

ML4Jets2026 on 09.16.2026

“Foundation Models Session”

University of Vienna, Austria



1

Authors:

Tianji Cai,
Anna Hallin,
Shivasankar K.A.,
Michael Krämer.



**What can representation
geometry reveal about
collider foundation models?**

*— Towards Scientific AI as a
New Microscope*



Presented by Tianji Cai

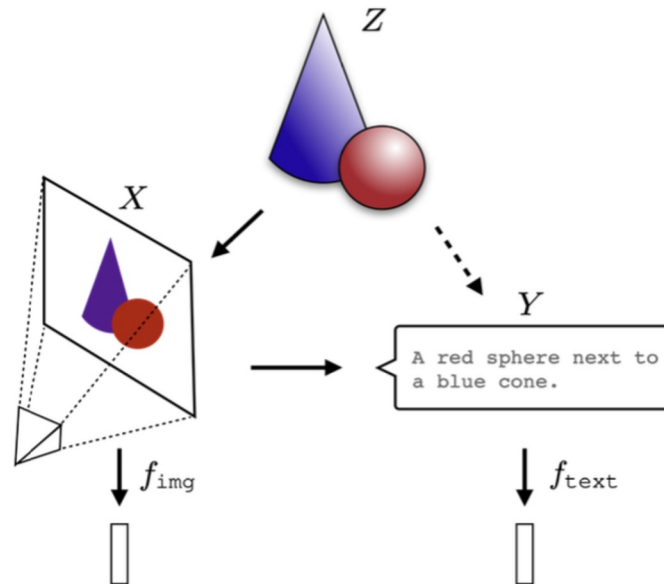
*School of Physical Science & Engineering
Tongji University, Shanghai, China*

The “Platonic Hypothesis” for us?

*Inspiration from our CS colleagues:
[2405.07987](#)*

The Platonic Representation Hypothesis

Neural networks, trained with different objectives on different data and modalities, are converging to a shared statistical model of reality in their representation spaces.



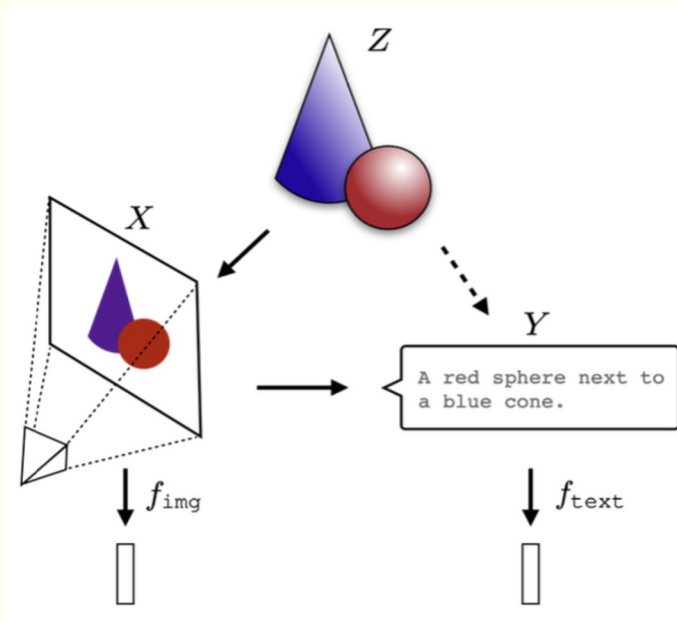
The “Platonic Hypothesis” for us?

Inspiration from our CS colleagues:

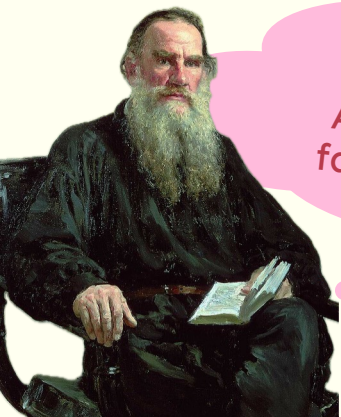
[2405.07987](#)

The Platonic Representation Hypothesis

Neural networks, trained with different objectives on different data and modalities, are converging to a shared statistical model of reality in their representation spaces.



3



Anna Karenina Principle:
All good models are alike; each failed model fails in its own way.

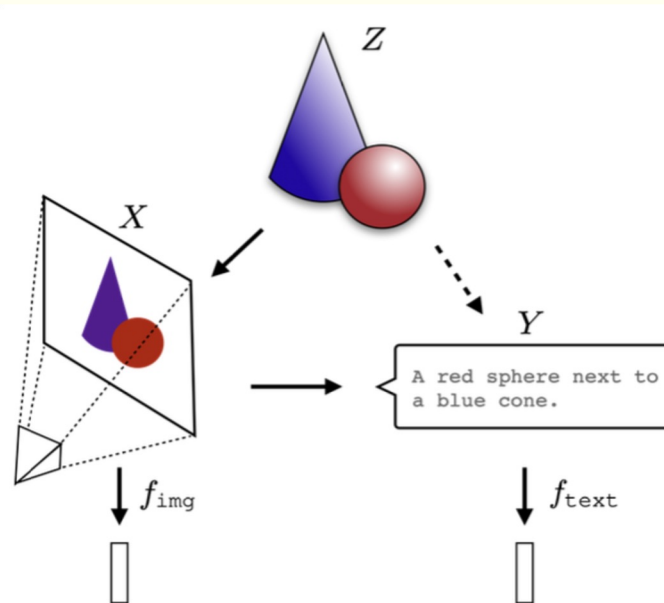
The “Platonic Hypothesis” for us?

*Inspiration from our CS colleagues:
[2405.07987](#)*

The Platonic Representation Hypothesis

Neural networks, trained with different objectives on different data and modalities, are converging to a shared statistical model of reality in their representation spaces.

Anna Karenina Principle:
All good models are alike; each failed model fails in its own way.



The “underlying reality” of our particle data is far richer than texts/images. It’s *Nature* itself!

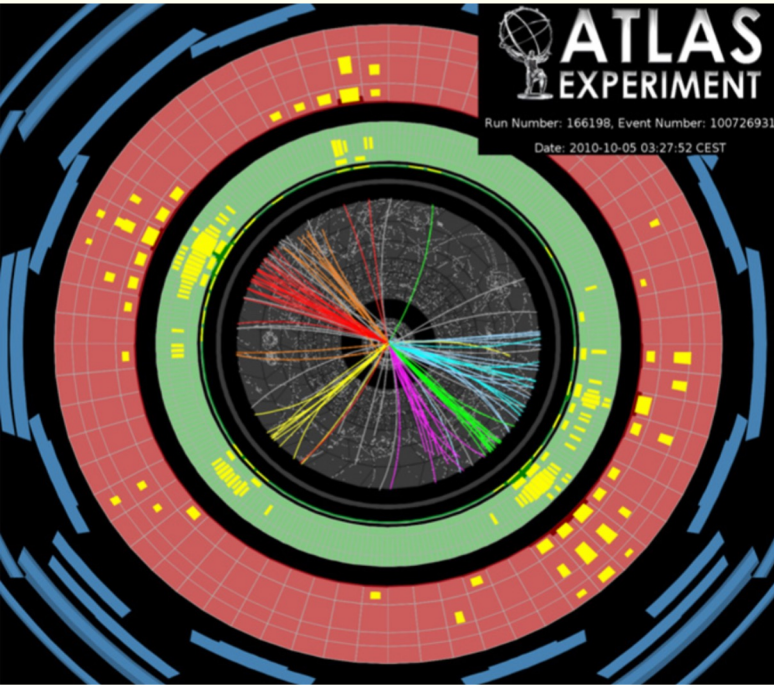
What then can our model representation spaces uncover?

Scientific(Ψ) AI as a New Microscope for Particle Physics

5

Scientific(Ψ) AI as a New Microscope for Particle Physics

6



**Good NN Models
for Collider Data
Analysis**

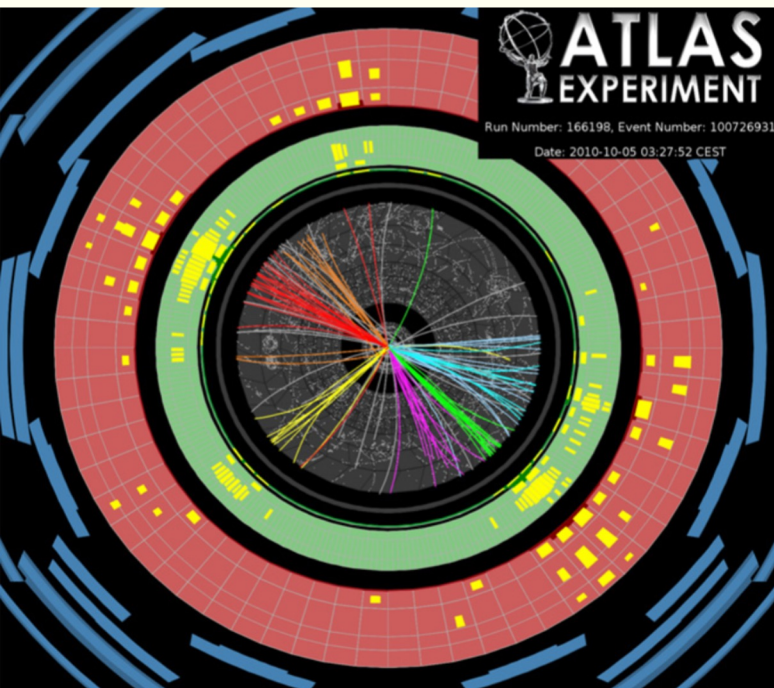
Now performance is a
prerequisite, not the focus.
Models themselves are our
discovery tool.

$$\begin{aligned}\mathcal{L} = & -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} \\ & + i\bar{\psi}\not{D}\psi \\ & + \chi_i y_{ij} \chi_j \phi + h.c. \\ & + |D_\mu \phi|^2 - V(\phi)\end{aligned}$$

*Or perhaps something
completely alien?*

Scientific(Ψ) AI as a New Microscope for Particle Physics

7



**Good NN Models
for Collider Data
Analysis**

Now performance is a
prerequisite, not the focus.
Models themselves are our
discovery tool.

$$\begin{aligned}\mathcal{L} = & -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} \\ & + i\bar{\psi}\not{D}\psi \\ & + \chi_i y_{ij} \chi_j \phi + h.c. \\ & + |D_\mu \phi|^2 - V(\phi)\end{aligned}$$

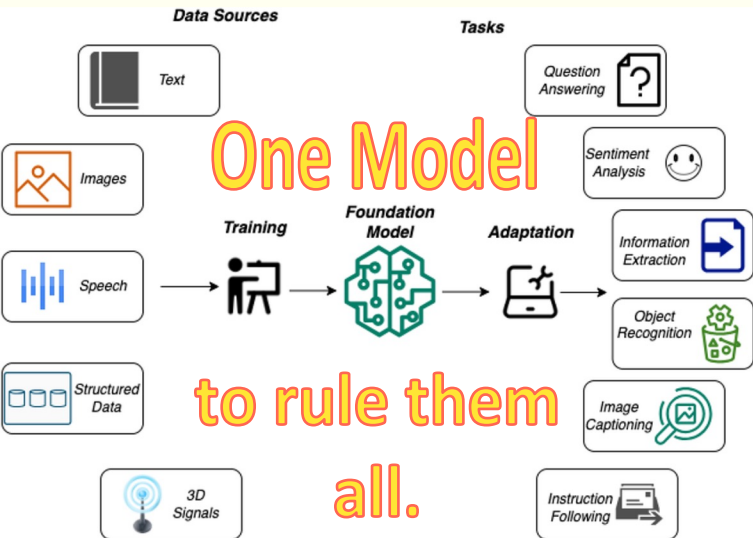
*Or perhaps something
completely alien?*

Before putting the “microscope” to use on our particle data, we first need to study the “properties” of our instrument in a more rigorous and systematic way.

