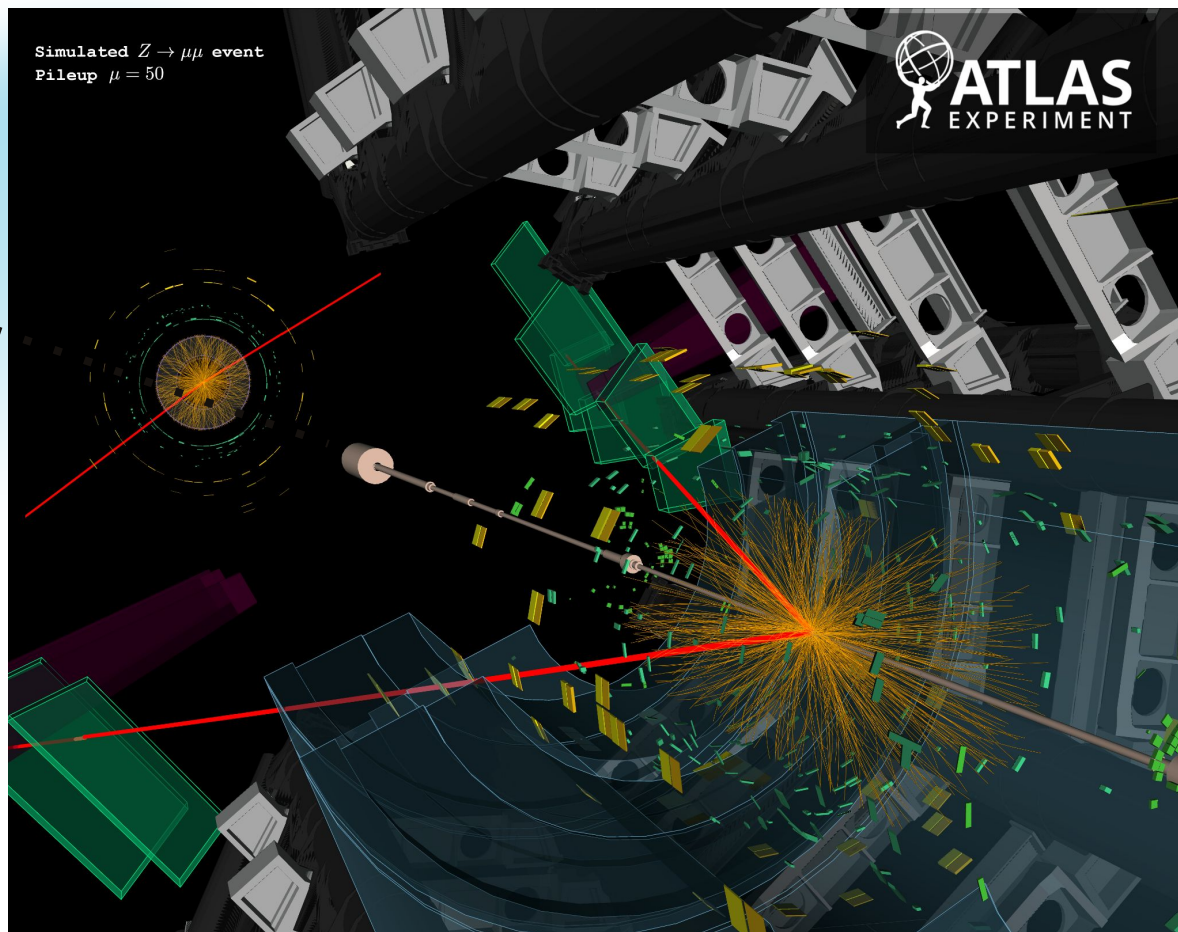
Optimizing Optimal Transport-based Pileup Mitigation



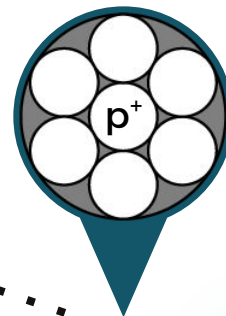


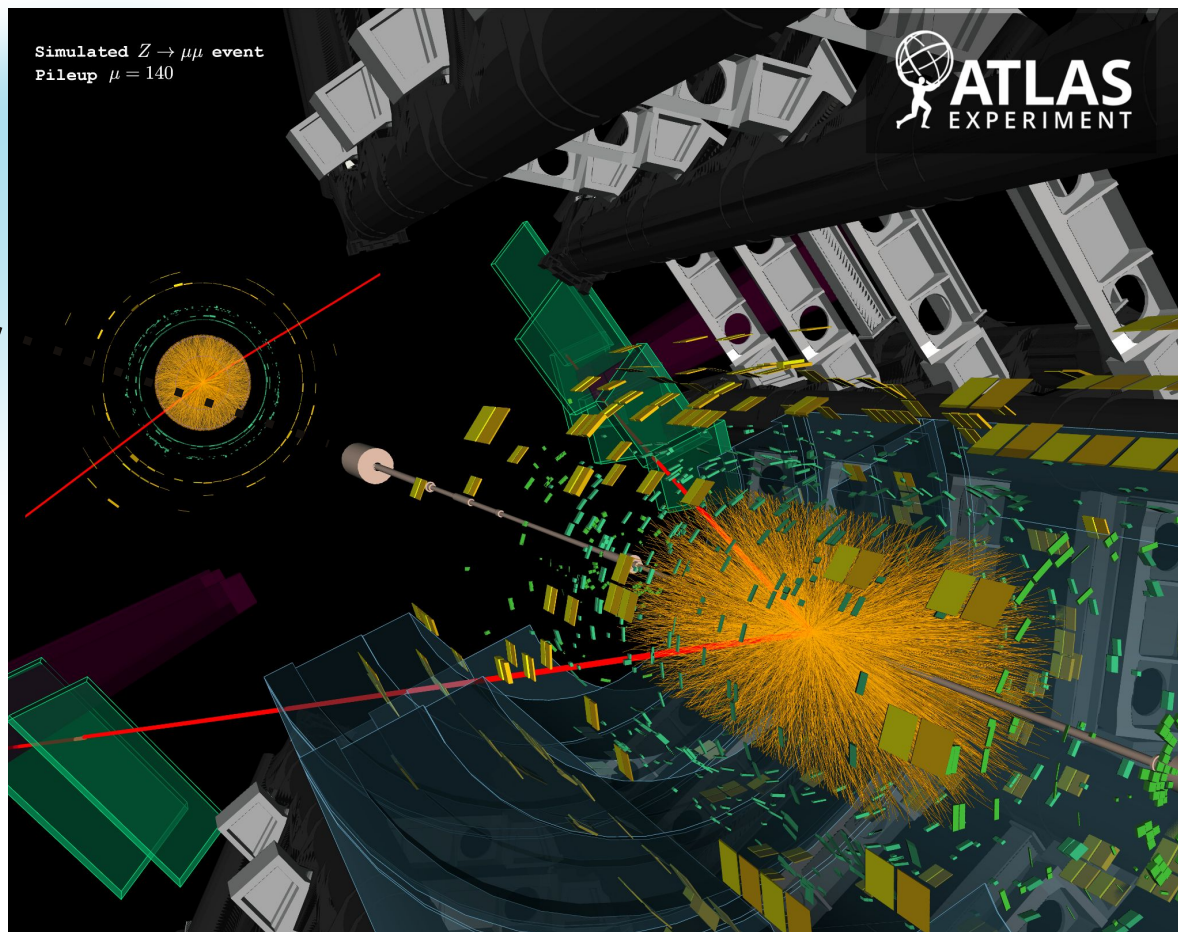
- Pileup Descriptors
1. Charged versus neutral
 2. In-time versus out-of-time



