

Toward Domain-Invariant Tokenization for Jet Foundation Models

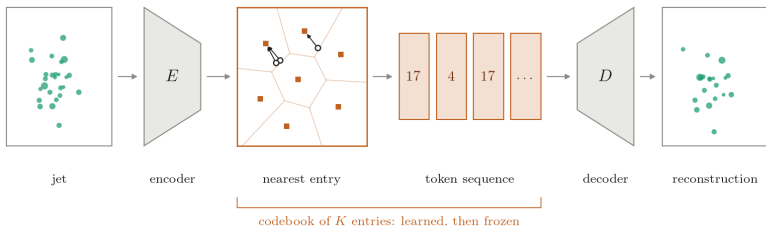
Giovanni Ottaviano, Chris Scheulen, Tobias Golling



Why discretize and what it costs

2 / 11

Foundation models for physics: one backbone, many systems. Two routes: continuous representations or discrete tokens



The price: q always returns a code, never “none of these fit”, and what the bottleneck discards is lost to every downstream task

This talk

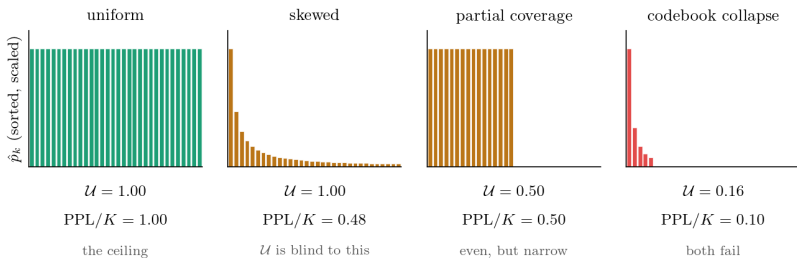
How do we check the vocabulary is doing its job, before training anything on top of it, and on physics it has never seen?

Evaluating the discrete bottleneck

3 / 11

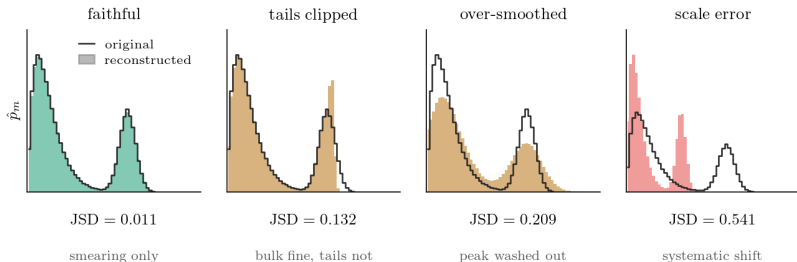
Setting: nearest-neighbour quantizer q over a codebook $\mathcal{C}$ of size $K \rightarrow$ empirical code usage $\hat{p}$ over a stream of latents

- **Utilization** $\mathcal{U} \in [0, 1]$: fraction of codes ever used $\rightarrow$ *coverage*
- **Normalized perplexity** $\text{PPL}/K \in [0, 1]$: effective vocabulary size $\rightarrow$ *evenness*



Setting: input x , reconstruction $\hat{x} = D(q(E(x)))$, physical quantity $\mathcal{O}$
→ histograms $\hat{p}$, $\hat{q}$ on shared bins

Goal: does the tokenizer preserve the *statistics* of $\mathcal{O}$, not just its point-wise values?



Jensen-Shannon distance

Symmetric in $\hat{p}, \hat{q}$ · bounded: $\text{JSD} \in [0, 1]$ · a true metric

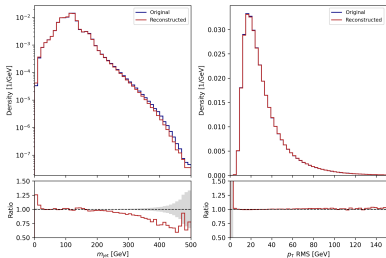
Setup: Same transformers-based VQ-VAE used for Omnijet- α FM

- ▶ **Codebook:** trained on JetClass, then frozen
- ▶ **Selection:** $p_{T_{\text{jet}}} \in [500, 1000]$ GeV, $\Delta R < 0.8$
- ▶ **Normalization:** applied unchanged to each dataset

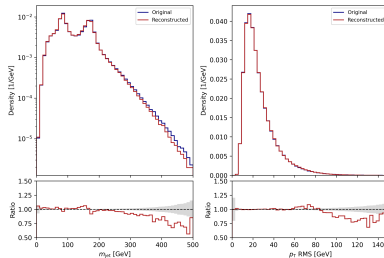
Dataset	$\mathcal{U}$ (%)	PPL/ K	jets
JetClass	86.99	0.748	20M
Aspen	87.05	0.623	$\sim 24\text{M}$
Rodem	87.51	0.693	$\sim 1.8\text{M}$

Coverage is flat across all three, evenness is not

Same frozen codebook and same round trip to compare original against reconstructed (in-domain vs. off-domain case).