→ **GOAL OF THIS WORK!**

So, what are our good models?

8



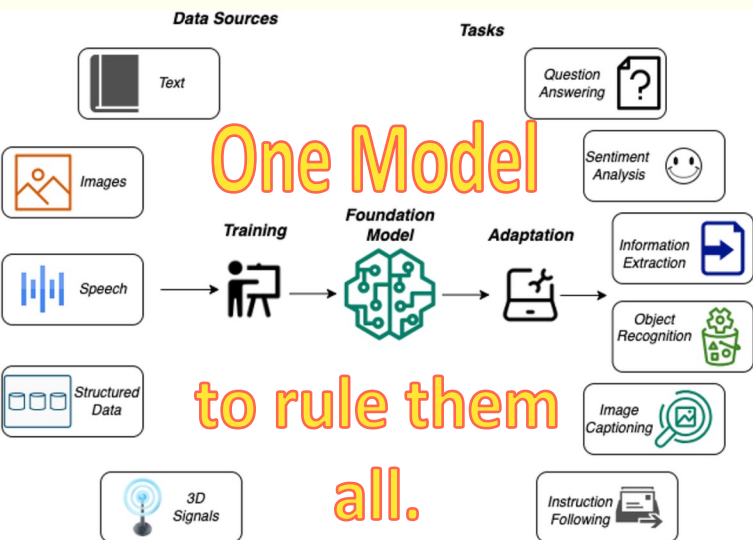
Foundation Models (FMs) for Jets:

- [OmniJet-a](#)
- [OmniLearn](#)
- [J-JEPA](#)
- [HEP-JEPA](#)
- ...

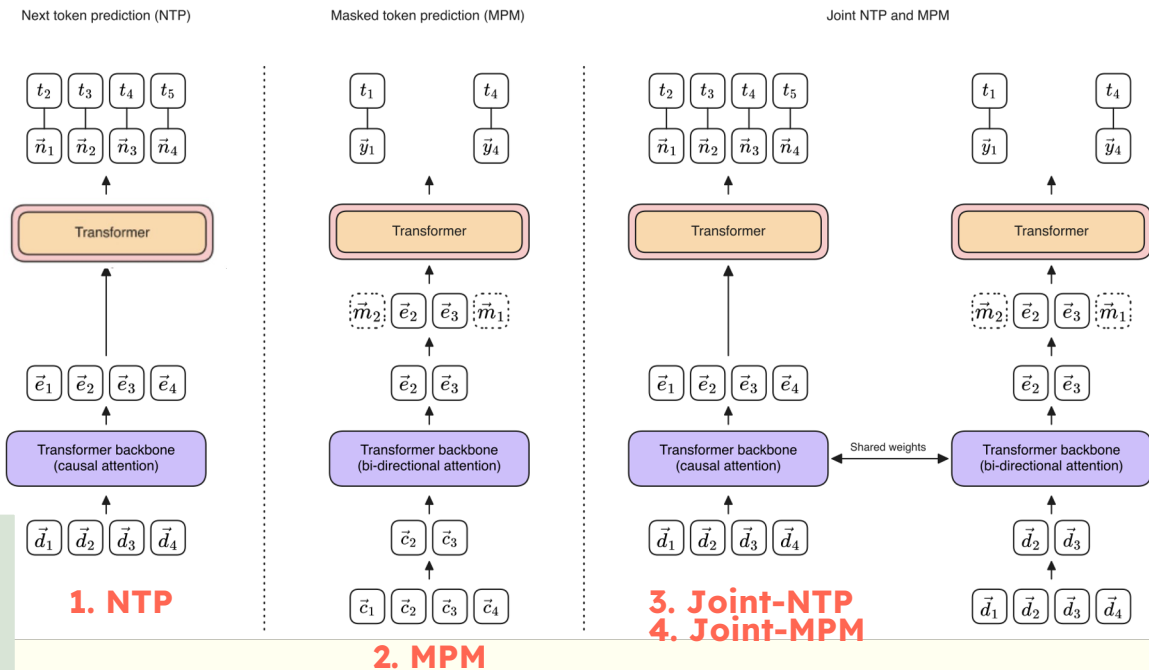
Presumably better at many downstream tasks, including jet tagging, jet generation, unfolding, anomaly detection, etc.

So, what are our good models?

9



Here we focus on the **OmniJet** family, which has **3** different pre-training objectives.



Foundation Models (FMs) for Jets:

- [OmniJet-a](#)
- [OmniLearn](#)
- [J-JEPA](#)
- [HEP-JEPA](#)
- ...

Presumably better at many downstream tasks, including jet tagging, jet generation, unfolding, anomaly detection, etc.

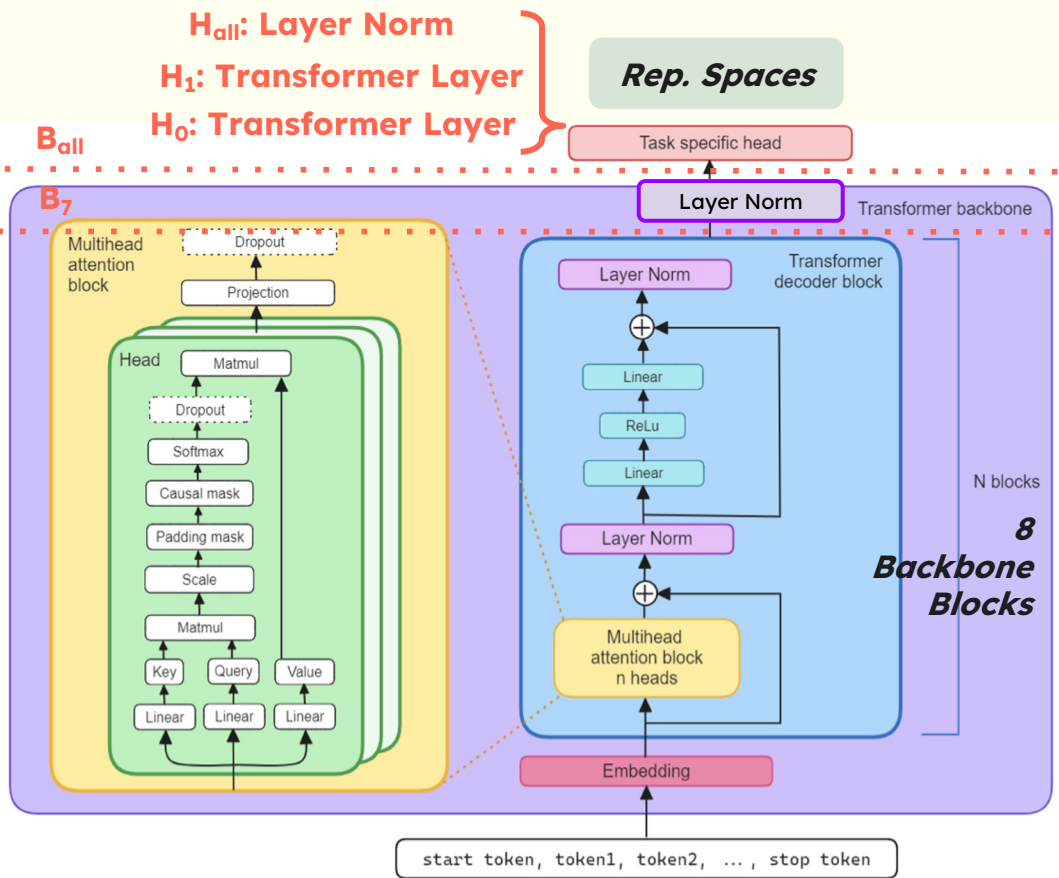
For more info about the OmniJet models, please refer to [2512.04149](#) & [2403.05618](#).

How do we extract & analyze the representation spaces?

10

How do we extract & analyze the representation spaces?

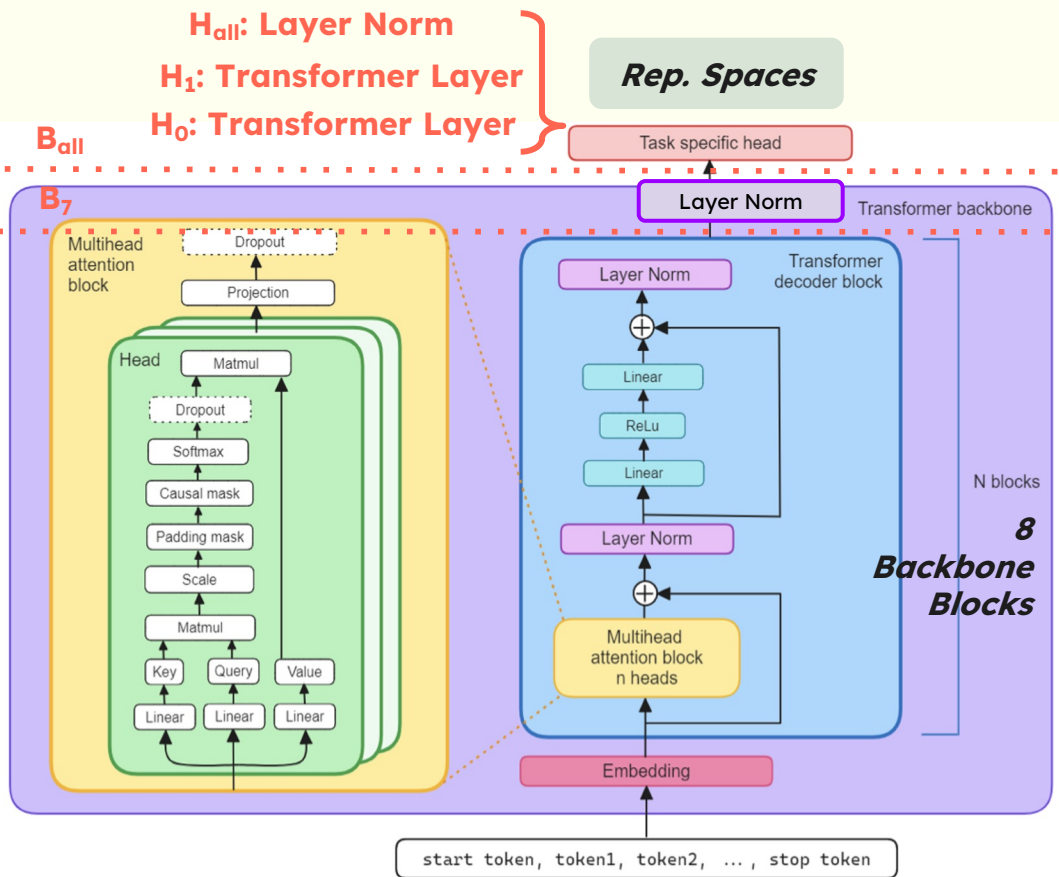
11



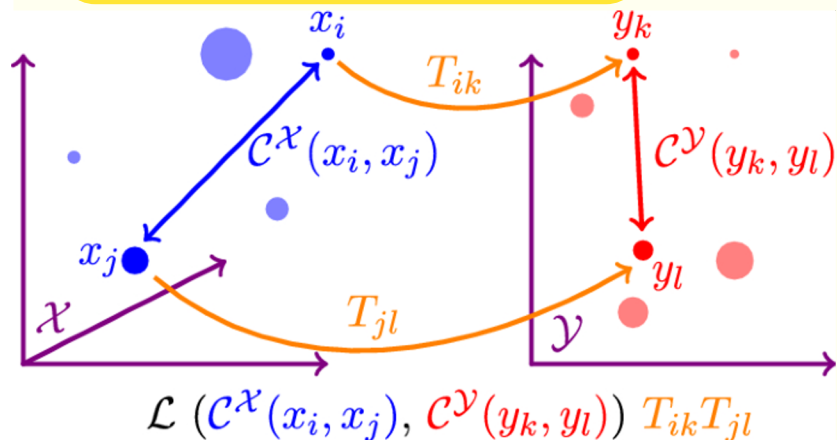
To construct rep. spaces, we extract the activation after each **backbone block (8; B_0 - B_7)** & its final **norm layer (1; B_{all})**, and each **head layer (2; H_0 - H_1)** & its final **norm layer (1; H_{all})**.