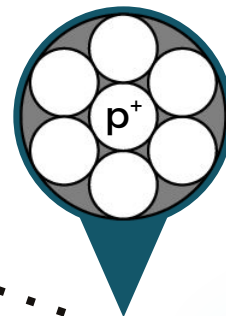


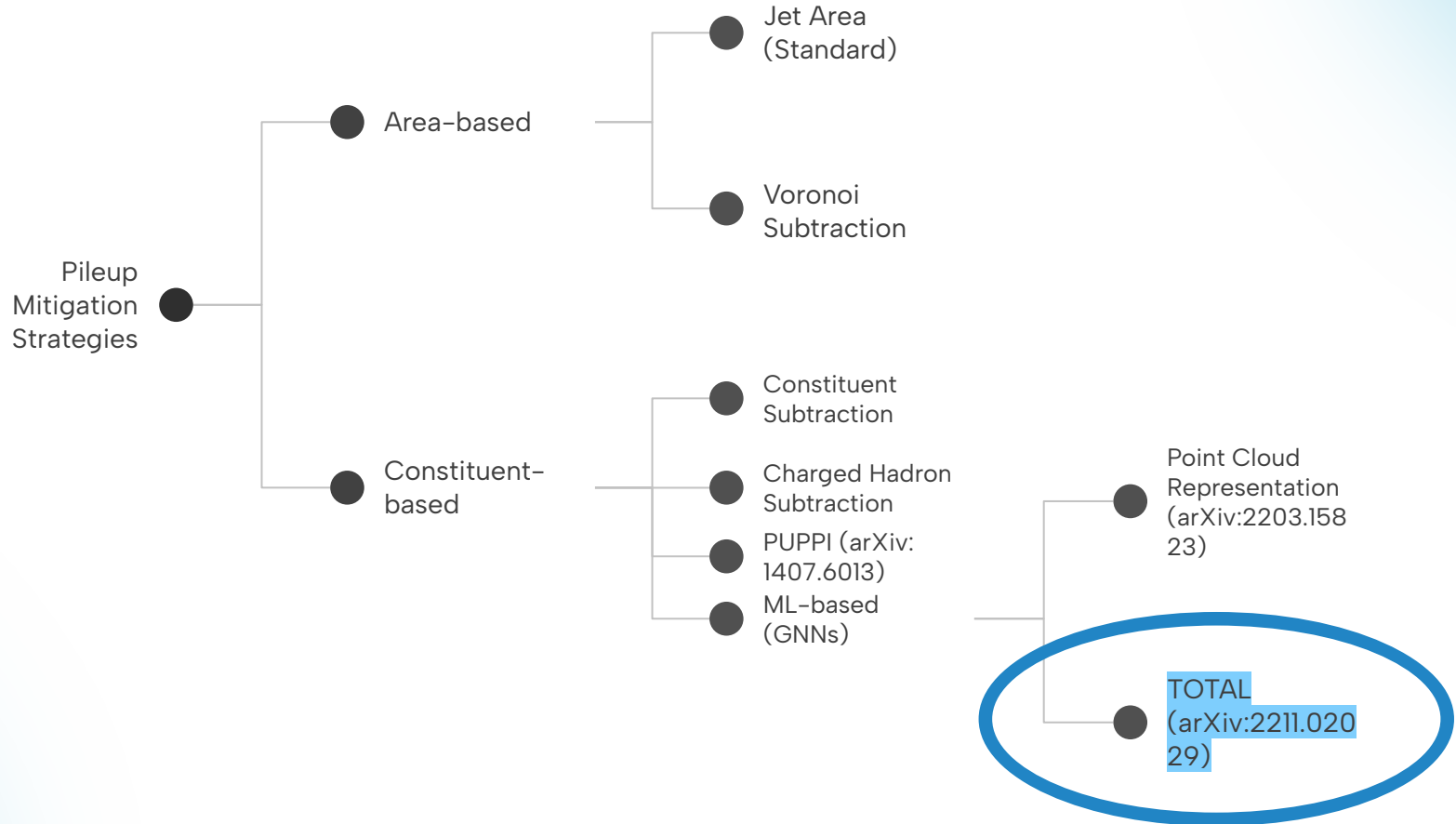
- Pileup Descriptors
1. Charged versus neutral
 2. In-time versus out-of-time





- Pileup Descriptors
1. Charged versus neutral
 2. In-time versus out-of-time

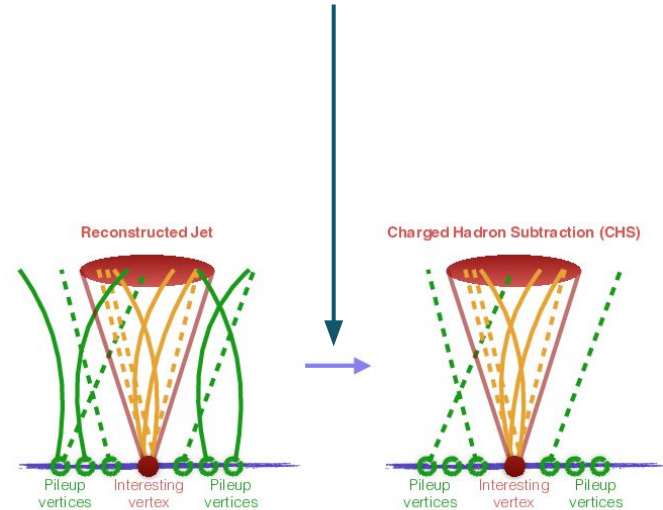




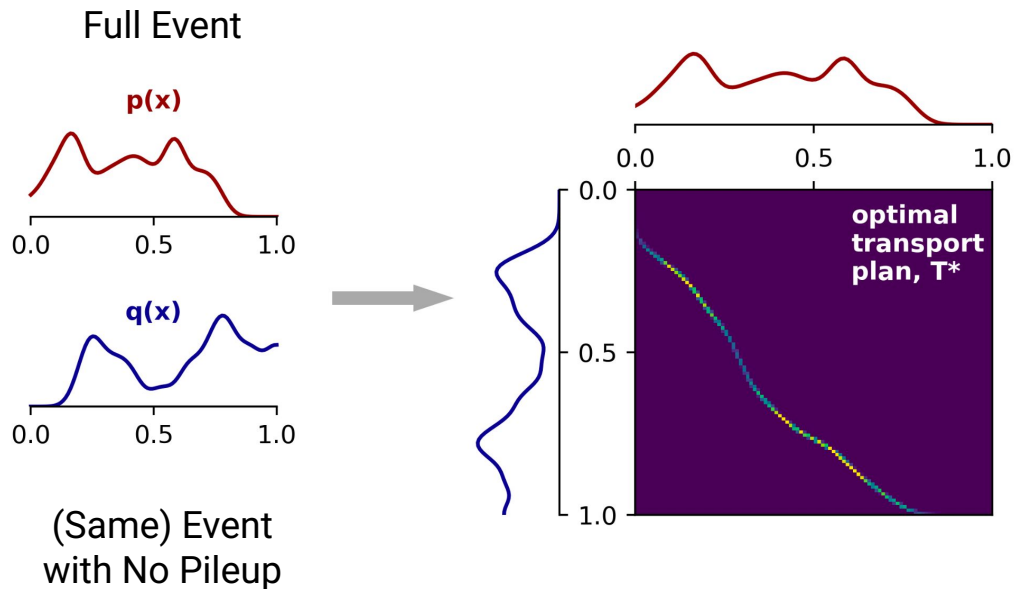
Charged Hadron Subtraction

- Benefits
 - Very effective at removing charged pileup due to track information
- Drawbacks
 - Inapplicable to neutral pileup
- (arXiv:2012.06271)

$$\text{CVF} = \frac{\sum_{\text{tracks,HS}} p_T^{\text{track}}}{\sum_{\text{tracks,HS}} p_T^{\text{track}} + \sum_{\text{tracks,PU}} p_T^{\text{track}}},$$



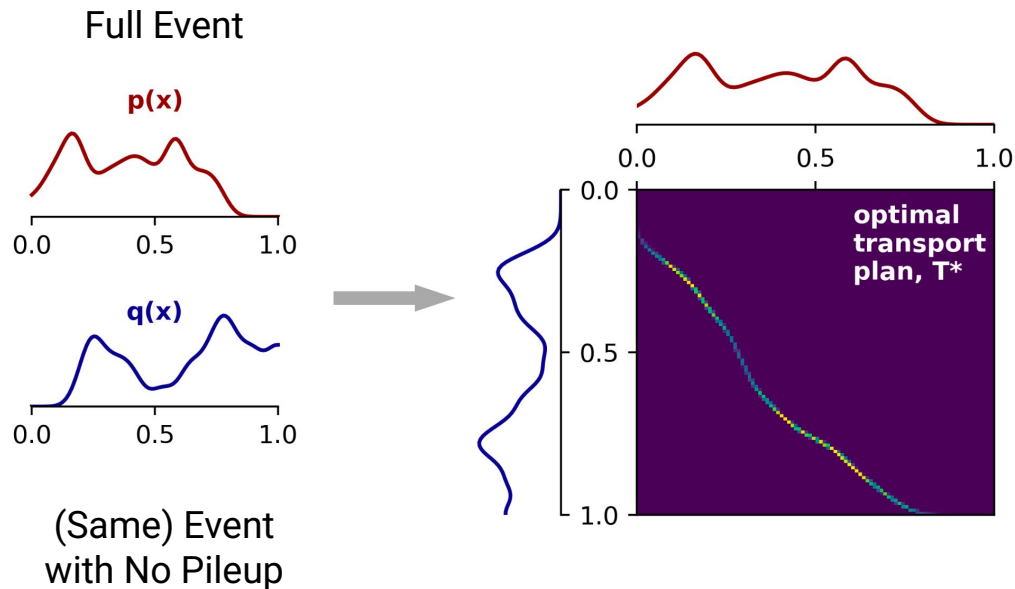
TOTAL: Training Optimal Transport with Attention Learning



$$\mathcal{L} = \text{SWD}(x'_p, x_{np}) + \lambda \text{MSE}(p_T^{\text{miss}}(x'_p), p_T^{\text{miss}}(x_{np}))$$

TOTAL: Training Optimal Transport with Attention Learning

- The probability density is intractable, but we can approximate the density
- Realizations of the density are accessible
- Optimal transport over the space of inputs allows for approximation



$$\mathcal{L} = \text{SWD}(x'_p, x_{np}) + \lambda \text{MSE}(p_T^{\text{miss}}(x'_p), p_T^{\text{miss}}(x_{np}))$$

TOTAL: Training Optimal Transport with Attention Learning

$$\mathcal{L} = \text{SWD}(x'_p, x_{np}) + \lambda \text{MSE}(p_T^{\text{miss}}(x'_p), p_T^{\text{miss}}(x_{np}))$$