JetClass: in-domain reference



Rodem: off-domain case

Reconstruction fidelity is a separated from codebook coverage, and Jensen-Shannon distance (JSD) is the tool we use to investigate it.

One token, three relative coordinates:

$$\left(\underbrace{\log z}_{z = p_{T,i}/p_{T,\text{jet}}}, \underbrace{\Delta\eta, \Delta\varphi}_{\text{w.r.t. the jet axis}} \right) \quad \text{vs.} \quad (\log p_{T,i}, \Delta\eta, \Delta\varphi)$$

Each entry is measured against the jet itself $\Rightarrow$ the tuple describes *what splitting happened*, not where the jet sits

- ▶ **Which variable** ($p_{T,i} \rightarrow z$): $p_{T,i}$ inherits the jet momentum, and where a jet's p_T sits is decided by *sample selection* (generation cut, skim). The ratio z is set by fragmentation, and QCD radiation is approximately scale invariant $\Rightarrow z$ is the near-universal entry
- ▶ **Which transform** ($z \rightarrow \log z$): among transforms of z , $\log$ spreads codebook occupancy far more evenly (ex. the median at 37.8% of the p01–p99 range against 13.2% for $\sqrt{z}$)¹

¹p01, p99: the 1st and 99th percentiles of the distribution.

Putting the energy scale back

The tuple is dimensionless $\Rightarrow$ the jet energy scale must be included in a different way

$$p_{T,i} = e^{\log z_i} \cdot p_{T,\text{jet}} \quad (\text{use jet-}p_T \text{ to re-encode constituents' } p_{T,i})$$

	Product quantization	Conditioning
Scale lives	inside the token stream	outside the vocabulary
Token	$(i_{\text{const}}, i_{\text{scale}})$	i_{const} only
Vocabulary	$K_c \times K_s$	K_c
Pros	self-contained sequence, scale is itself generable	zero rate cost; structure only codebook; easy off-envelope extrapolation
Cons	larger vocabulary; untrained embedding for off-envelope, $\mathcal{U}$ collapses	must enter encoder <i>and</i> decoder, then be ablated for side-channel collapse

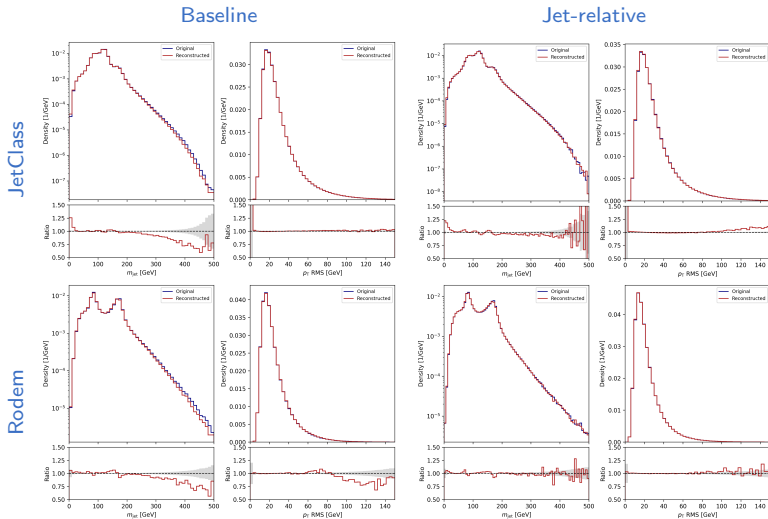
- ▶ Same codebook recipe, same frozen-transfer protocol
- ▶ Same selection: $p_{T,\text{jet}} \in [500, 1000]$ GeV, $\Delta R < 0.8$
- ▶ Only the coordinates change: $\log p_{T,i} \rightarrow \log z$

Dataset	$\Delta \mathcal{U}$	$\mathcal{U}$ (%)	$\Delta \text{PPL}/K$	PPL/K
JetClass	+10.4	97.39	+0.071	0.819
Aspen	+10.4	97.42	+0.239	0.862
Rodem	+9.8	97.35	+0.159	0.852

Coverage gains ~ 10 percentage points everywhere, and evenness improves most off-domain $\rightarrow$ Aspen and Rodem now exceed JetClass, despite the codebook being fitted on it.