How do we extract & analyze the representation spaces?

12



The Gromov-Wasserstein Distance (GW) to compare the intrinsic geometric structure of different rep. spaces.



minimize over all (i, k)(j, l) combinations

- 500 vs 500 jets (uniform weight) at inference for rep. spaces X/Y.
- Average particle rep. vectors to get jet rep. vector.
- Use cosine distance between jet rep. vectors in the rep. space, i.e., $C(x_i, x_j) = 1 - x_i \cdot x_j / \|x_i\| \|x_j\|$. L is the square loss.

To construct rep. spaces, we extract the activation after each **backbone block** (8; B_0-B_7) & its final **norm layer** (1; B_{all}), and each **head layer** (2; H_0-H_1) & its final **norm layer** (1; H_{all}).

Note: 3 seeds per model (error bands) and 4 sets of 500 jets (too small to notice) are always used to estimate result variations.

13

Q1

How does the rep. space of the same model change w.r.t the internal model blocks at the end of training?

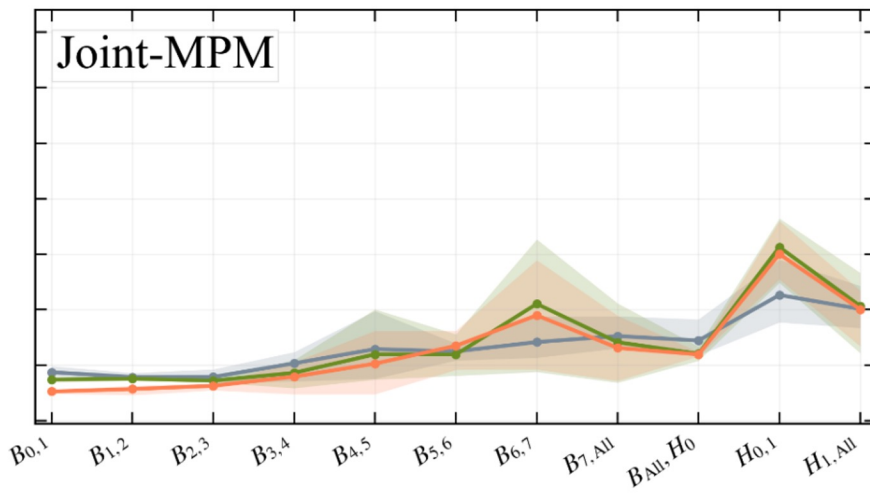
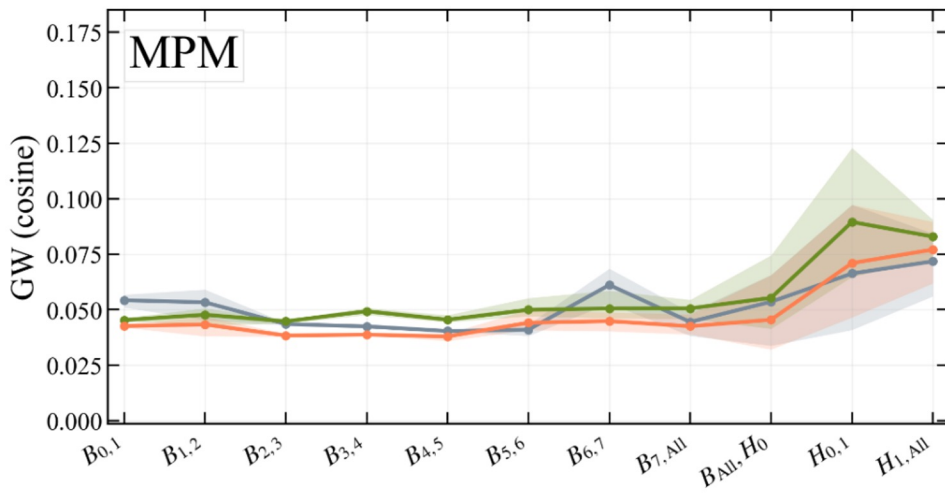
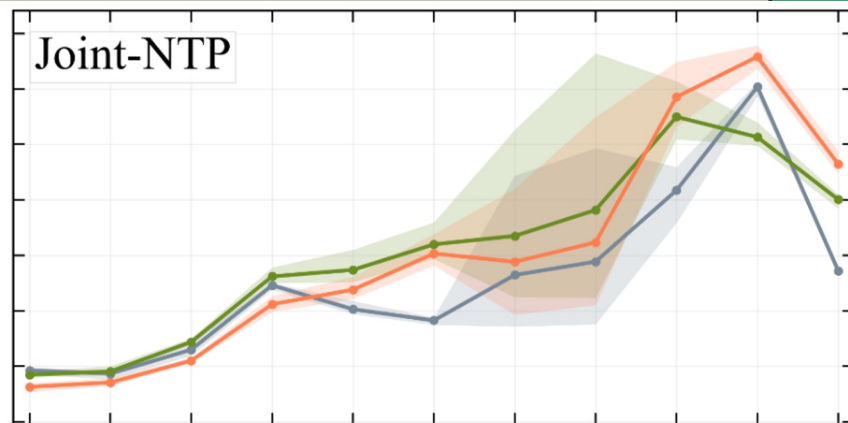
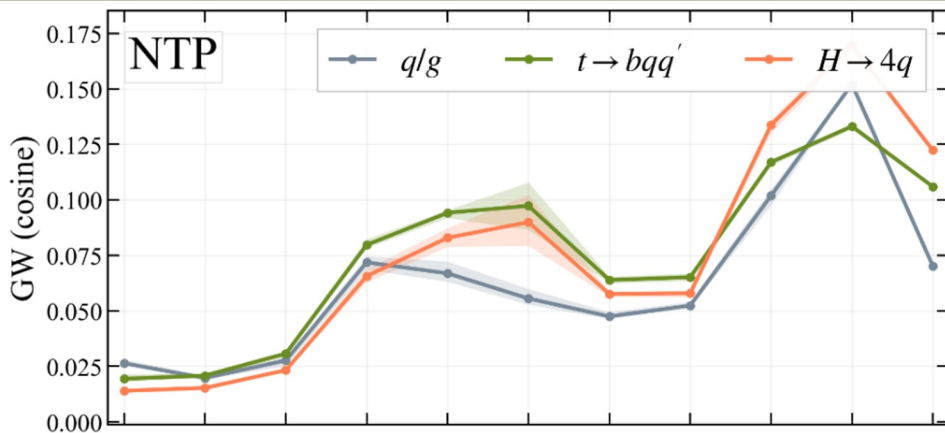
End of Training: 2.7 times the entire JetClass.

**Preliminary results only;
Proceed with caution ⚠**

Block Evolution of Same Jet Type

$$\text{GW}(\mathcal{A}^{b_i}, \mathcal{A}^{b_{i+1}})|_{s=\text{EoT}, c=\text{fix}}$$

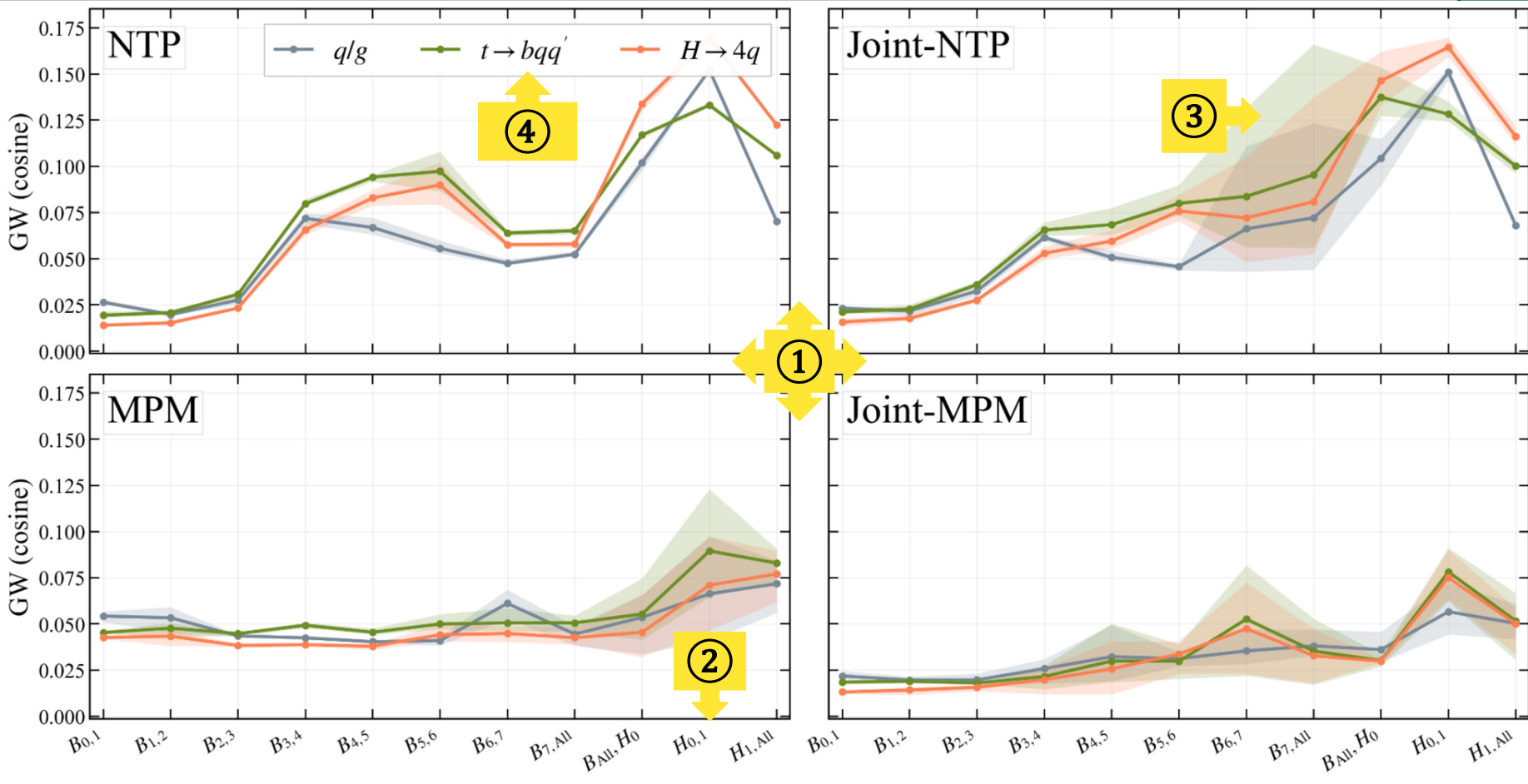
14



Block Evolution of Same Jet Type

$$\text{GW}(\mathcal{A}^{b_i}, \mathcal{A}^{b_{i+1}})|_{s=\text{EoT}, c=\text{fix}}$$