TOTAL: Training Optimal Transport with Attention Learning

$$\mathcal{L} = \text{SWD}(x'_p, x_{np})$$

- Wasserstein distance (WD): Finds the transport function that keeps hard scattering particles and removes those from simultaneous vertices
- Sliced WD to compensate for poor scaling of computational costs of calculating WD at high dimensions

$$x'_p = \omega x_p$$

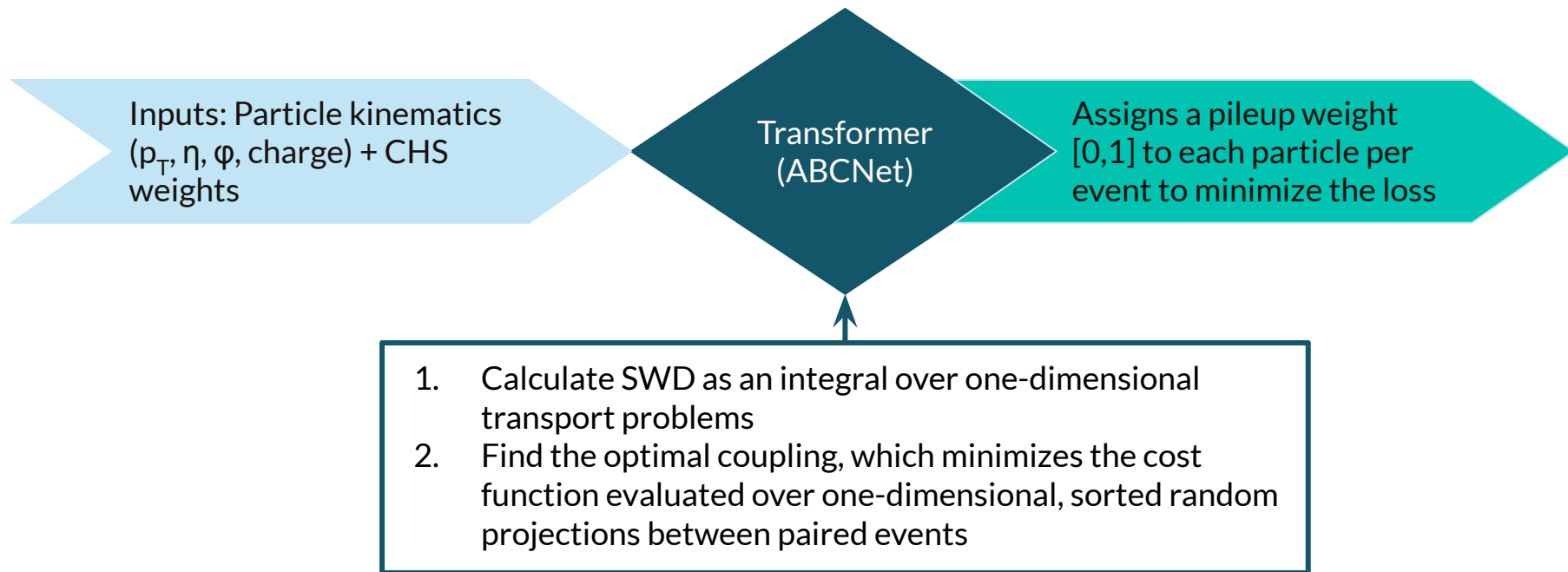

TOTAL: Training Optimal Transport with Attention Learning

$\mathcal{L} =$

- Scaled Mean Square Error of missing p_T
- Forces energy conservation between the pure and full samples

$$+ \lambda \text{MSE}(p_T^{\text{miss}}(x'_p), p_T^{\text{miss}}(x_{np}))$$

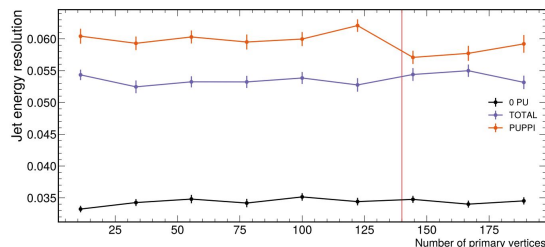
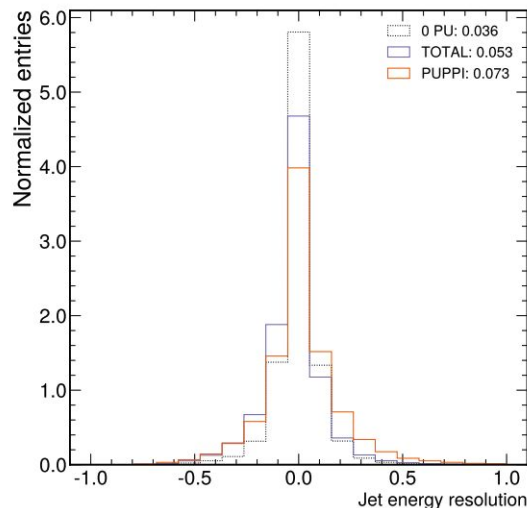
TOTAL: Training Optimal Transport with Attention Learning



TOTAL: Training Optimal Transport with Attention Learning

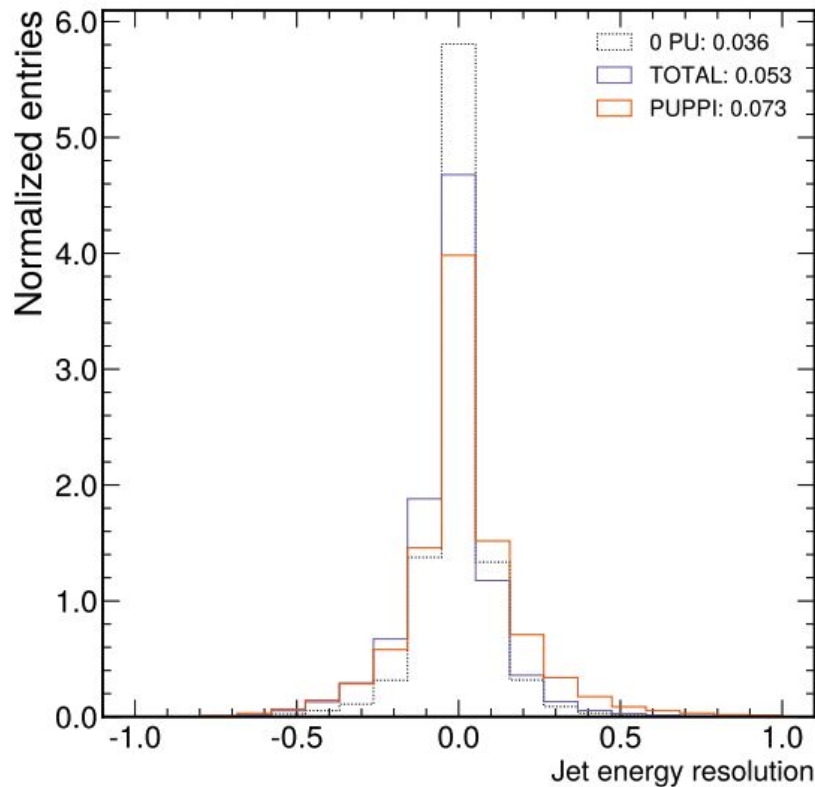
- + Outperforms traditional and ML-based alternatives
- + Relies on global event descriptions without MC corrections
- + Robustly learns pileup characteristics as a transport function

- Requires direct matching of events
- Limited due to supervision

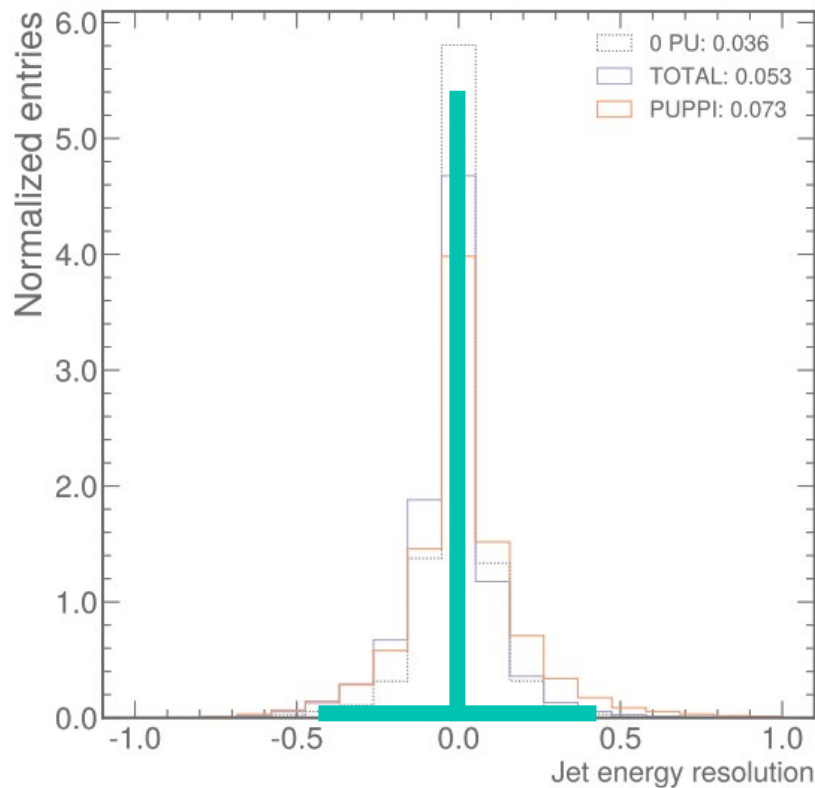


(arXiv:2211.02029)

