

Economical Jet Taggers

Equivariant, slim, and quantized

Antoine Petitjean

September 18th, 2026

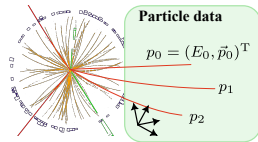


UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386

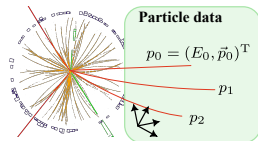


Funded by
the European Union

Jet tagging



Jet tagging

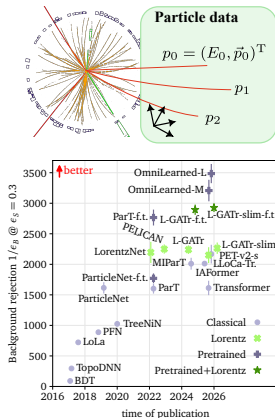


Point cloud $(p_1, \dots, p_n) =$ set of vectors v_k with additional scalar features δ_k

$$v = \begin{pmatrix} (E_0 & p_{x,0} & p_{y,0} & p_{z,0}) \\ \vdots \\ (E_n & p_{x,n} & p_{y,n} & p_{z,n}) \end{pmatrix}$$

$$\delta = \begin{pmatrix} \text{PID}_0 & \Delta\phi_0 & \Delta\eta_0 & \dots \\ \vdots & & & \\ \text{PID}_n & \Delta\phi_n & \Delta\eta_n & \dots \end{pmatrix}$$

Jet tagging

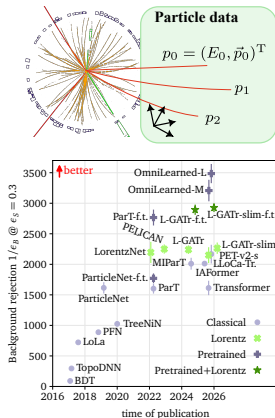


Point cloud $(p_1, \dots, p_n) =$ set of vectors v_k with additional scalar features δ_k

$$v = \begin{pmatrix} (E_0 & p_{x,0} & p_{y,0} & p_{z,0}) \\ \vdots \\ (E_n & p_{x,n} & p_{y,n} & p_{z,n}) \end{pmatrix}$$

$$\delta = \begin{pmatrix} \text{PID}_0 & \Delta\phi_0 & \Delta\eta_0 & \dots \\ \vdots & & & \\ \text{PID}_n & \Delta\phi_n & \Delta\eta_n & \dots \end{pmatrix}$$

Jet tagging



Point cloud $(p_1, \dots, p_n) =$ set of vectors v_k with additional scalar features δ_k

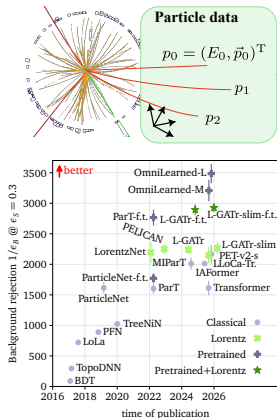
$$v = \begin{pmatrix} (E_0 & p_{x,0} & p_{y,0} & p_{z,0}) \\ \vdots \\ (E_n & p_{x,n} & p_{y,n} & p_{z,n}) \end{pmatrix}$$

$$\delta = \begin{pmatrix} \text{PID}_0 & \Delta\phi_0 & \Delta\eta_0 & \dots \\ \vdots & & & \\ \text{PID}_n & \Delta\phi_n & \Delta\eta_n & \dots \end{pmatrix}$$

$$\begin{array}{ccc} \begin{pmatrix} \delta \\ v \end{pmatrix} & \xrightarrow{\phi} & \begin{pmatrix} \delta' \\ v' \end{pmatrix} \\ \downarrow \wedge & \curvearrowright & \downarrow \wedge \\ \begin{pmatrix} \delta \\ \Lambda v \end{pmatrix} & \xrightarrow{\phi} & \begin{pmatrix} \delta' \\ \Lambda v' \end{pmatrix} \end{array}$$