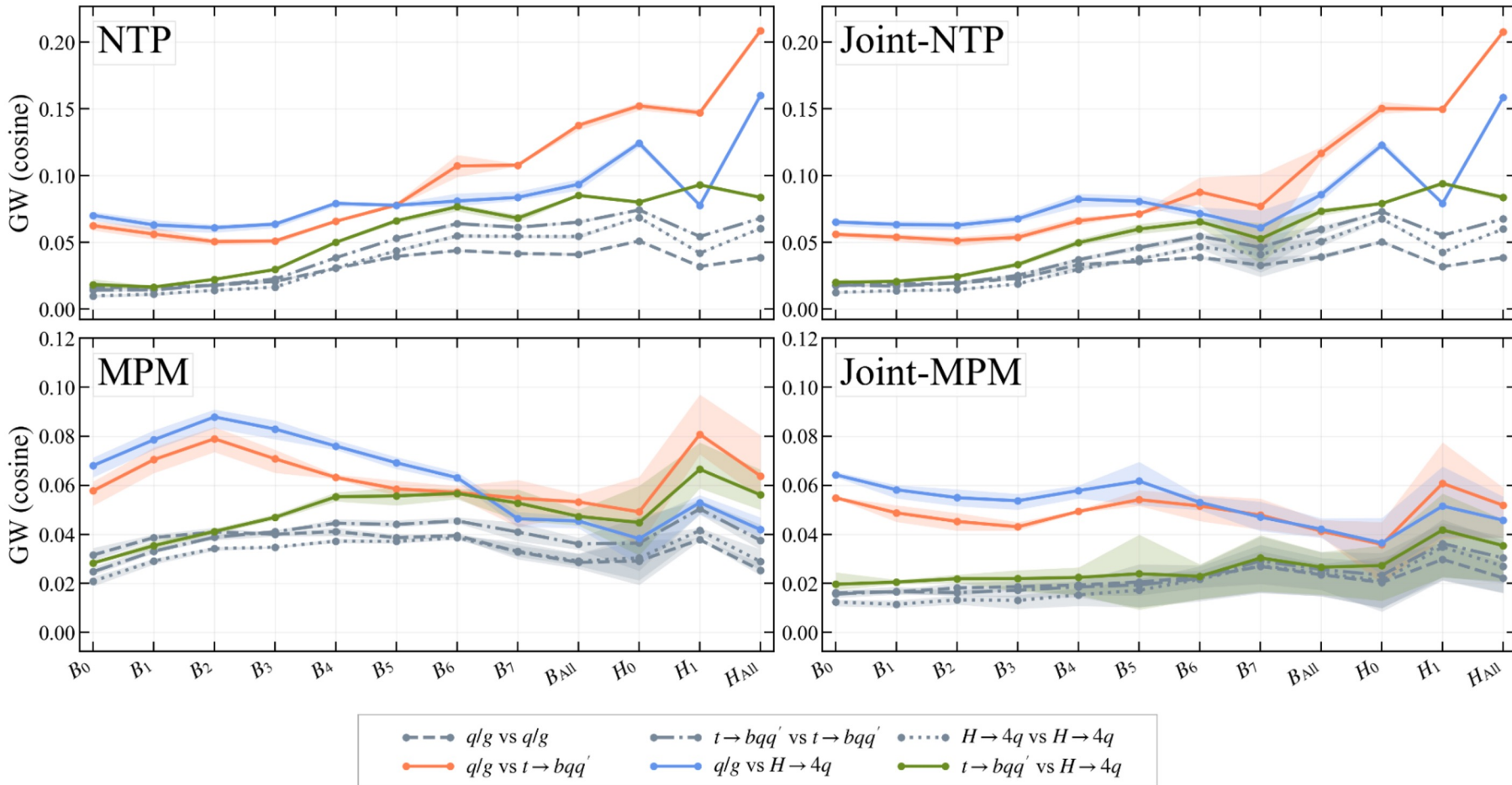
15



... of Differences between Diff Jet Types

$$\text{GW}(c_1, c_2)|_{\mathcal{A}^{b_i}, s=\text{EoT}}$$

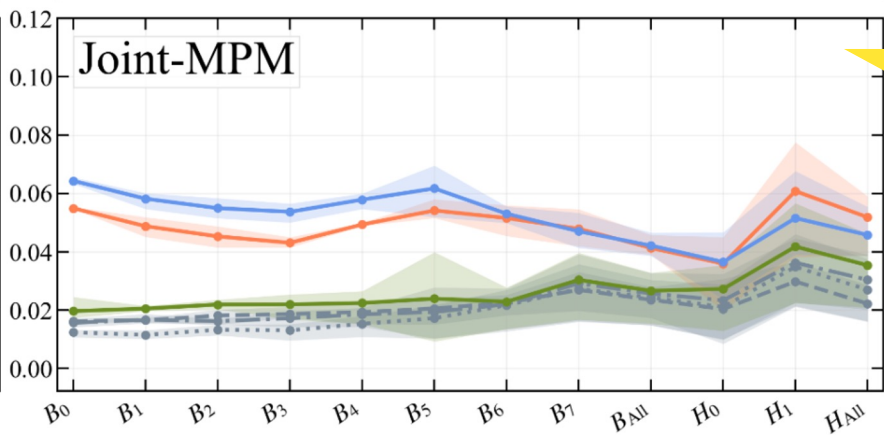
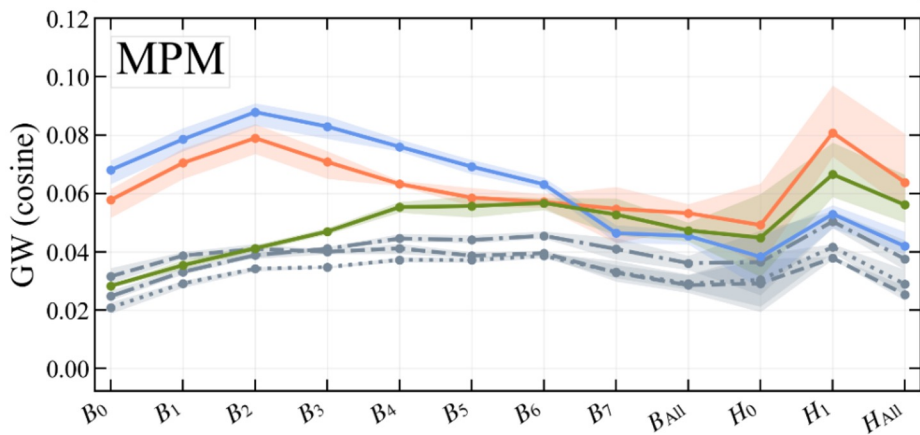
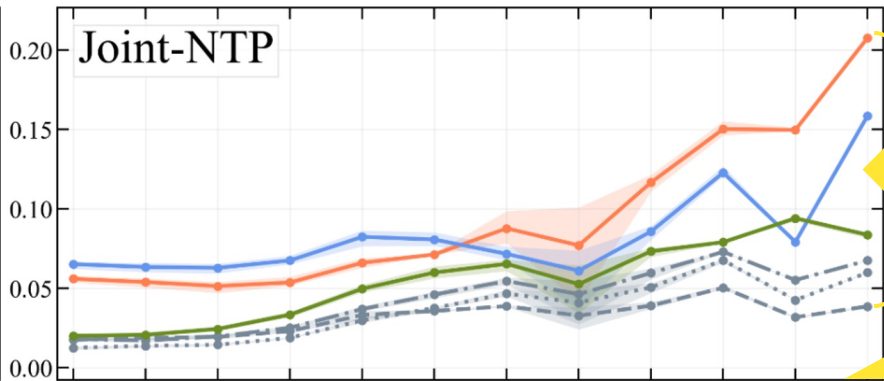
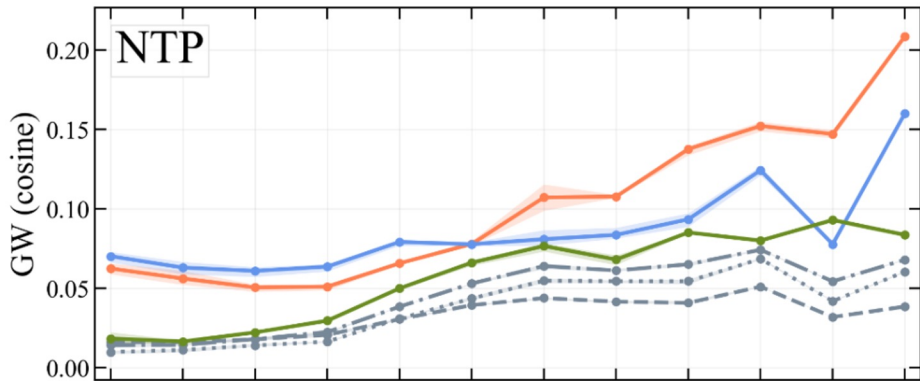
16



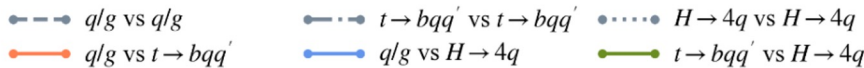
... of Differences between Diff Jet Types