Reconstructed statistics: coordinates vs. domain

10 / 11



Across: the coordinate change. **Down:** in-domain to off-domain.

- ▶ The tokenizer is frozen before the backbone → whatever it discards is unavailable on every downstream task
- ▶ Measured with $\mathcal{U}$, PPL/ K and JSD on reconstructed physical statistics (necessary, but not sufficient)
- ▶ Jet-relative coordinates: z removes a sample-selection variable (set only by physical arguments), log spreads occupancy far more evenly than z
- ▶ Scale re-enters through the constituent: product quantization or conditioning

Thank you!

ottaviano@lpsm.paris — Poster 50



slides & links

Backup

The Kullback–Leibler divergence

Definition: for p, q on $\{1, \dots, K\}$ with $p \ll q$,

$$\text{KL}(p \parallel q) = \sum_{k=1}^K p_k \log_2 \frac{p_k}{q_k} = \underbrace{\mathbb{E}_{k \sim p} \left[\log_2 \frac{p_k}{q_k} \right]}_{\text{log-likelihood ratio under } p}$$

Reading: expected number of extra bits needed to encode samples from p using a code optimized for q

- ▶ **Non-negative:** $\text{KL}(p \parallel q) \geq 0$, with equality iff $p = q$ (Jensen)
- ▶ **Not a metric:** $\text{KL}(p \parallel q) \neq \text{KL}(q \parallel p)$, and the triangle inequality fails
- ▶ **Unbounded:** $\text{KL}(p \parallel q) = +\infty$ if $q_k = 0 < p_k$ for some k

Entropy decomposition

$$\text{KL}(p \parallel q) = \underbrace{- \sum_k p_k \log_2 q_k}_{\text{cross-entropy } H(p, q)} - \underbrace{\left(- \sum_k p_k \log_2 p_k \right)}_{\text{entropy } H(p)}$$

Setting: encoder E , codebook $\mathcal{C} = \{\beta_1, \dots, \beta_K\}$, nearest-neighbour quantizer

$$q(\mathbf{z}) = \arg \min_{1 \leq k \leq K} \|\mathbf{z} - \beta_k\|_2, \quad \mathbf{z} \in \mathbb{R}^d$$

Goal: quantify how well the K codes are exploited on a stream of latents $(\mathbf{z}_i)_{1 \leq i \leq N}$

► **Empirical code distribution:** $\hat{p}_k = \frac{1}{N} \sum_{i=1}^N 1\{q(\mathbf{z}_i) = k\}$

► **Failure mode:** index collapse, i.e. $\hat{p}$ concentrated on $\ll K$ codes

Codebook metrics

$$\underbrace{\mathcal{U}(\hat{p}) = \frac{1}{K} \sum_{k=1}^K 1\{\hat{p}_k > 0\}}_{\text{utilization} \in [0,1]}$$

$$\underbrace{\text{PPL}(\hat{p}) = 2^{(-\sum_k \hat{p}_k \log_2 \hat{p}_k)}}_{\text{perplexity} \in [1, K]}$$

Setting: input x , reconstruction $\hat{x} = D(q(E(x)))$, physical quantity $\mathcal{O}$

Histograms on shared bins $B_1, \dots, B_M$:

$$\hat{p}_m = \frac{1}{N} \sum_{i=1}^N 1\{\mathcal{O}(x)_i \in B_m\}, \quad \hat{q}_m = \frac{1}{N} \sum_{i=1}^N 1\{\mathcal{O}(\hat{x})_i \in B_m\}$$

Goal: does the tokenizer preserve the *statistics* of $\mathcal{O}$, not just its point-wise values?

- ▶ **Mixture:** $\bar{m} = \frac{1}{2}(\hat{p} + \hat{q})$, with $\bar{m}_k \geq \frac{1}{2}\hat{p}_k$ and $\bar{m}_k \geq \frac{1}{2}\hat{q}_k$ for all k
- ▶ **Why not KL:** asymmetric and $= +\infty$ as soon as $\hat{q}_k = 0 < \hat{p}_k$

Jensen-Shannon distance

$$\text{JSD}(\hat{p}, \hat{q}) = \underbrace{\sqrt{\frac{1}{2}\text{KL}(\hat{p} \parallel \bar{m}) + \frac{1}{2}\text{KL}(\hat{q} \parallel \bar{m})}}_{\text{symmetric, bounded: JSD} \in [0,1]}$$