Lorentz equivariance

Jet tagging



Point cloud $(p_1, \dots, p_n) =$ set of vectors v_k with additional scalar features $\mathfrak{s}_k$

$$v = \begin{pmatrix} (E_0 & p_{x,0} & p_{y,0} & p_{z,0}) \\ \vdots \\ (E_n & p_{x,n} & p_{y,n} & p_{z,n}) \end{pmatrix}_{n \times 1 \times 4}$$

$$\mathfrak{s} = \begin{pmatrix} \text{PID}_0 & \Delta\phi_0 & \Delta\eta_0 & \dots \\ \vdots \\ \text{PID}_n & \Delta\phi_0 & \Delta\eta_0 & \dots \end{pmatrix}_{n \times m \times 1}$$

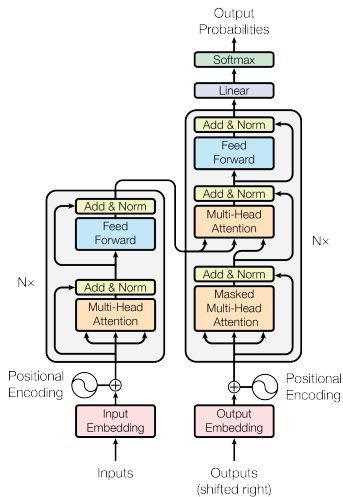
$$\begin{pmatrix} \mathfrak{s} \\ v \end{pmatrix} \xrightarrow{\phi} \begin{pmatrix} \mathfrak{s}' \\ v' \end{pmatrix}$$

$$\Lambda \downarrow \quad \quad \quad \downarrow \Lambda$$

$$\begin{pmatrix} \mathfrak{s} \\ \Lambda v \end{pmatrix} \xrightarrow{\phi} \begin{pmatrix} \mathfrak{s}' \\ \Lambda v' \end{pmatrix}$$

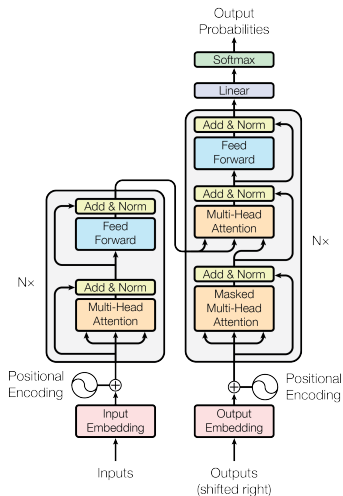
Lorentz equivariance

Standard transformer



Vaswani et al., 2017

Standard transformer



Vaswani et al., 2017

Linear

$$\text{Linear}(x) = Wx + b$$

Attention

$$A_{ij} = \frac{q_i \cdot k_j}{\sqrt{c}}$$

$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) v_j$$

Other operations

- Activation function
- Layer norm

Linear

$$\text{Linear}(x) = Wx + b$$

Attention

$$A_{ij} = \frac{q_i \cdot k_j}{\sqrt{c}}$$
$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) v_j$$

Linear

$$\text{Linear}(x) = Wx + b$$

Attention

$$A_{ij} = \frac{q_i \cdot k_j}{\sqrt{c}}$$
$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) v_j$$

Linear

$$\text{Linear}\begin{pmatrix} s \\ v \end{pmatrix} = \begin{pmatrix} w_s s + b_s \\ w_v v \end{pmatrix}$$

Linear

$$\text{Linear}(x) = Wx + b$$

Attention

$$A_{ij} = \frac{q_i \cdot k_j}{\sqrt{c}}$$
$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) v_j$$

Linear

$$\text{Linear} \begin{pmatrix} s \\ v \end{pmatrix} = \begin{pmatrix} w_s s + b_s \\ w_v v \end{pmatrix}$$