$$\text{GW}(c_1, c_2)|_{\mathcal{A}^{b_i}, s=\text{EoT}}$$

17



①



①

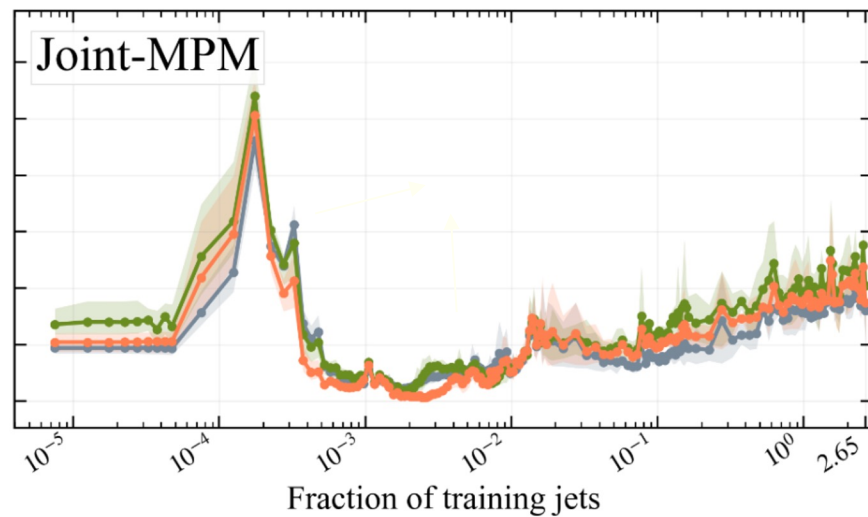
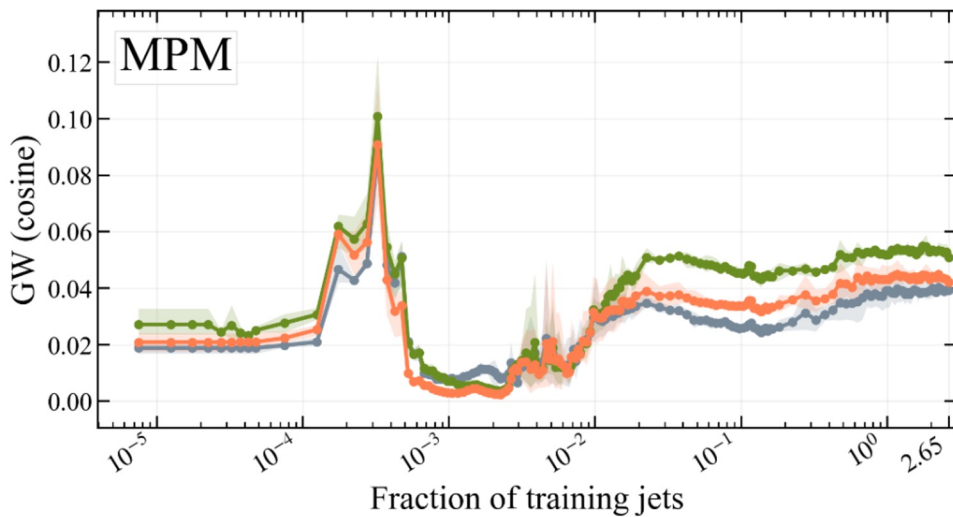
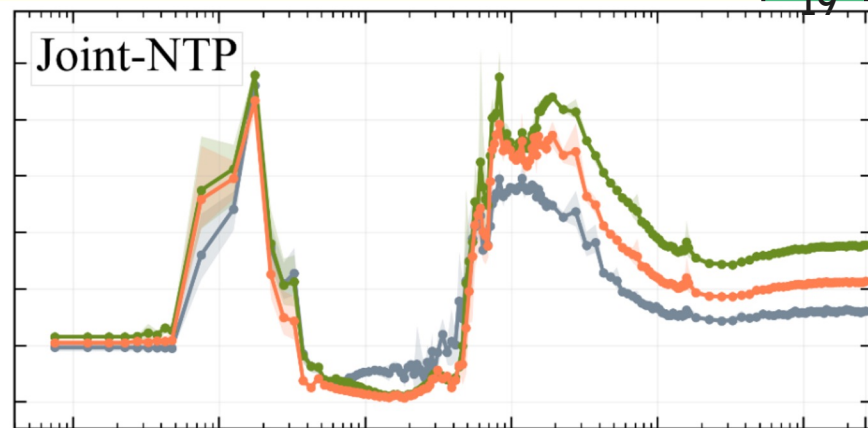
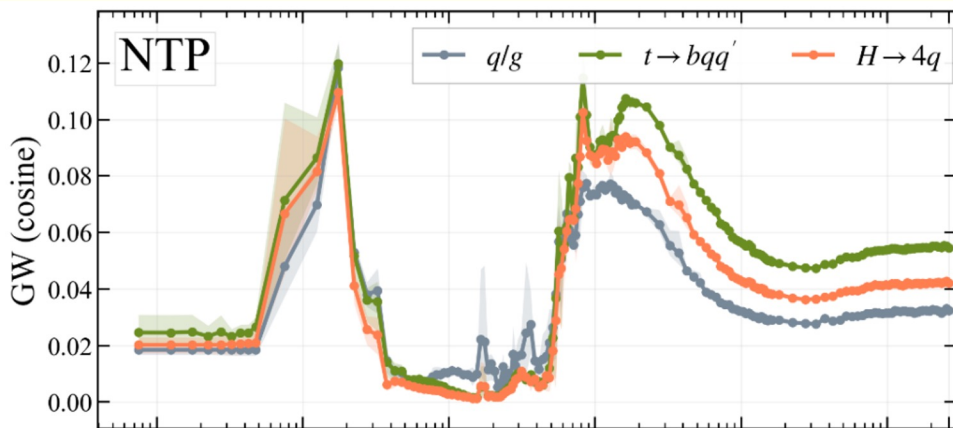
②

Q2

How does the “final” rep. space (H_1) of the same model change w.r.t training?

Training Evolution of Same Jet Type $\text{GW}(\mathcal{A}_{s_i}, \mathcal{A}_{s_{i+1}})|_{b_i=\text{final}, c=\text{fix}}$

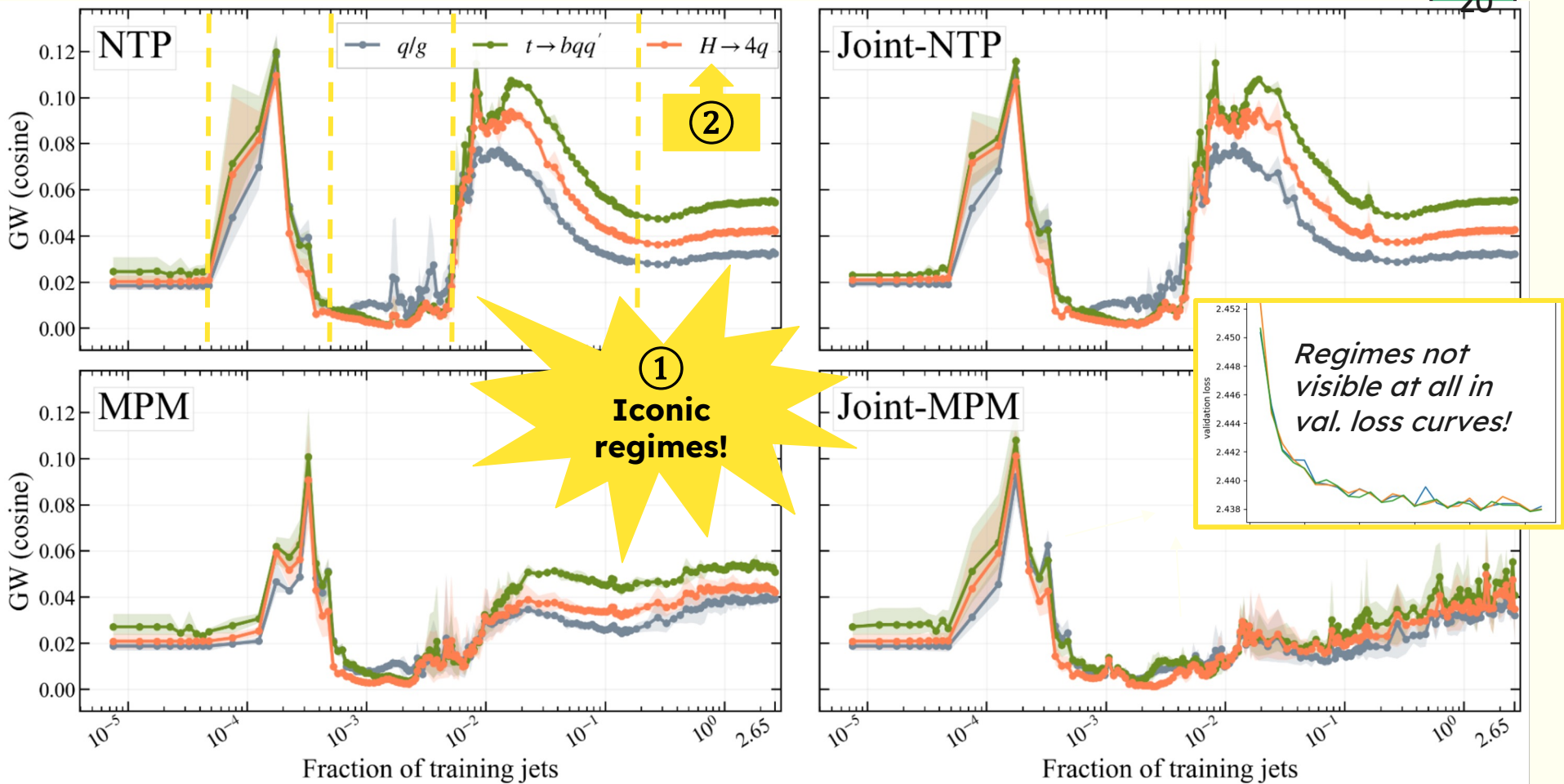
19



Training Evolution of Same Jet Type

$$\text{GW}(\mathcal{A}_{s_i}, \mathcal{A}_{s_{i+1}}) |_{b_i=\text{final}, c=\text{fix}}$$