Product quantization: two indices, two embedding tables:

$$\text{emb} = \underbrace{E_c(i_{\text{const}})}_{K_c \times d} + \underbrace{E_s(i_{\text{scale}})}_{K_s \times d} \quad \text{or} \quad \text{emb} = W[E_c(i_{\text{const}}); E_s(i_{\text{scale}})]$$

These are the same function class: $[a; b]W = aW_1 + bW_2 \Rightarrow$ a dimension budget, not a modelling decision

- ▶ **Sum:** keeps d fixed, ties both tables to it
- ▶ **Concat:** splits d between the two, makes the provenance explicit

Conditioning: same sum, but the scale term is *continuous*:

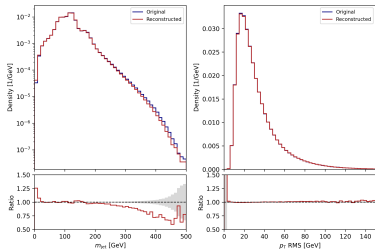
$$\text{emb} = E_c(i_{\text{const}}) + \underbrace{g(\log p_{T,i})}_{\text{no cbook, no bins}} \quad \text{or} \quad h = \text{model}(\text{tokens}, \text{cond} = \log p_{T,i})$$

	Origin	Detector	R	Constituents	Jet p_T [GeV]
JetClass	MG5 + Pythia	CMS (Delphes)	0.8	PF	500–1000
Aspen	CMS 2016 data	CMS	0.8	PF, PUPPI	≥ 300
JetSet	Powheg + Pythia	ATLAS (Geant4)	0.4	tracks	20–500
Rodem	MG5 + Pythia	ATLAS (Delphes)	1.0	PF	450–1200

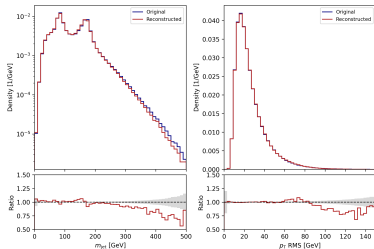
- ▶ **JetClass**: the reference benchmark (100M train / 5M val / 20M test), 10 classes, $|\eta| < 2$. **No pileup, no measurement uncertainty**.
- ▶ **Aspen**: CMS open data (~ 178 M unlabelled jets), 16.4 fb^{-1} . Trigger bias near turn-on and detector artefacts (visible geometry at $|\eta| \simeq 1.6$, dead cells)
- ▶ **JetSet**: ATLAS open data, ~ 199 M jets from $t\bar{t}$ at $\sqrt{s} = 13.6 \text{ TeV}$. The only **full Geant4** simulation, and the only one built from **charged tracks alone** (no neutral component)
- ▶ **Rodem**: large-radius jets, hence ~ 50.5 constituents/jet, and **no soft- p_T threshold**, the spectrum runs smooth to the bottom, with 5.09% below 0.5 GeV

Setup: single JetClass-trained codebook, applied frozen to all four datasets
 $(p_{T,\text{jet}} \in [500, 1000] \text{ GeV}, \Delta R < 0.8, \text{ same input normalization})$

Dataset	$\mathcal{U}$ (%)	PPL/ K	jets
JetClass	86.99	0.748	20M
Aspen	87.05	0.623	$\sim 24\text{M}$
JetSet	68.55	0.333	$\sim 25\text{M}$
Rodem	87.51	0.693	$\sim 3\text{M}$



JetClass: in-domain reference



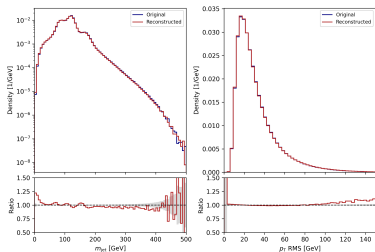
Rodem: original vs. reconstructed

Full results: the jet-relative tuple

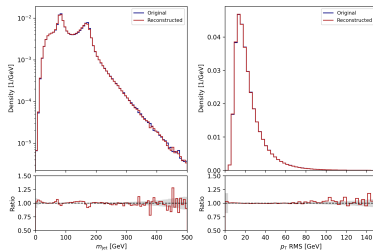
19 / 11

Same setup, same frozen JetClass codebook, only the coordinates change
($p_{T,\text{jet}} \in [500, 1000]$ GeV, $\Delta R < 0.8$)

Dataset	$\mathcal{U}$ (%)	Δ	PPL/ K	Δ
JetClass	97.39	+10.4	0.819	+0.071
Aspen	97.42	+10.4	0.862	+0.239
JetSet	68.76	+0.2	0.367	+0.034
Rodem	97.35	+9.8	0.852	+0.159



JetClass: in-domain reference



Rodem: original vs. reconstructed