TOTAL: Training Optimal Transport with Attention Learning

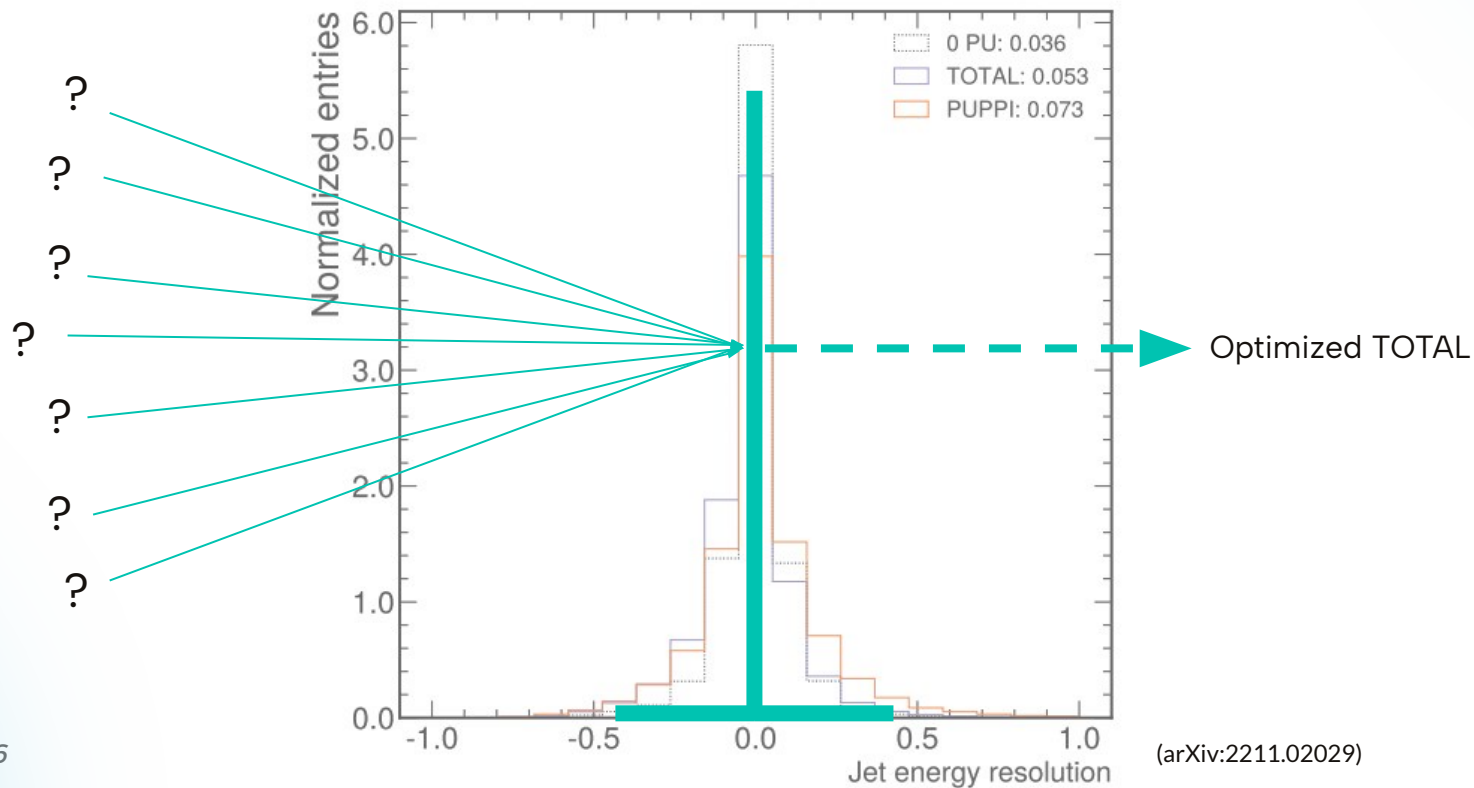


TOTAL: Training Optimal Transport with Attention Learning



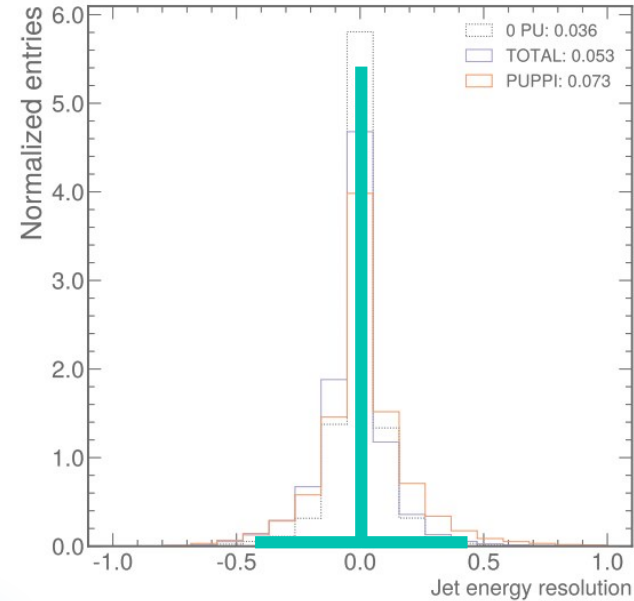
(arXiv:2211.02029)

TOTAL: Training Optimal Transport with Attention Learning



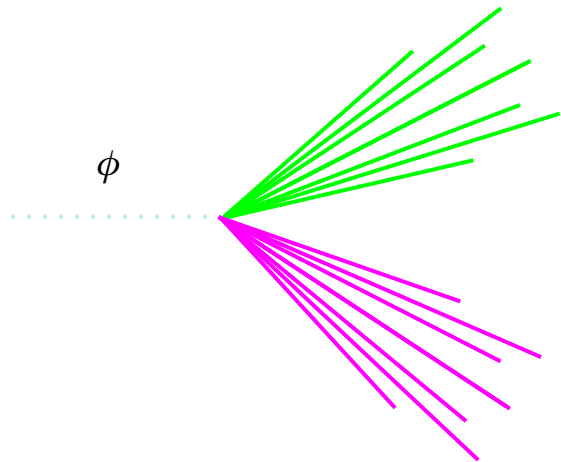
TOTAL Upgrades

1. Physics-based constraints
2. Alternative loss functions



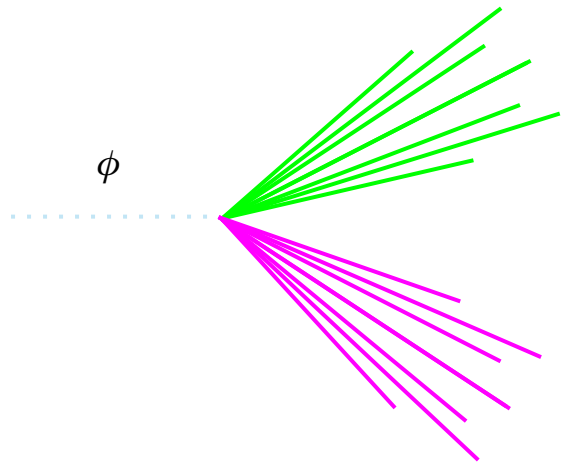
(arXiv:2211.02029)

Physics Example: High p_T Jets



- PUMML Dataset:
<https://zenodo.org/records/2652034>
- Process: q - q bar
light-quark-initiated jets from the
from the decay of a Higgs-like
scalar
 - Pileup was generated by
overlaying soft QCD on top
of signal