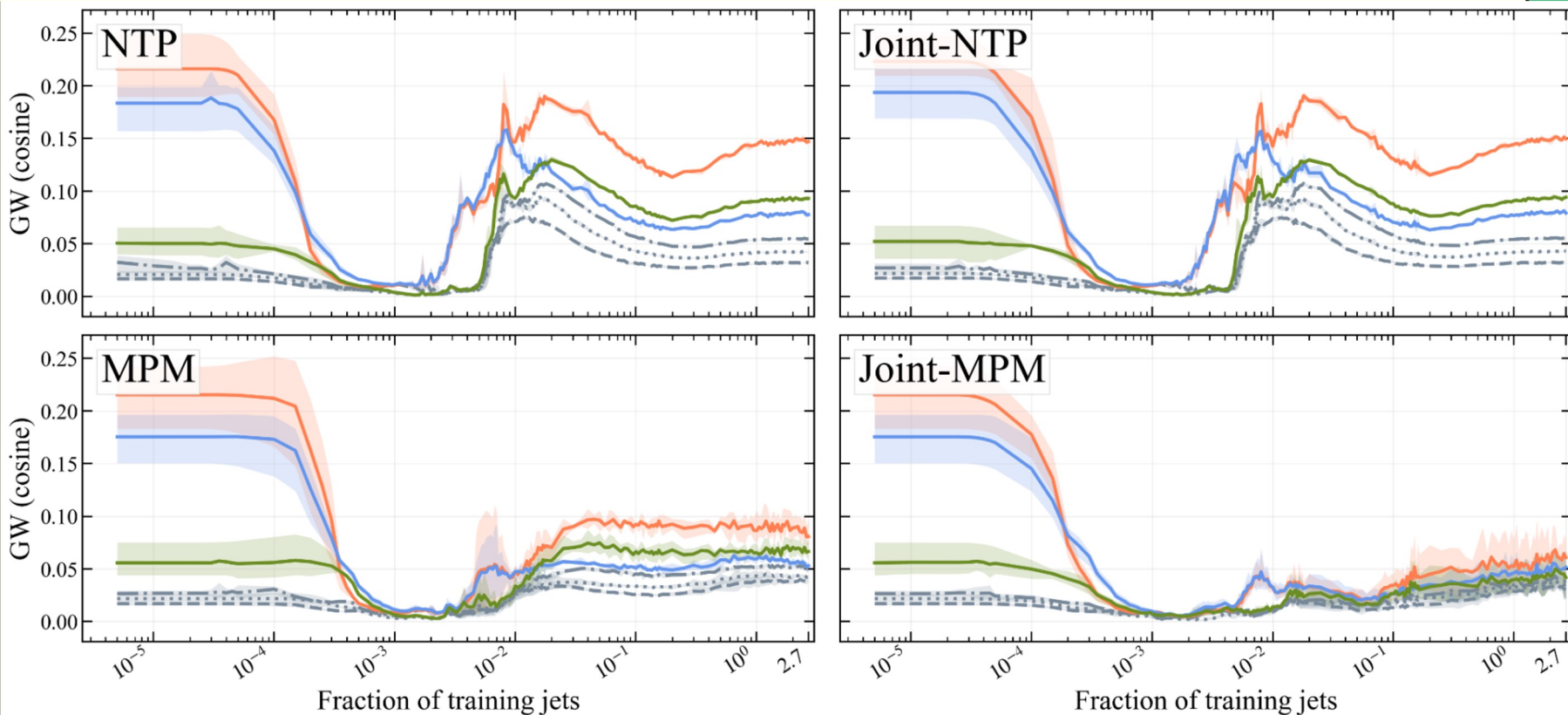
20



... of Differences between Diff Jet Types

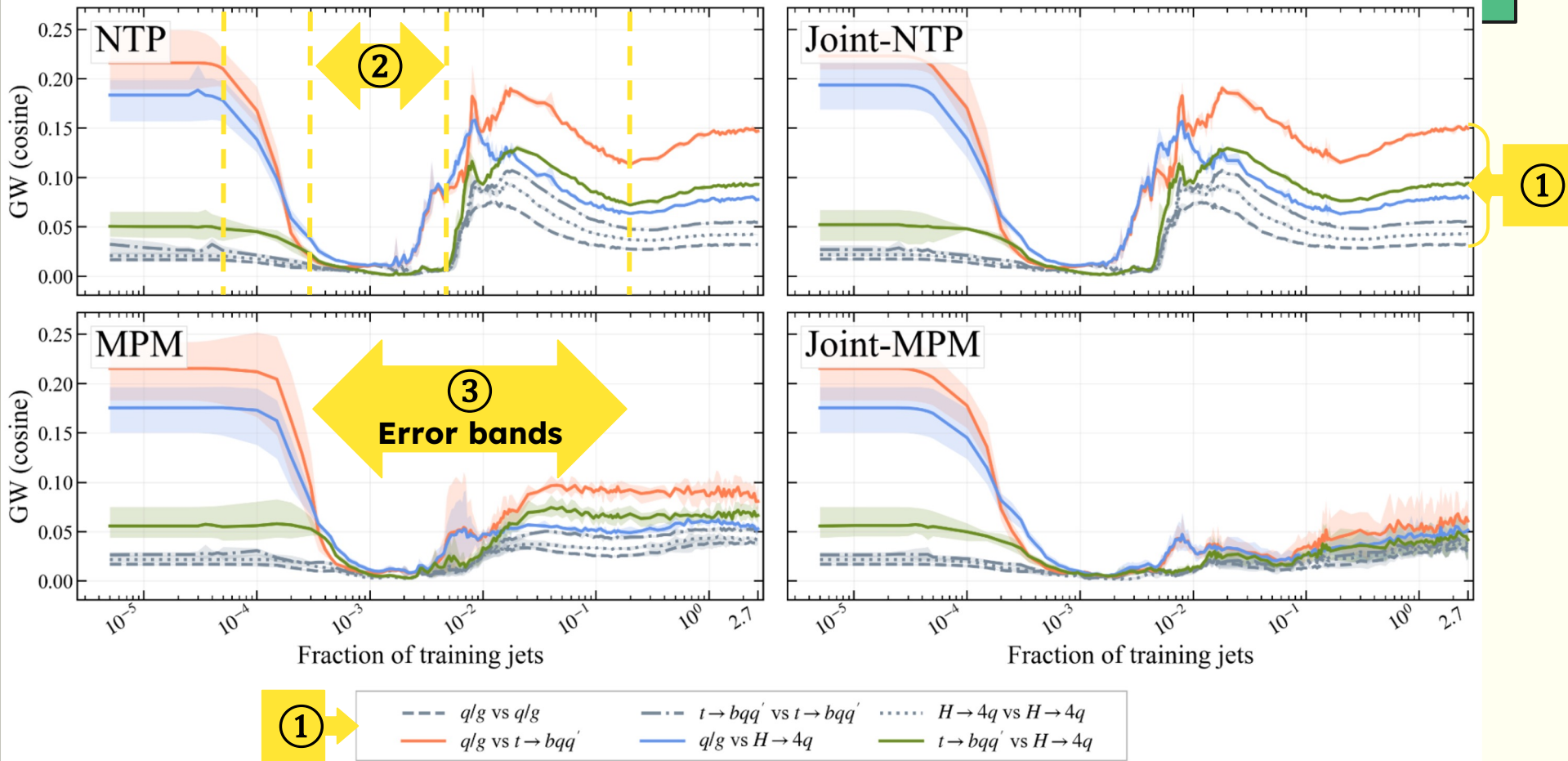
$$\text{GW}(c_1, c_2) |_{\mathcal{A}_{s_i}, b_i = b_{\text{final}}}$$

21



... of Differences between Diff Jet Types

$$GW(c_1, c_2) |_{\mathcal{A}_{s_i}, b_i=b_{\text{final}}}$$



Q3

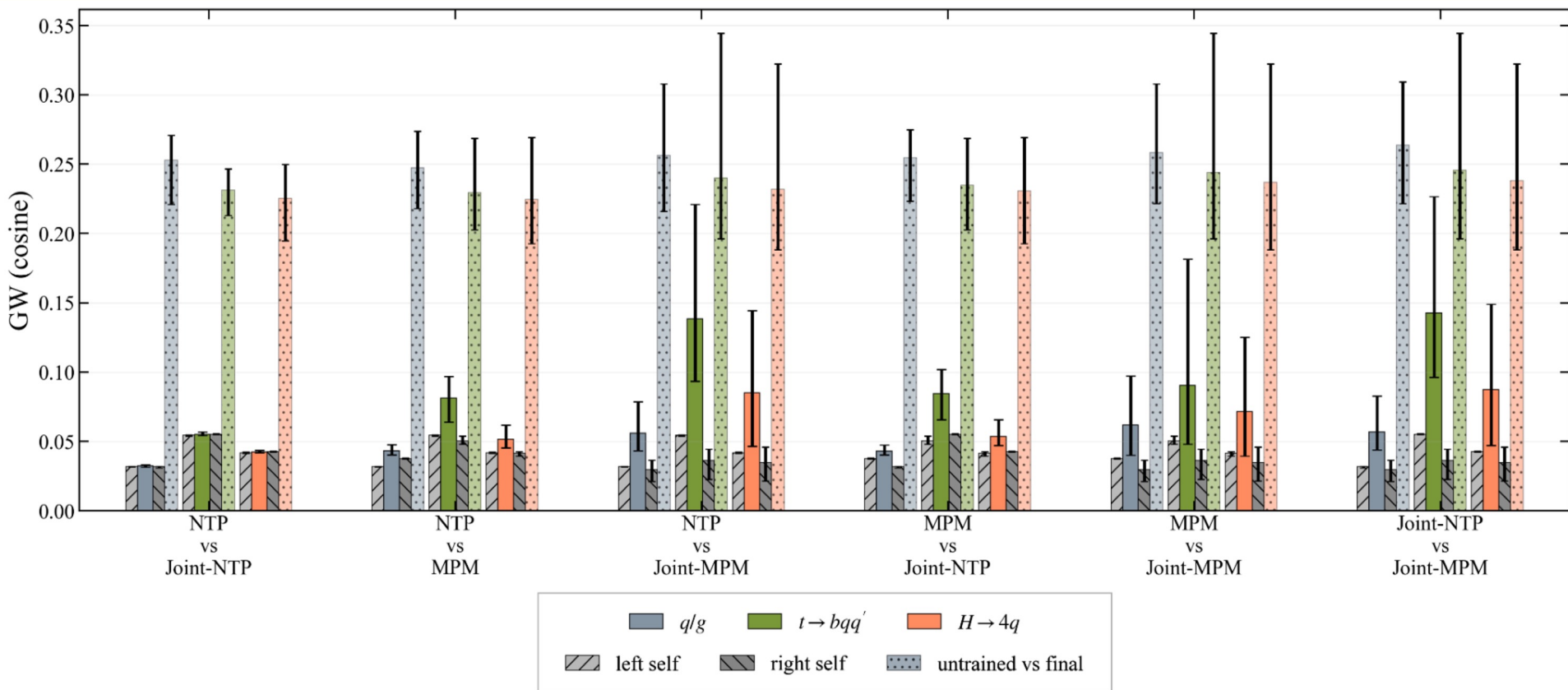
How do different models organize the same jet type in their respective rep. spaces?

$\text{GW}(\mathcal{A}, \mathcal{B})|_{b_i=\text{final}, s=\text{EoT}, c=\text{fix}}$

GW Difference between Two Models for Same jet Type

24

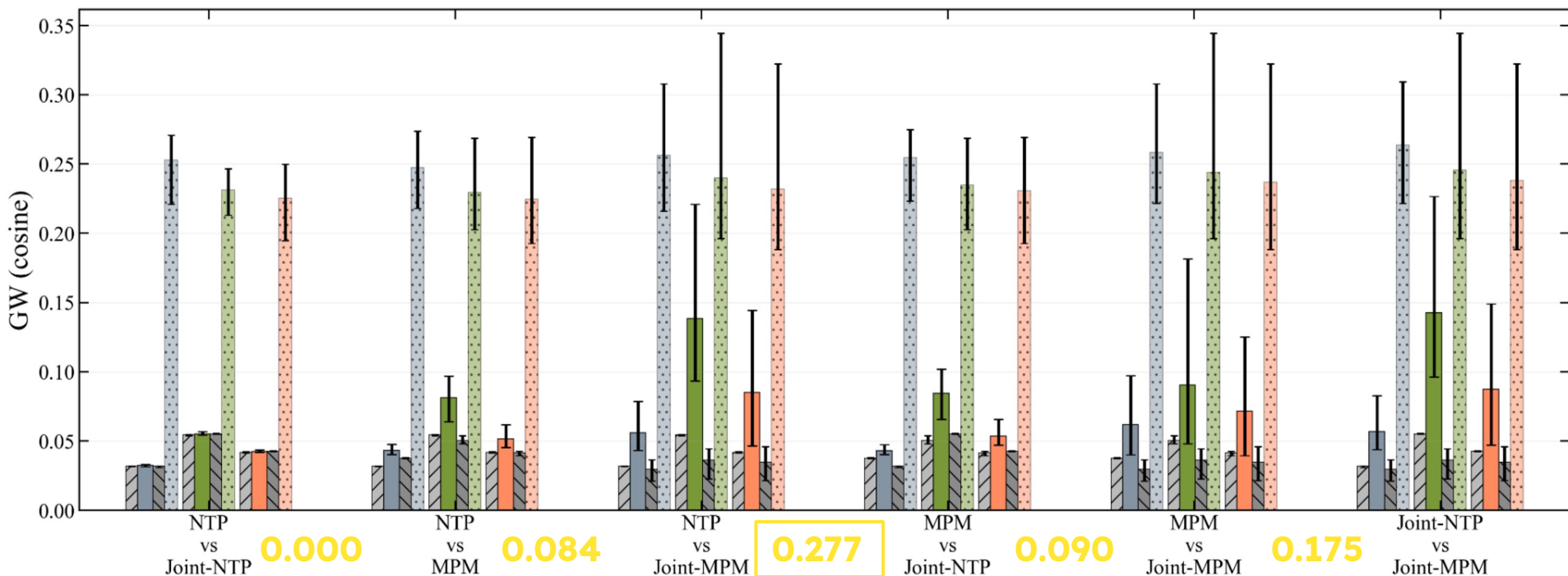
Error bars: min & max of $3 \times 3 = 9$ seed combinations ($6 \times 6 = 36$ for “*untrained vs final*”).