- $(n, m, 1) \mapsto (n, c_s, 1)$
- $(n, 1, 4) \mapsto (n, c_v, 4)$

Linear

$$\text{Linear}(x) = Wx + b$$

Attention

$$A_{ij} = \frac{q_i \cdot k_j}{\sqrt{c}}$$
$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) v_j$$

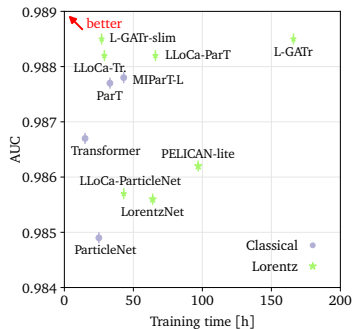
Linear

$$\text{Linear} \begin{pmatrix} \delta \\ v \end{pmatrix} = \begin{pmatrix} w_\delta \delta + b_\delta \\ w_v v \end{pmatrix}$$

Attention

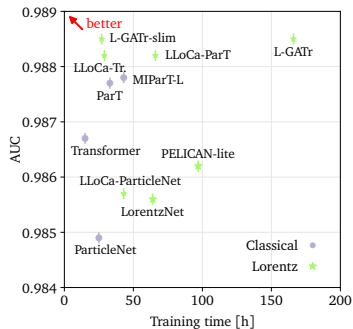
$$A_{ij} = \frac{q_{\delta,i} \cdot k_{\delta,j} + \langle q_{v,i}, k_{v,j} \rangle}{\sqrt{c_\delta + 4c_v}}$$
$$\text{out}_i = \sum_j \text{Softmax}_j(A_{ij}) \begin{pmatrix} v_\delta \\ v_v \end{pmatrix}_j$$