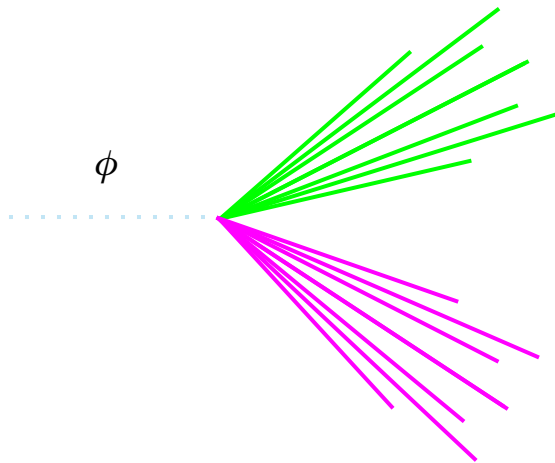
Physics Example: High p_T Jets

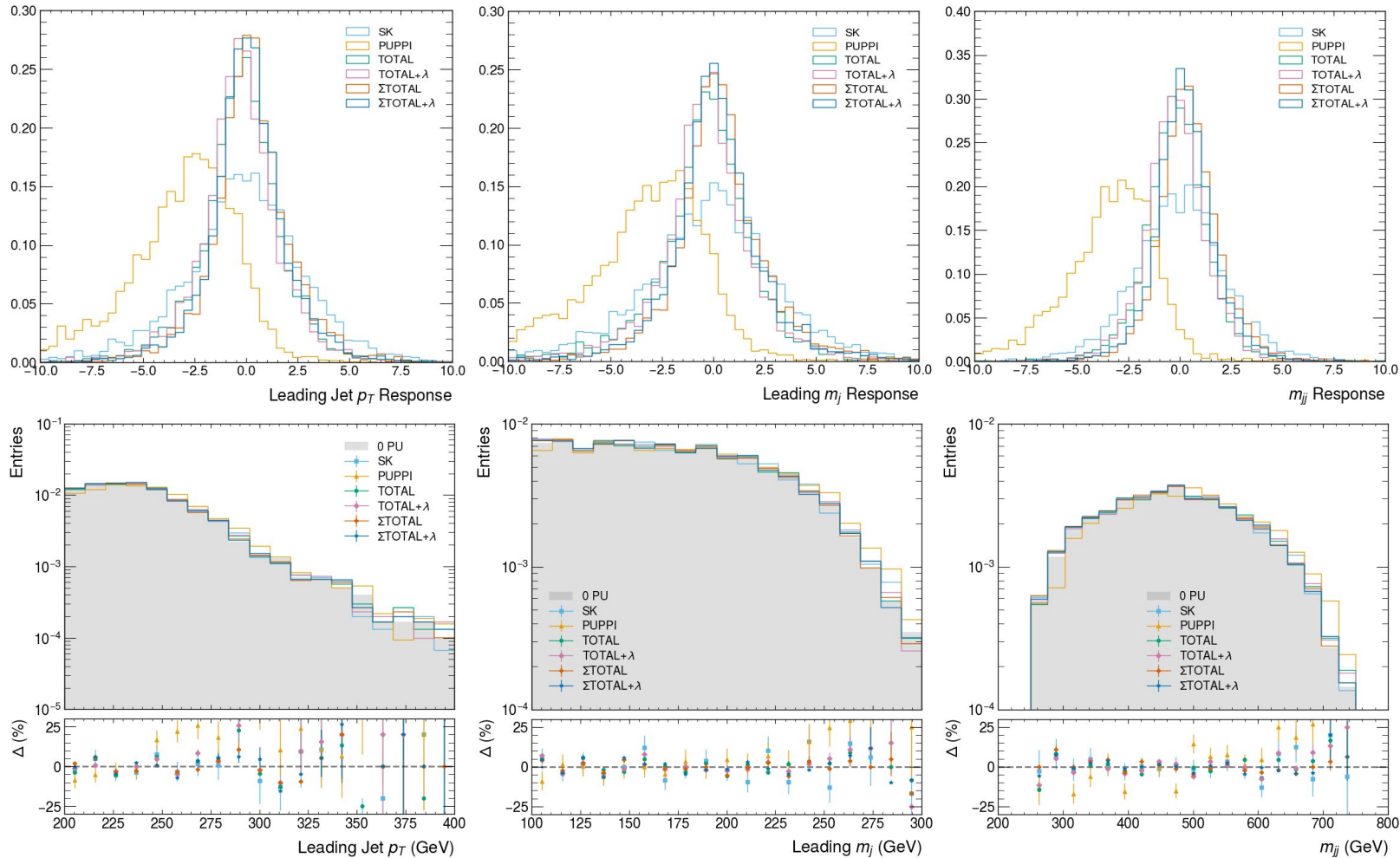


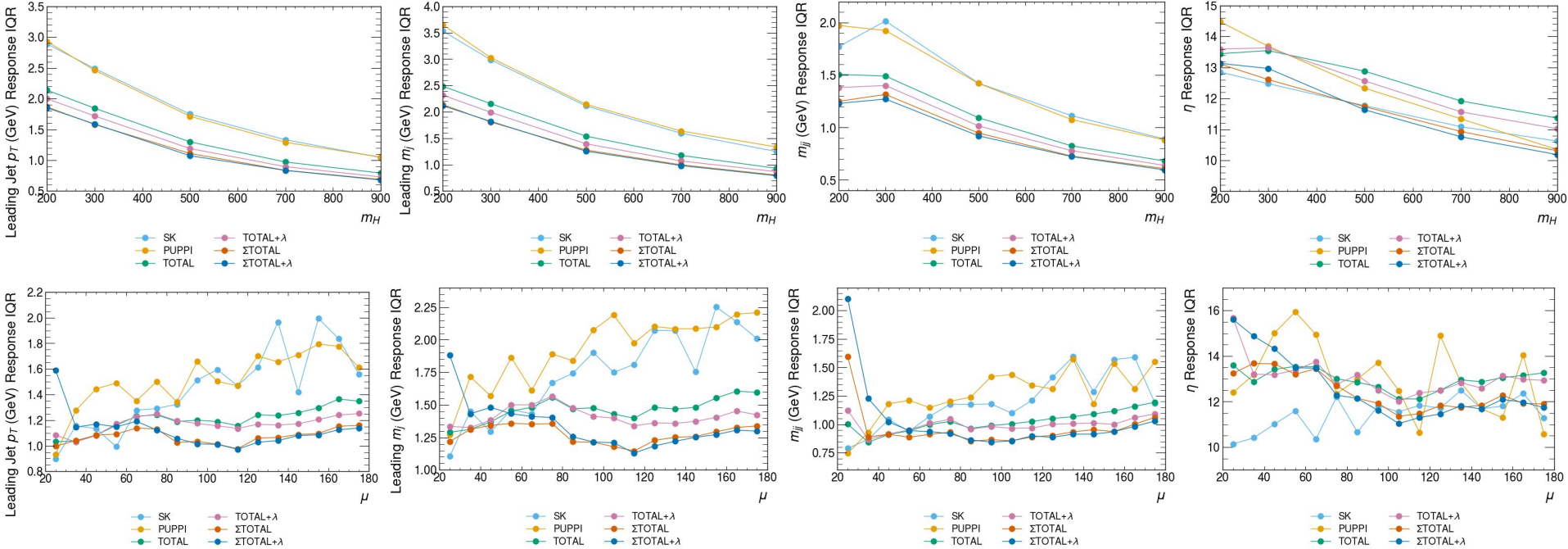
- PUMML Dataset:
<https://zenodo.org/records/2652034>
- Datasets
 - $\mu = 140$, Δm : Set pileup vertex count, varied scalar mass
 - $\Delta\mu$, $m = 500$ GeV: Varied pileup vertex (PV) count, set scalar mass

Physics-based Constraints

1. **Event-level constraints (λ):** Add MET and H_T constraints to the loss to control and regularize event-level energy scales
2. **Constituent splitting (Σ):** Sum two separate losses for charged and neutral constituents
3. **Activation function*:** Adjust from using a hard sigmoid to a sigmoid activation function

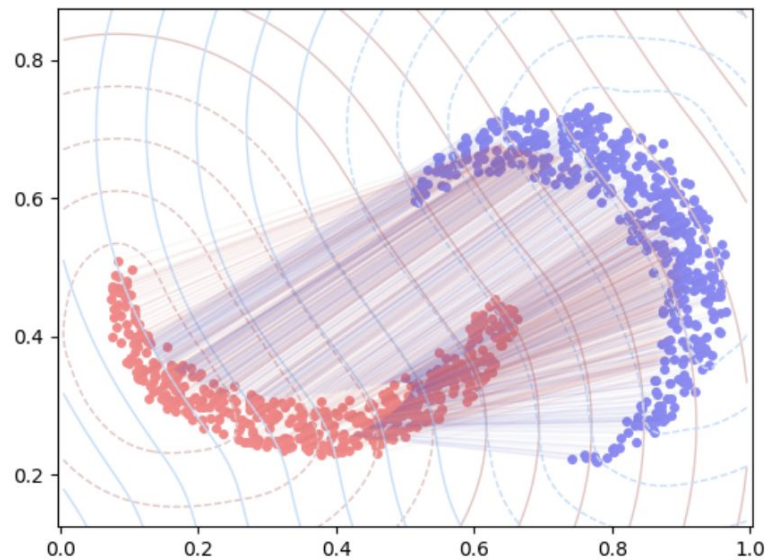


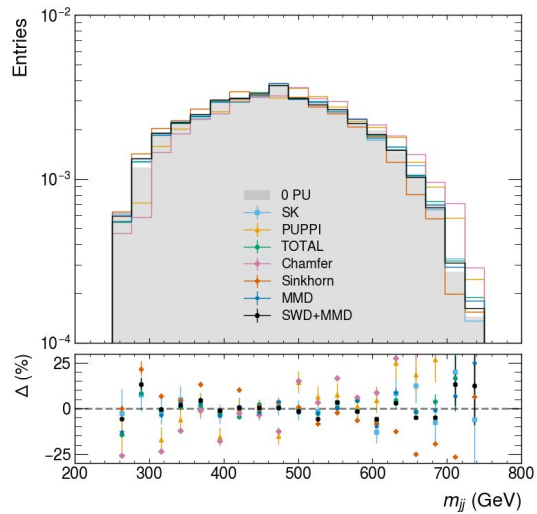
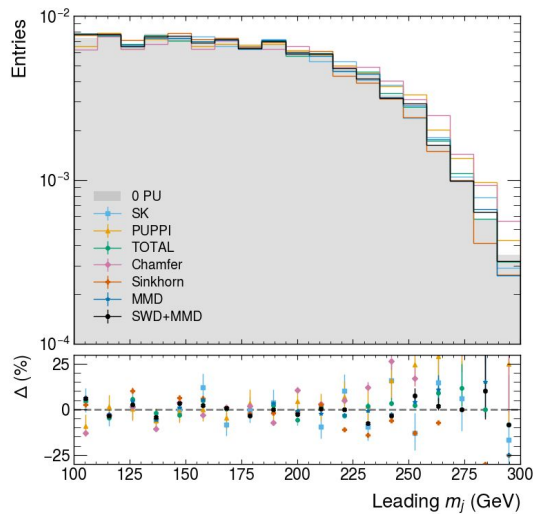
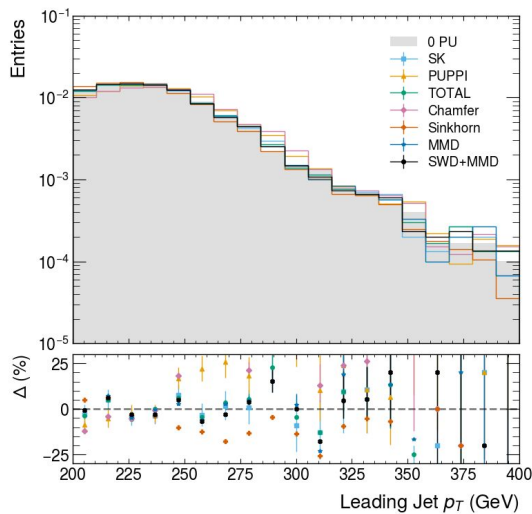
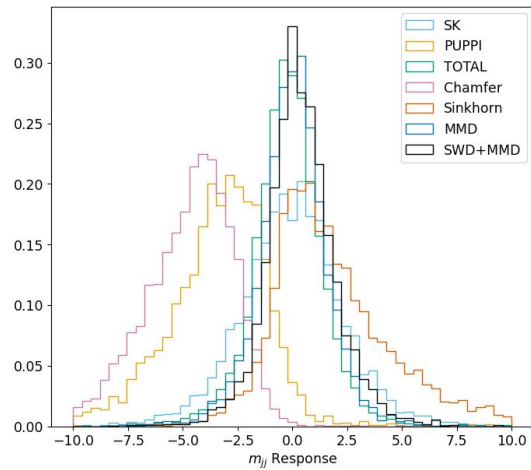
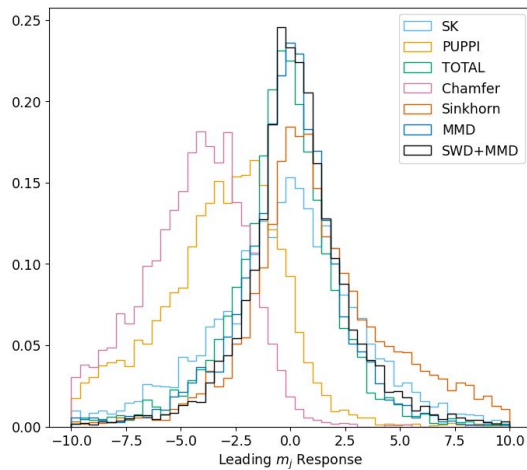
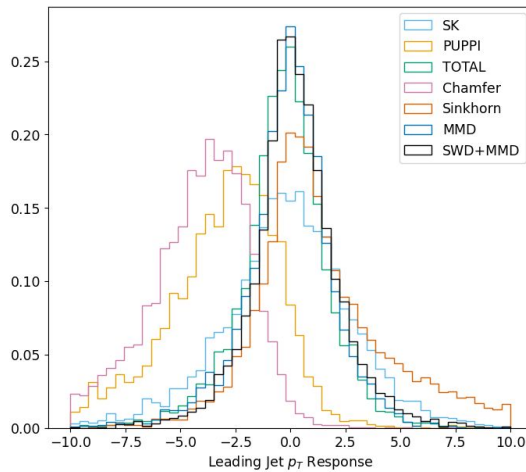