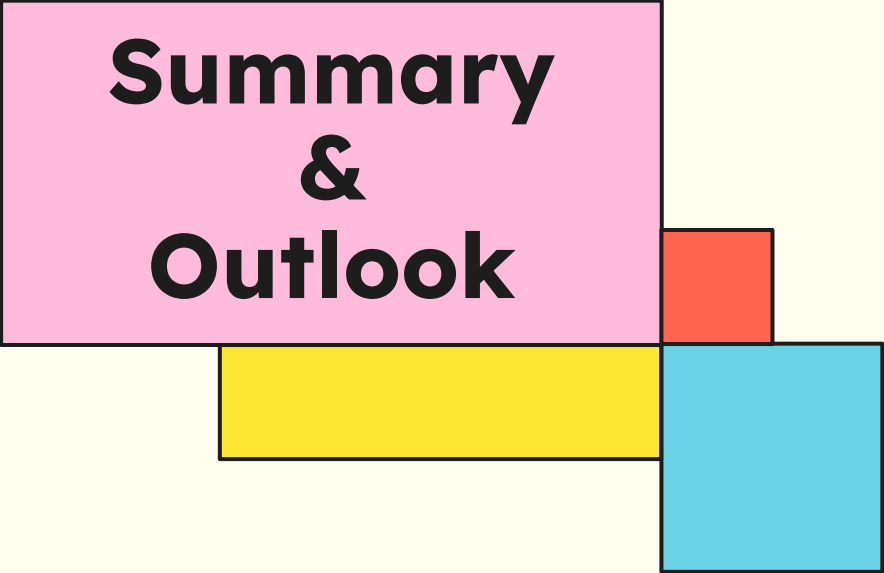
GW Difference between Two Models for Same jet Type

25

Error bars: min & max of $3 \times 3 = 9$ seed combinations ($6 \times 6 = 36$ for “*untrained vs final*”).



$$D_{\text{rel}}(A, B) = \frac{d_{\text{GW}}(A, B) - d_{\text{same}}}{d_{\text{null}} - d_{\text{same}}}$$



Summary & Outlook

(Not Yet Conclusive) Conclusions

27

- The “final” representation space of all models share similar geometric structure across jet classes. → *The Platonic Representation Hypothesis seems to be working here!* 🤔
- Geometric structure of each model is to first order the same for each jet type.
- Different jet types have larger geometric differences compared to the same jet type, with t vs qcd consistently having the highest GW.
- GW reveals a striking multi-regime structure during training, which is totally invisible in the usual validation loss curves.
- The penultimate blocks of the models seem to undergo the largest geometric rearrangement.

(Not Yet Conclusive) Conclusions

28

- The “final” representation space of all models share similar geometric structure across jet classes. → *The Platonic Representation Hypothesis seems to be working here!* 🤖
- Geometric structure of each model is to first order the same for each jet type.
- Different jet types have larger geometric differences compared to the same jet type, with t vs qcd consistently having the highest GW.
- GW reveals a striking multi-regime structure during training, which is totally invisible in the usual validation loss curves.
- The penultimate blocks of the models seem to undergo the largest geometric rearrangement.

Questions beget questions.
And so many more questions!

We are only picking up a few curious shells on the shore; the ocean remains unexplored...

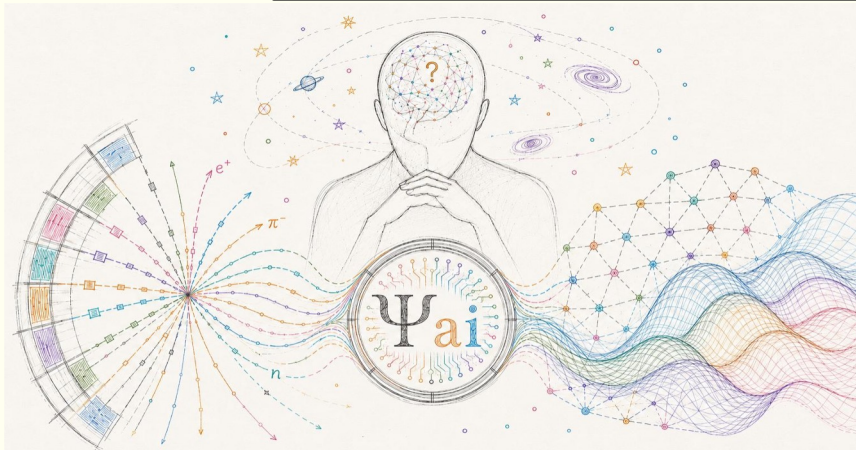


Contact Emails:

tianji_cai@tongji.edu.cn /
tianjiiresearch@gmail.com.

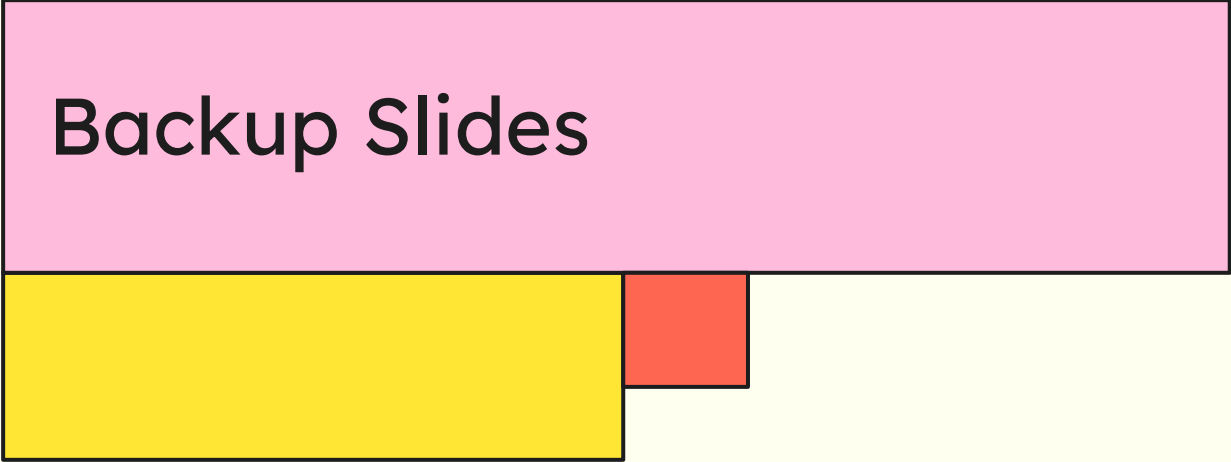
Thanks for your
attention!

Q&A?



Towards Scientific AI, Towards Physical Insight

Backup Slides



Foundation models have demonstrated remarkable performance in collider physics, yet little is understood about the geometric principles underlying their neural representations. We propose representation geometry as a new direction for *Scientific AI*, where machine learning models serve not only as powerful predictors but also as “microscopes” for revealing the intrinsic structure of the underlying physical data. In this work, we study the OmniJet family of collider foundation models pre-trained with three different self-supervised objectives on the JetClass benchmark dataset. Using the Gromov–Wasserstein distance, we compare the metric geometry of representation spaces across network blocks, training steps, and pre-training schemes. Our preliminary results uncover highly structured geometric evolution both across network depth and throughout training, while suggesting that learned representations exhibit both common structural features and systematic variations induced by pre-training objectives. These studies establish a framework for investigating whether and to what extent collider foundation models converge toward common geometric structures, laying the groundwork for future empirical tests of the Platonic Representation Hypothesis in a scientific domain whose underlying reality is governed by fundamental physical laws.

Object	Notation	Examples
Model	Curly capital letters	$\mathcal{A}, \mathcal{B}, \mathcal{C}, \dots$
Model Head	H = L for linear, T for transformer in subscript	$\mathcal{A}_L, \mathcal{A}_T, \dots$
Model Pre-training	NTP, MPM, NTP+MPM in subscript after head	$\mathcal{A}_{L\text{-NTP}}, \mathcal{A}_{T\text{-MPM}}, \dots$
Initializations	Numbers added to the head in subscript	$\mathcal{A}_{L1}, \mathcal{A}_{L2}, \mathcal{A}_{T1}, \dots$
Model Block	b_i in the superscript	$\mathcal{A}^{b_1}, \mathcal{A}^{b_2} \dots$
Training Step	s as subscript (EoT for end of training)	$\mathcal{A}_{s=1000}, \mathcal{A}_{s=10^5}, \mathcal{A}_{s=\text{EoT}} \dots$
Inference Data	Lower letter c ($c_1, c_2, \dots$)	$c = \text{JetClass}, \text{QCD}, \text{T}, \dots$

Model Architecture

NTP			
#	Name	Type	Params
0	backbone	Backbone/Transformer	1.6 M
1	backbone.input_projection	Project/Add	640
2	backbone.input_projection.embed_part	Linear	512
3	backbone.input_projection.embed_jet	Linear	128
4	backbone.transformer	Transformer	1.6 M
5	backbone.transformer.blocks	ModuleList	1.6 M
6	backbone.transformer.blocks.0	EncoderBlock	198 K
9	backbone.transformer.blocks.1	EncoderBlock	198 K
12	backbone.transformer.blocks.2	EncoderBlock	198 K
15	backbone.transformer.blocks.3	EncoderBlock	198 K
18	backbone.transformer.blocks.4	EncoderBlock	198 K
21	backbone.transformer.blocks.5	EncoderBlock	198 K
24	backbone.transformer.blocks.6	EncoderBlock	198 K
27	backbone.transformer.blocks.7	EncoderBlock	198 K
30	backbone.transformer.final_norm	LayerNorm	256
31	head_gen	TokenPredictionHead	1.3 M
32	head_gen.transformer_blocks	ModuleList	265 K
34	head_gen.transformer_blocks.blocks	ModuleList	265 K
35	head_gen.transformer_blocks.blocks.0	EncoderBlock	132 K
36	head_gen.transformer_blocks.blocks.1	EncoderBlock	132 K
37	head_gen.transformer_blocks.final_norm	LayerNorm	256
38	head_gen.unembedding_mlp	Identity	0
39	head_gen.unembedding_matrix	Linear	1.1 M

MPM			
#	Name	Type	Params
0	backbone	Backbone/Transformer	1.6 M
1	backbone.input_projection	Project/Add	640
2	backbone.input_projection.embed_part	Linear	512
3	backbone.input_projection.embed_jet	Linear	128
4	backbone.transformer	Transformer	1.6 M
5	backbone.transformer.blocks	ModuleList	1.6 M
6	backbone.transformer.blocks.0	EncoderBlock	198 K
9	backbone.transformer.blocks.1	EncoderBlock	198 K
12	backbone.transformer.blocks.2	EncoderBlock	198 K
15	backbone.transformer.blocks.3	EncoderBlock	198 K
18	backbone.transformer.blocks.4	EncoderBlock	198 K
21	backbone.transformer.blocks.5	EncoderBlock	198 K
24	backbone.transformer.blocks.6	EncoderBlock	198 K
27	backbone.transformer.blocks.7	EncoderBlock	198 K
30	backbone.transformer.final_norm	LayerNorm	256
31	head_mpm	TokenPredictionHead	1.3 M
32	head_mpm.transformer_blocks	ModuleList	265 K
34	head_mpm.transformer_blocks.blocks	ModuleList	265 K
35	head_mpm.transformer_blocks.blocks.0	EncoderBlock	132 K
36	head_mpm.transformer_blocks.blocks.1	EncoderBlock	132 K
37	head_mpm.transformer_blocks.final_norm	LayerNorm	256
38	head_mpm.unembedding_mlp	Identity	0
39	head_mpm.unembedding_matrix	Linear	1.1 M