SOTA performance of L-GATr-slim

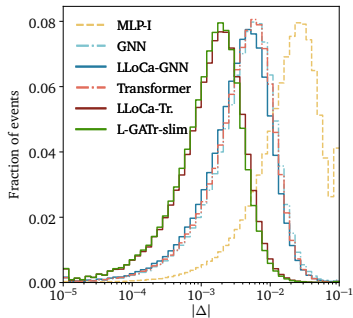


Jet tagging

SOTA performance of L-GATr-slim

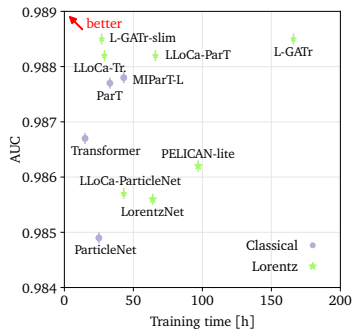


Jet tagging

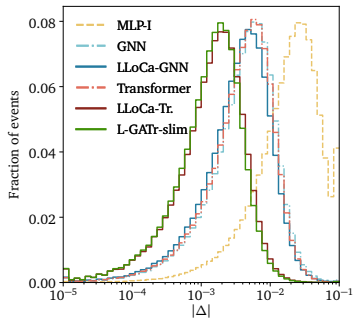


Amplitude regression

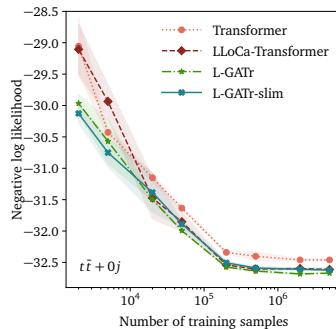
SOTA performance of L-GATr-slim



Jet tagging

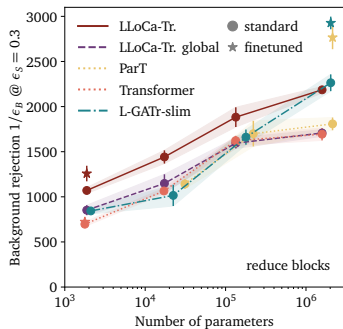


Amplitude regression



Event generation

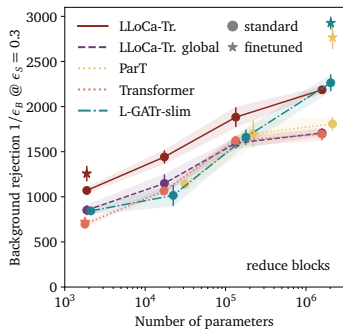
Top taggers from $O(10^6)$ down to $O(10^3)$ parameters



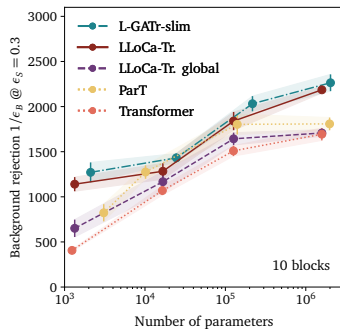
Fewer blocks: $10 \rightarrow 1$

Downscaling

Top taggers from $O(10^6)$ down to $O(10^3)$ parameters

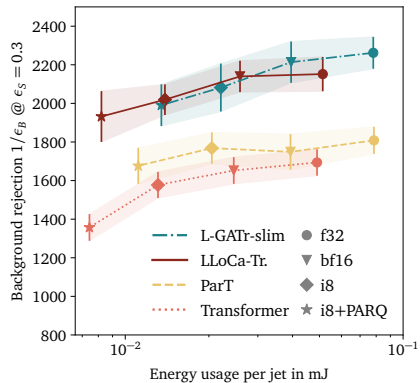


Fewer blocks: $10 \rightarrow 1$

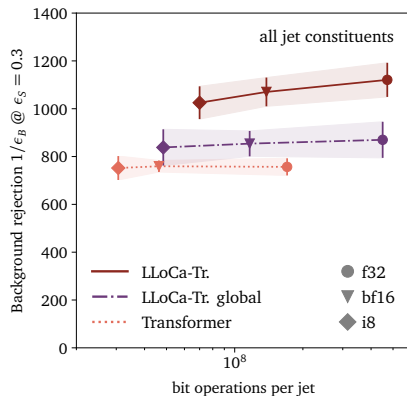
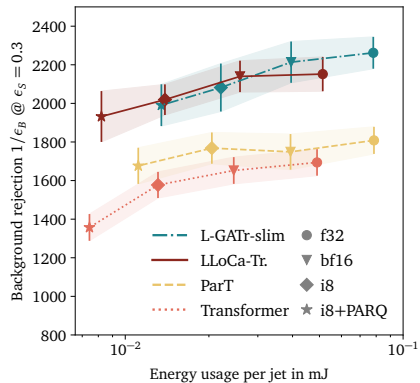


Narrower blocks: reduce latent dims

Quantization savings



Quantization savings



- **Simple equivariance**
 - L-GATr-slim = transformer with scalars and vectors streams
 - L-GATr performance at a much lower cost

- **Simple equivariance**
 - L-GATr-slim = transformer with scalars and vectors streams
 - L-GATr performance at a much lower cost
- **Down-scaling**
 - Shrink to 1k parameters
 - Quantization to `int8`

- **Simple equivariance**
 - L-GATr-slim = transformer with scalars and vectors streams
 - L-GATr performance at a much lower cost
- **Down-scaling**
 - Shrink to 1k parameters
 - Quantization to `int8`
- **Next steps**
 - Better L-GATr-slim performance at 1 block
 - Different quantization scheme (HGQ)
 - FPGA implementation

Conclusion

- **Simple equivariance**

- L-GATr-slim = transformer with scalars and vectors streams
- L-GATr performance at a much lower cost

- **Down-scaling**

- Shrink to 1k parameters
- Quantization to `int8`

- **Next steps**

- Better L-GATr-slim performance at 1 block
- Different quantization scheme (HGQ)
- FPGA implementation



`pip install lgatr`



`arXiv:2512.17011`

Backup: Quantization

Matrix multiplications dominate the cost:
run them at lower precision.