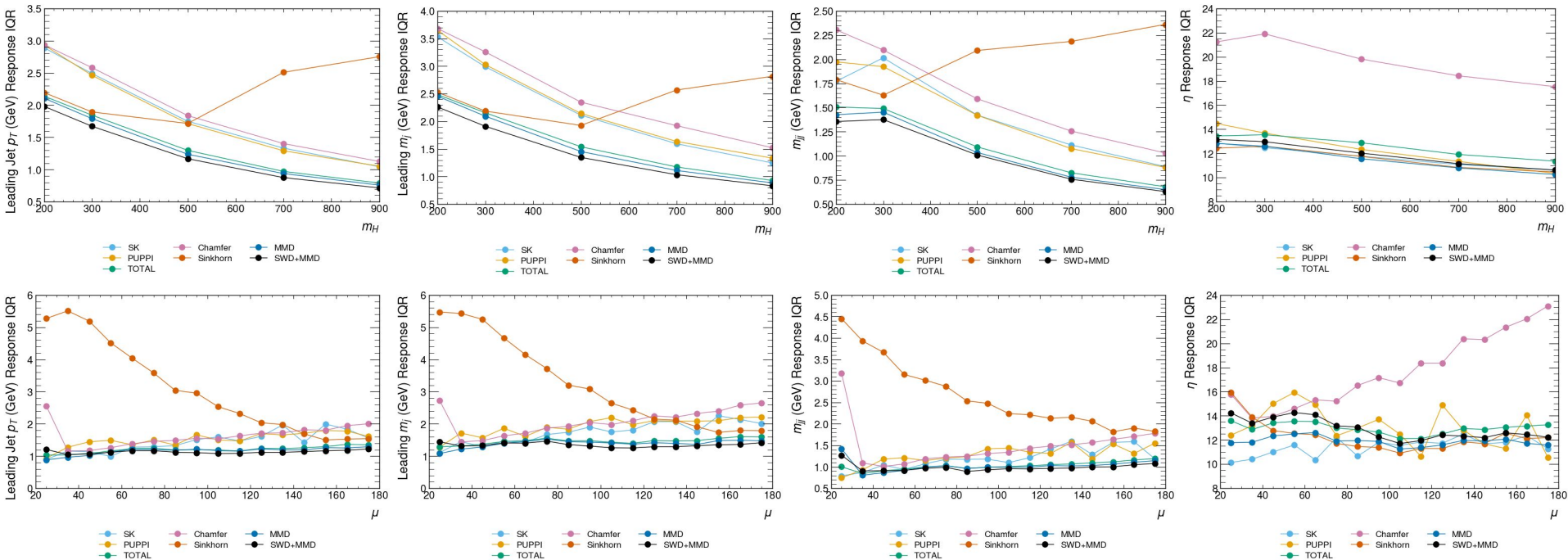


(Alternative) Loss Functions

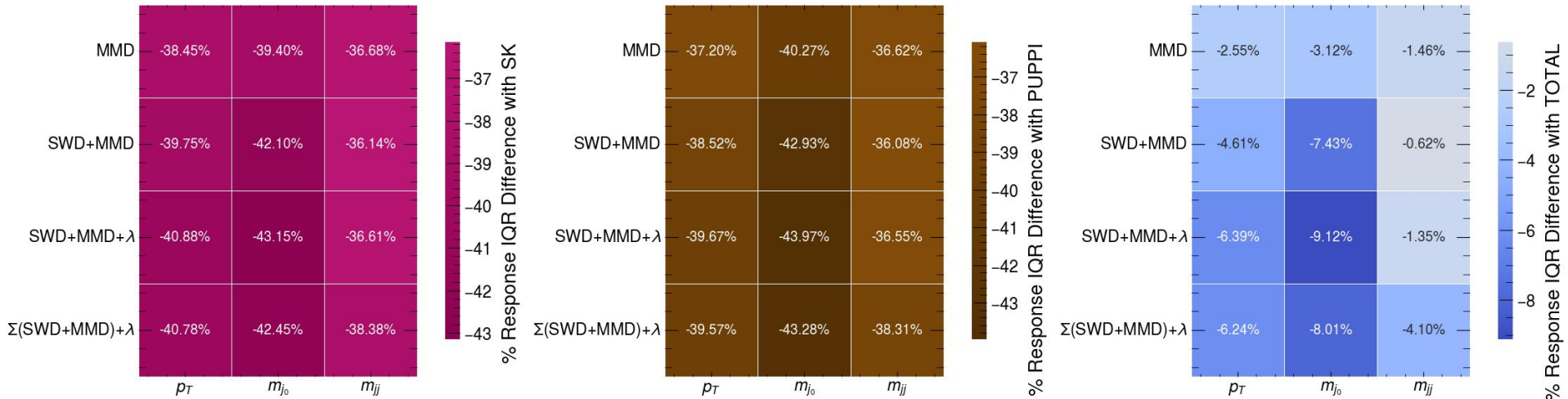
1. **Sliced Wasserstein Distance (SWD)**
2. Chamfer distance
3. Sinkhorn divergence
4. Maximum Mean Discrepancy (MMD)
 - a. MMD + SWD

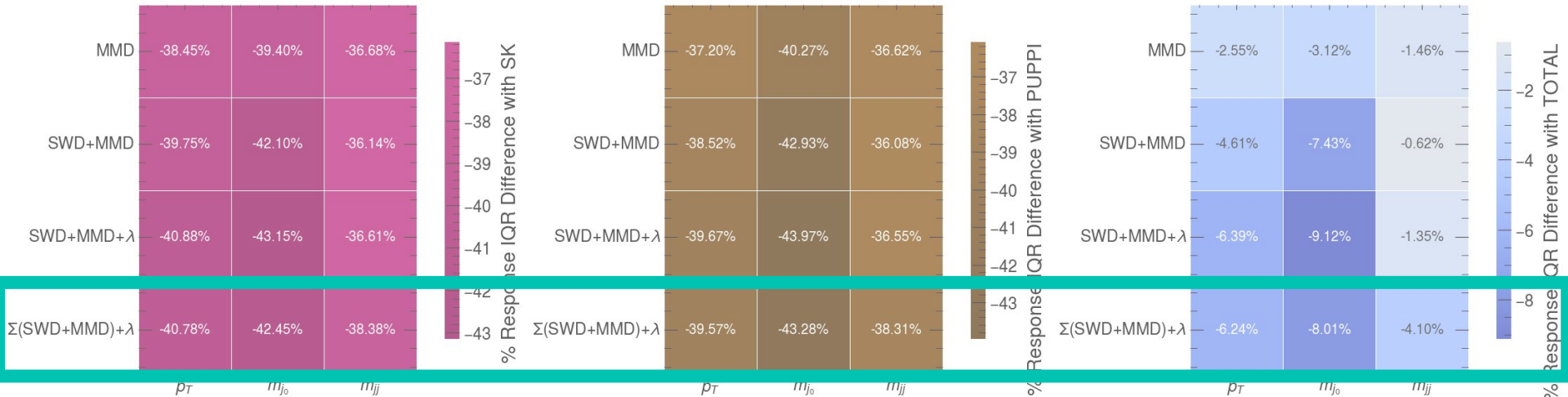






So what is
the *total*
TOTAL?