- inputs: zero-point quantization

$$x_q = \text{clip}\left(\left\lfloor \frac{x}{s} \right\rfloor + z\right)$$

- weights: quantization-aware training
with proximal gradient methods
(PARQ/STE)

	E_{add} [pJ]	E_{mul} [pJ]
float32	0.38	1.31
bfloat16	0.11	0.21
int8	0.007	0.07

Energy per operation, A100 GPU

Backup: weight quantization as a regularizer

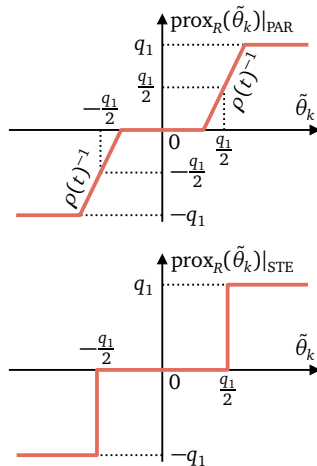
Push weights to ternary $\{0, \pm q\}$.

Constraint via proximal gradient

$$\mathcal{L}_{\text{QAT}} = \mathcal{L} + \lambda \sum_k R(\theta_k)$$

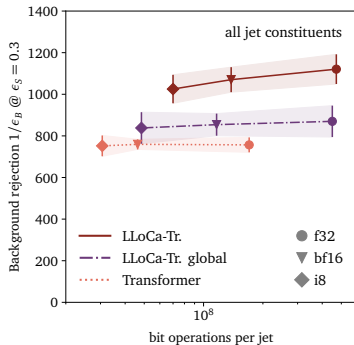
applied as a post-update map $\text{prox}_R(\tilde{\theta}_k)$

- **PARQ**: anneal map identity $\rightarrow$ step
- **STE**: hard step, special case

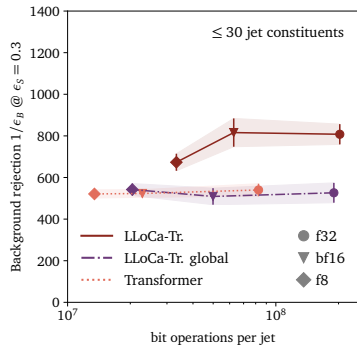


Backup: towards trigger-level tagging

Minimal LLoCa-Transformer: 1 block, 16-dim latent, static int8



All constituents



30 leading constituents

Backup: L-GATr-slim building blocks

$$\text{GLU}\begin{pmatrix} \mathfrak{z} \\ \mathfrak{v} \end{pmatrix} = \begin{pmatrix} \text{GELU}(\text{Linear}(\mathfrak{z})) \cdot \text{Linear}(\mathfrak{z}) \\ \text{GELU}(\langle \text{Linear}(\mathfrak{v}), \text{Linear}(\mathfrak{v}) \rangle) \cdot \text{Linear}(\mathfrak{v}) \end{pmatrix}$$

$$\text{RMSNorm}\begin{pmatrix} \mathfrak{z} \\ \mathfrak{v} \end{pmatrix} = \left(\frac{1}{c_{\mathfrak{v}} + c_{\mathfrak{z}}} \left(\sum_{c=1}^{c_{\mathfrak{v}}} |\langle \mathfrak{v}_c, \mathfrak{v}_c \rangle|^2 + \sum_{c=1}^{c_{\mathfrak{z}}} \mathfrak{z}_c^2 \right) + \epsilon \right)^{-1/2} \begin{pmatrix} \mathfrak{z} \\ \mathfrak{v} \end{pmatrix}$$

Backup: LLoCa