Key Takeaways

Stay tuned!
Final results to be released soon
arXiv 2607.XXXXXX

01

TOTAL is a data-driven, optimal
transport-based pileup mitigation technique

02

The addition of physics-based constraints
and a new loss function boosts TOTAL's
performance

03

Optimized TOTAL performs ~7% better than
TOTAL and ~40% better than PUPPI and SK

Appendices

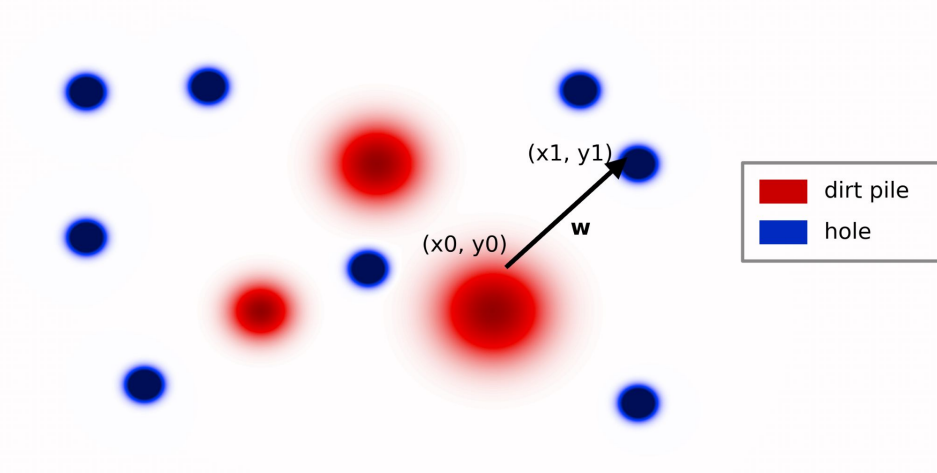
Earth Mover's Distance $= W_1$

- Assumption: Total volume of the holes = total volume of the dirt piles
- Piles as the probability density function of P and holes as the probability density function of Q
- Per unit transportation cost:

$$C(x_0, y_0, x_1, y_1) = (x_0 - x_1)^2 + (y_0 - y_1)^2$$

- Transportation Plan:

$$T(x_0, y_0, x_1, y_1) = w$$



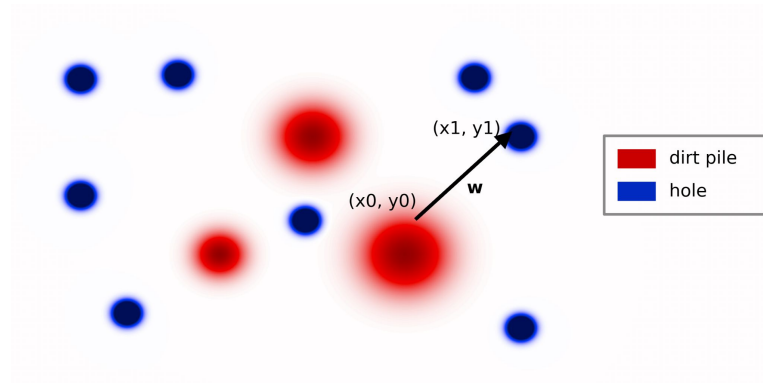
Earth Mover's Distance = W_1

- Assumption: Total volume of the holes = total volume of the dirt piles
- Piles as the probability density function of P and holes as the probability density function of Q
- Per unit transportation cost:

$$C(x_0, y_0, x_1, y_1) = (x_0 - x_1)^2 + (y_0 - y_1)^2$$

- Transportation Plan:

$$T(x_0, y_0, x_1, y_1) = w$$

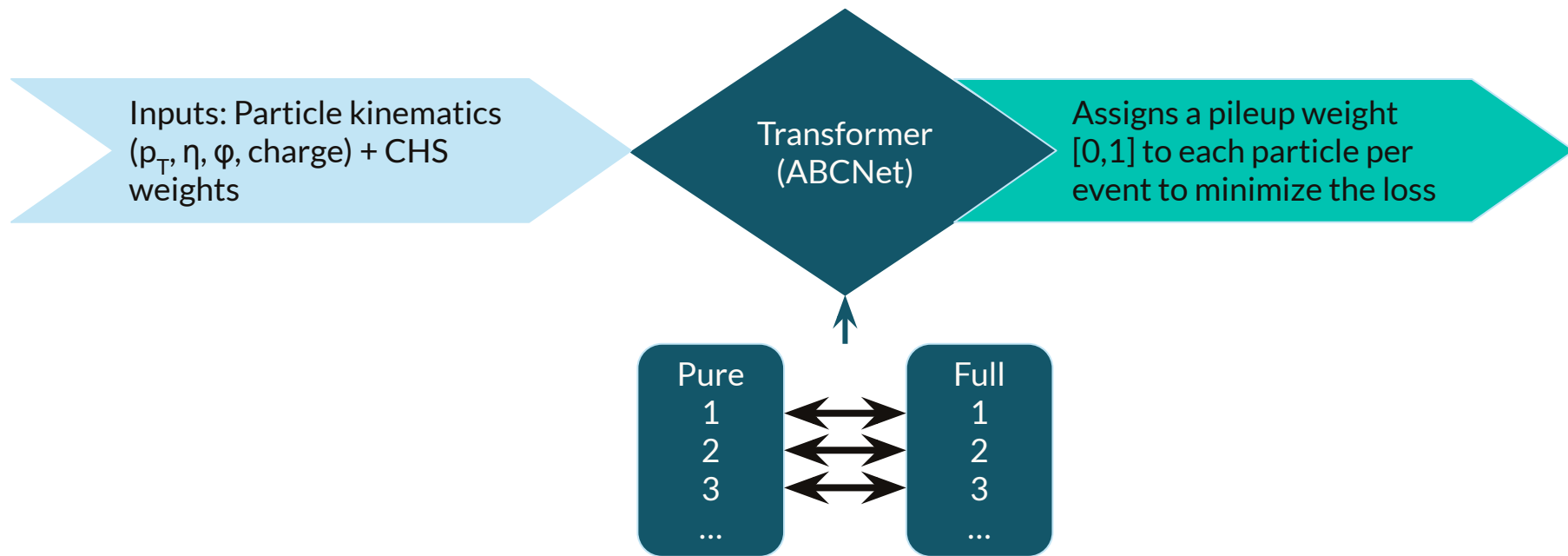


$$\int \int T(x_0, y_0, x, y) dx dy = p(x_0, y_0)$$

$$\int \int T(x, y, x_1, y_1) dx dy = q(x_1, y_1)$$

$$\text{Total Cost} = \int \int \int \int C(x_0, y_0, x_1, y_1) \cdot T(x_0, y_0, x_1, y_1) dx_0 dy_0 dx_1 dy_1$$

TOTAL: Training Optimal Transport with Attention Learning



TOTAL: Training Optimal Transport with Attention Learning

$$\mathcal{L} = \text{SWD}(x'_p, x_{np}) + \lambda \text{MSE}(p_T^{\text{miss}}(x'_p), p_T^{\text{miss}}(x_{np}))$$

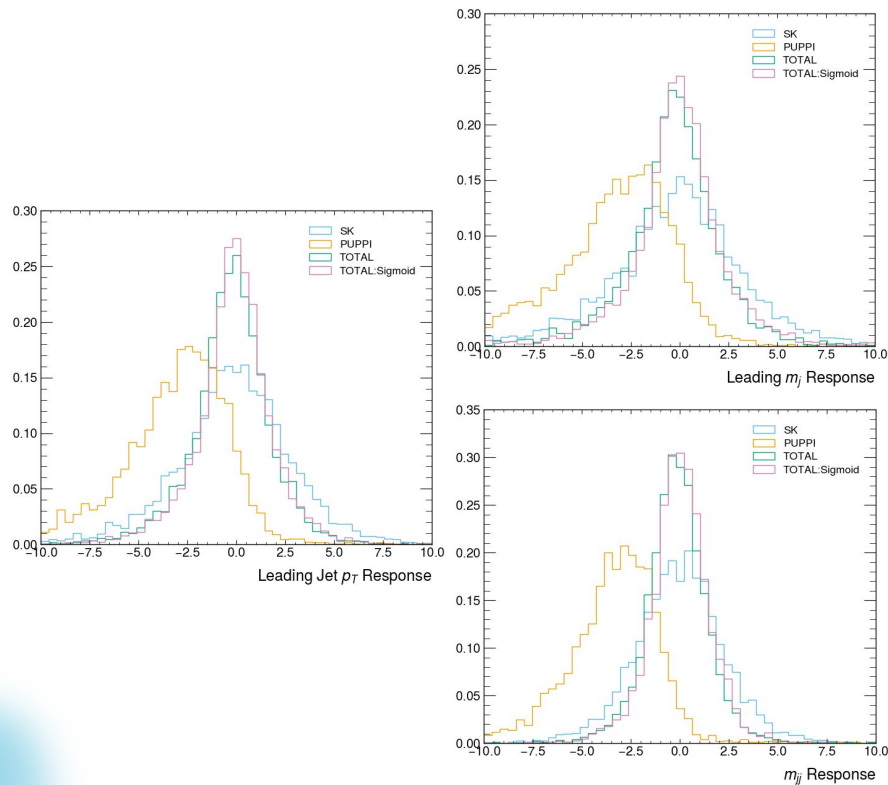


Modification for jet-based dataset (PUMML)

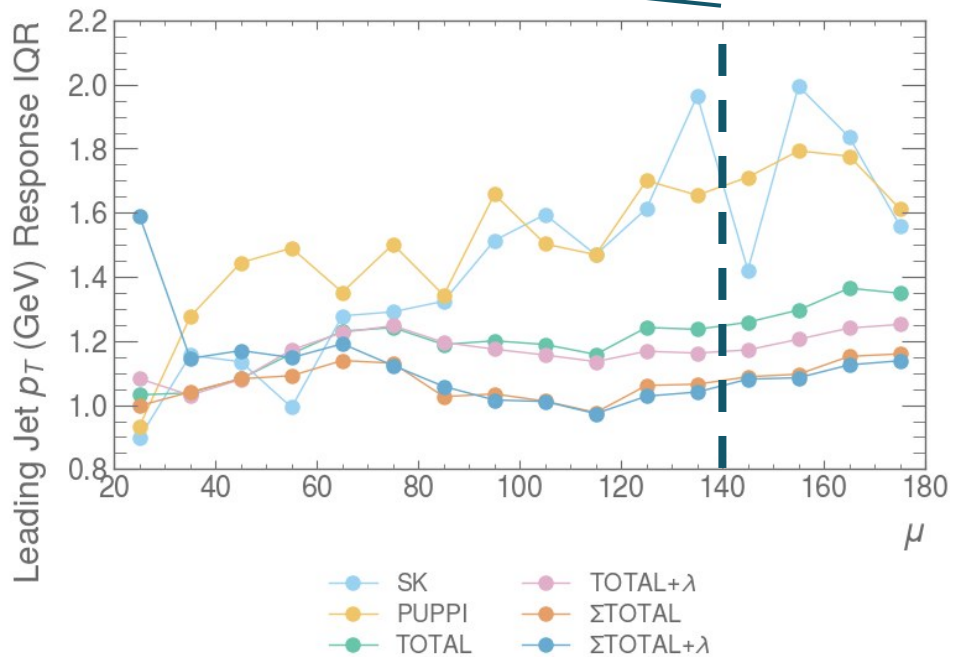
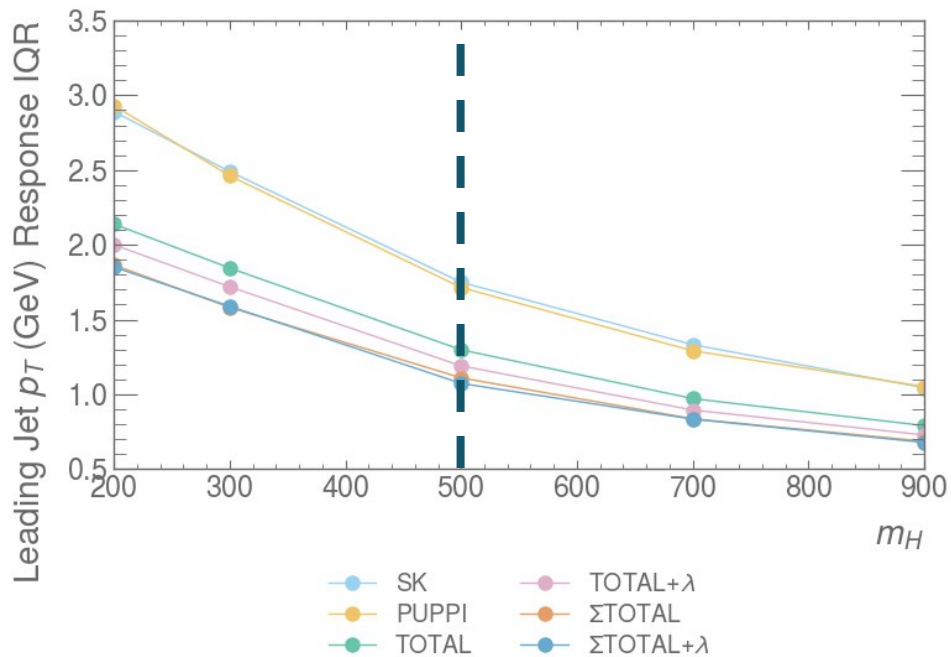
$$\mathcal{L} = \text{SWD}(x'_p, x_{np}) + \lambda \text{MSE}(E_T(x'_p), E_T(x_{np}))$$

Physics-based Constraints: Activation Function

- Hard Sigmoid
 - Fast as the function is piecewise-linear and quantizes easily for low-precision inference
 - Forces a bimodal output distribution for the classifier labels
 - Events with differing pileup distributions can appear more similar due to the reduced resolution



*Initial training value



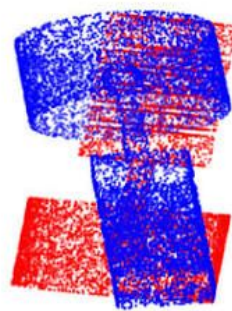
Alternative Loss Functions

1. Chamfer distance

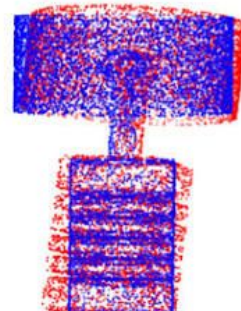
- Nearest-neighbor matching function
- For every point in the first set, it finds and measures the distance to the closest point in the second set and vice versa before summing these distances together
- Simpler and faster than SWD

$$c(A, B) = \frac{1}{|A|} \sum_{a \in A} \min_{b \in B} d(a, b),$$

$$d_C(A, B) = c(A, B) + c(B, A).$$



Rotation angle 6.58
Chamfer distance 0.12



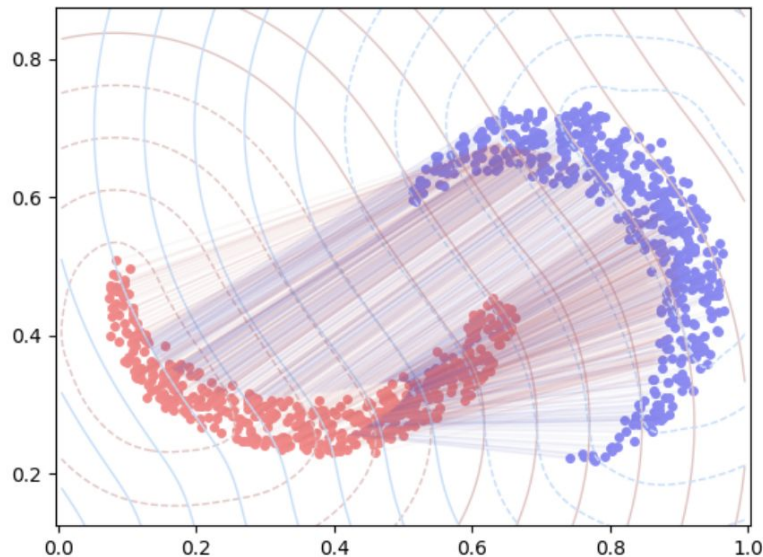
Rotation angle 2.98
Chamfer distance 0.22



Rotation angle 4.59
Chamfer distance 0.12

Alternative Loss Functions

1. Chamfer distance
2. **Sinkhorn divergence**
 - a. Computes a regularized transport coupling, allowing the metric to respect density in a superior fashion to the others including the projections of SWD.
 - b. Well-conditioned, differentiable gradients
 - c. Interpolates between true OT and MMD by ε



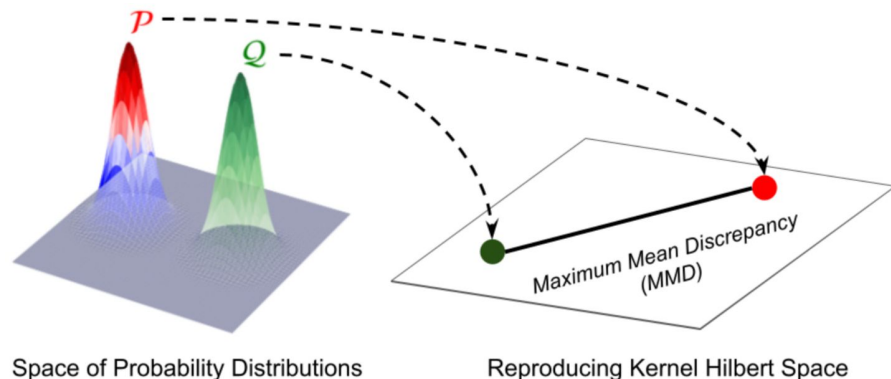
$$\text{OT}_\varepsilon(\mu, \nu) = \inf_{\pi \in \Pi(\mu, \nu)} \int_{X \times X} c(x, y) d\pi(x, y) + \varepsilon \text{KL}(\pi \| \mu \otimes \nu)$$

$$\text{KL}(\pi \| \rho) = \int \frac{d\pi}{d\rho} \log \frac{d\pi}{d\rho} - 1 d\rho$$

$$S_\varepsilon(\mu, \nu) := \text{OT}_\varepsilon(\mu, \nu) - \frac{1}{2} \text{OT}_\varepsilon(\mu, \mu) - \frac{1}{2} \text{OT}_\varepsilon(\nu, \nu)$$

Alternative Loss Functions

1. Chamfer distance
2. Sinkhorn divergence
3. **Maximum Mean Discrepancy (MMD)**
 - a. A kernel-based integral probability metric that maps data points from two distinct sets into a high-dimensional space known as a reproducing kernel Hilbert space (RKHS) using a kernel function
 - b. A closed-form estimator that does not rely on random projections
 - c. Highly sensitive to tuning of the kernel bandwidth



$$\mathcal{L}_D(\mathbf{X}^s, \mathbf{X}^t) = \text{MMD}(\mathbf{X}^s, \mathbf{X}^t) = \left\| \frac{1}{n^s} \sum_{i=1}^{n^s} \phi(\mathbf{x}_i^s) - \frac{1}{n^t} \sum_{j=1}^{n^t} \phi(\mathbf{x}_j^t) \right\|_{\mathcal{H}}^2$$

ϕ : mapping function
 $\mathbf{X}^s$: feature matrix in source domain
 $\mathbf{X}^t$: feature matrix in target domain