shared

separate

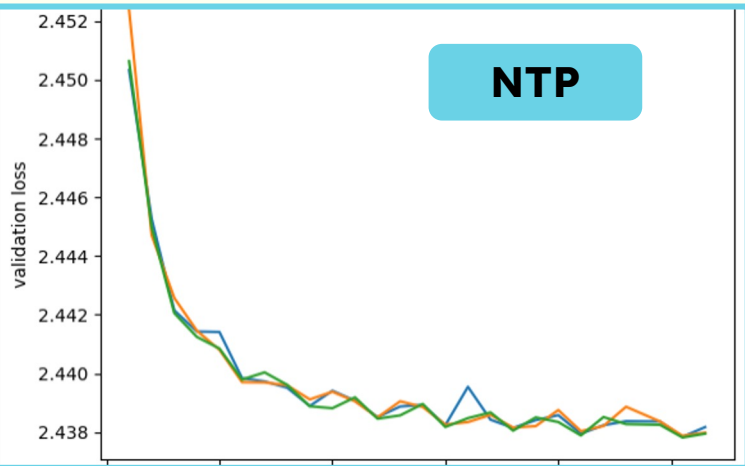
separate

Joint			
#	Name	Type	Params
0	backbone	Backbone/Transformer	1.6 M
1	backbone.input_projection	Project/Add	640
2	backbone.input_projection.embed_part	Linear	512
3	backbone.input_projection.embed_jet	Linear	128
4	backbone.transformer	Transformer	1.6 M
5	backbone.transformer.blocks	ModuleList	1.6 M
6	backbone.transformer.blocks.0	EncoderBlock	198 K
9	backbone.transformer.blocks.1	EncoderBlock	198 K
12	backbone.transformer.blocks.2	EncoderBlock	198 K
15	backbone.transformer.blocks.3	EncoderBlock	198 K
18	backbone.transformer.blocks.4	EncoderBlock	198 K
21	backbone.transformer.blocks.5	EncoderBlock	198 K
24	backbone.transformer.blocks.6	EncoderBlock	198 K
27	backbone.transformer.blocks.7	EncoderBlock	198 K
30	backbone.transformer.final_norm	LayerNorm	256
31	head_mpm	TokenPredictionHead	1.3 M
32	head_mpm.transformer_blocks	ModuleList	265 K
34	head_mpm.transformer_blocks.blocks	ModuleList	265 K
35	head_mpm.transformer_blocks.blocks.0	EncoderBlock	132 K
36	head_mpm.transformer_blocks.blocks.1	EncoderBlock	132 K
37	head_mpm.transformer_blocks.final_norm	LayerNorm	256
38	head_mpm.unembedding_mlp	Identity	0
39	head_mpm.unembedding_matrix	Linear	1.1 M
40	head_gen	TokenPredictionHead	1.3 M
41	head_gen.transformer_blocks	ModuleList	265 K
43	head_gen.transformer_blocks.blocks	ModuleList	265 K
44	head_gen.transformer_blocks.blocks.0	EncoderBlock	132 K
45	head_gen.transformer_blocks.blocks.1	EncoderBlock	132 K
46	head_gen.transformer_blocks.final_norm	LayerNorm	256
47	head_gen.unembedding_mlp	Identity	0
48	head_gen.unembedding_matrix	Linear	1.1 M

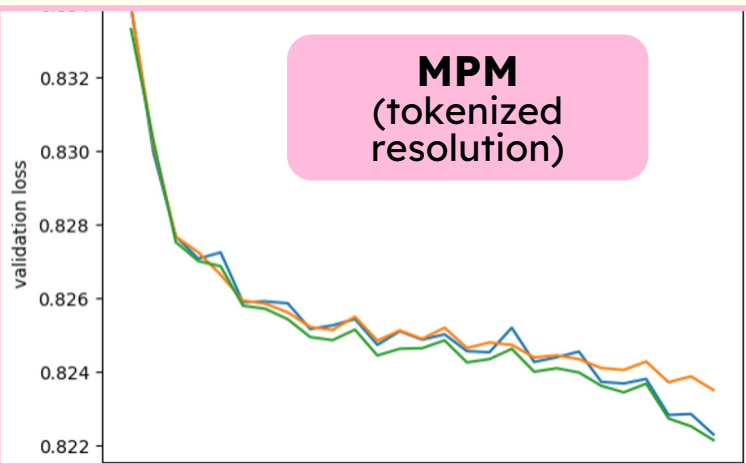
Validation Loss Curves During Training

34

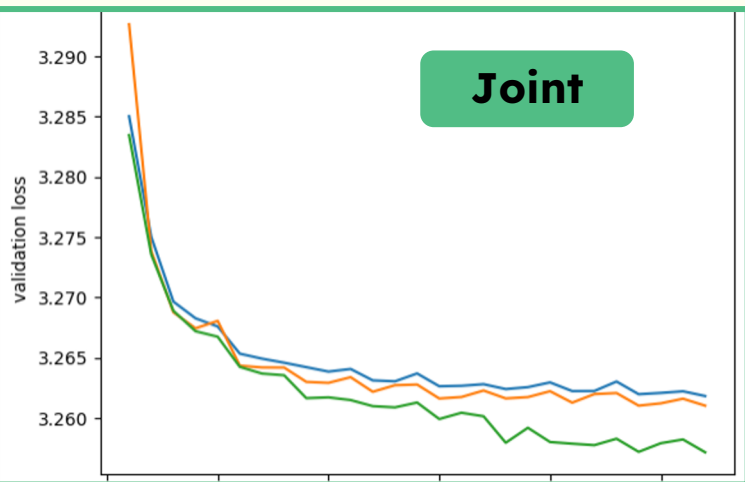
NTP



MPM
(tokenized
resolution)



Joint



Adaptive Checkpointing Schedule:

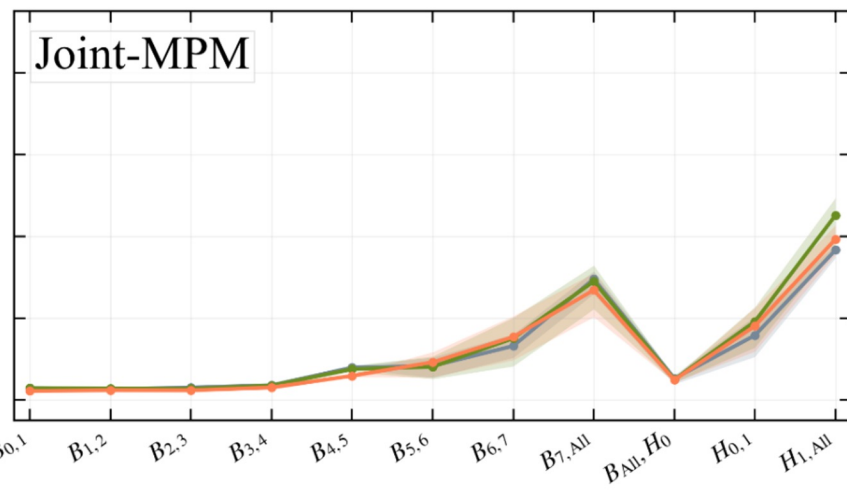
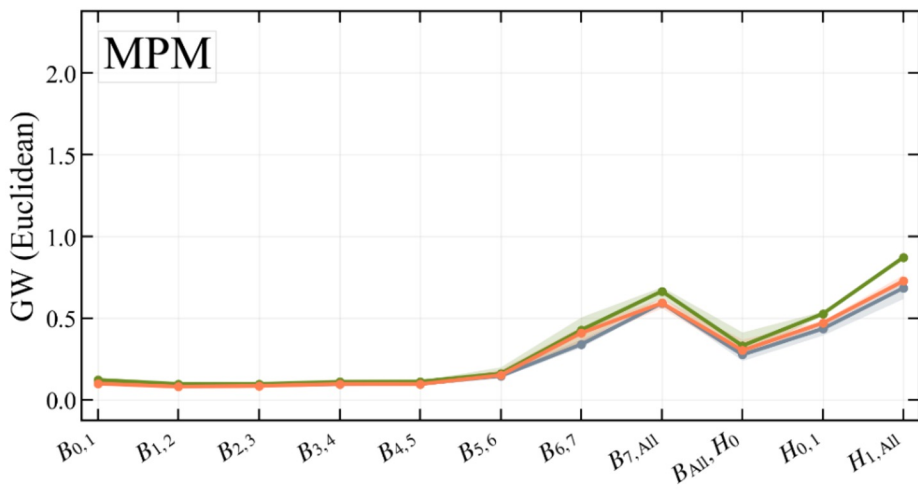
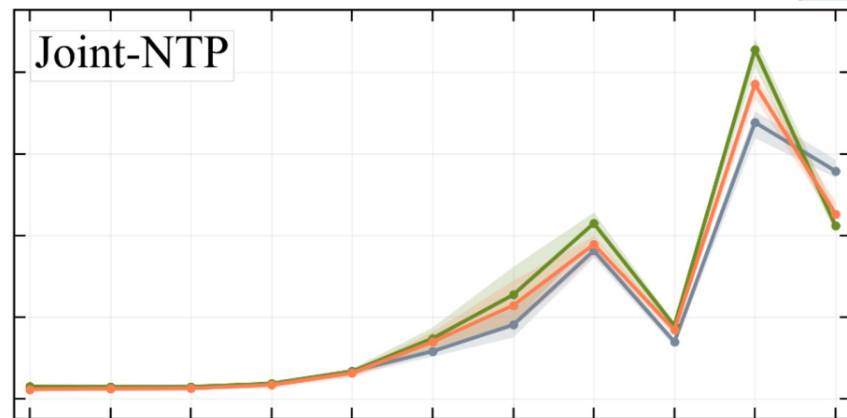
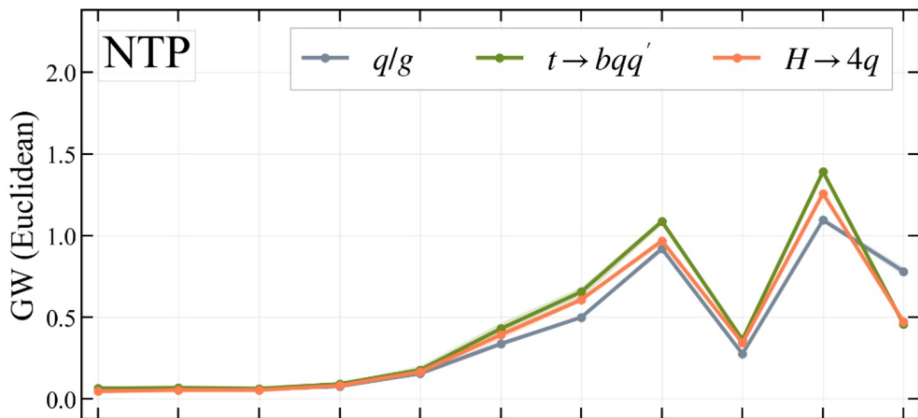
- Step 0-10 → Every 1 step: 10 ckpts, 500 jets/ckpt.
- Step 10-200 → Every 10 steps: 20 ckpts, 5k jets/ckpt.
- Step 200-600 → Every 20 steps: 20 ckpts, 10k jets/ckpt.
- Step 600-1600 → Every 50 steps: 20 ckpts; 20k jets/ckpt.
- Step 1600-3600 → Every 100 steps: 20 ckpts; 50k jets/ckpt.
- Step 3600-33600 → Every 1000 steps: 30 ckpts; 500k jets/ckpt.
- Step 33600-333.6k → Every 10k steps: 30 ckpts; 5M jets/ckpt.
- Step 333.6k-540k → Every 20k steps: 10M jets/ckpt

540k steps means 266,705,000 jets. JetClass has 100M training jets, meaning that each jet is seen ~2.7 times.

Block Evolution of Same Jet Type

$$\text{GW}(\mathcal{A}^{b_i}, \mathcal{A}^{b_{i+1}})|_{s=\text{EoT}, c=\text{fix}}$$

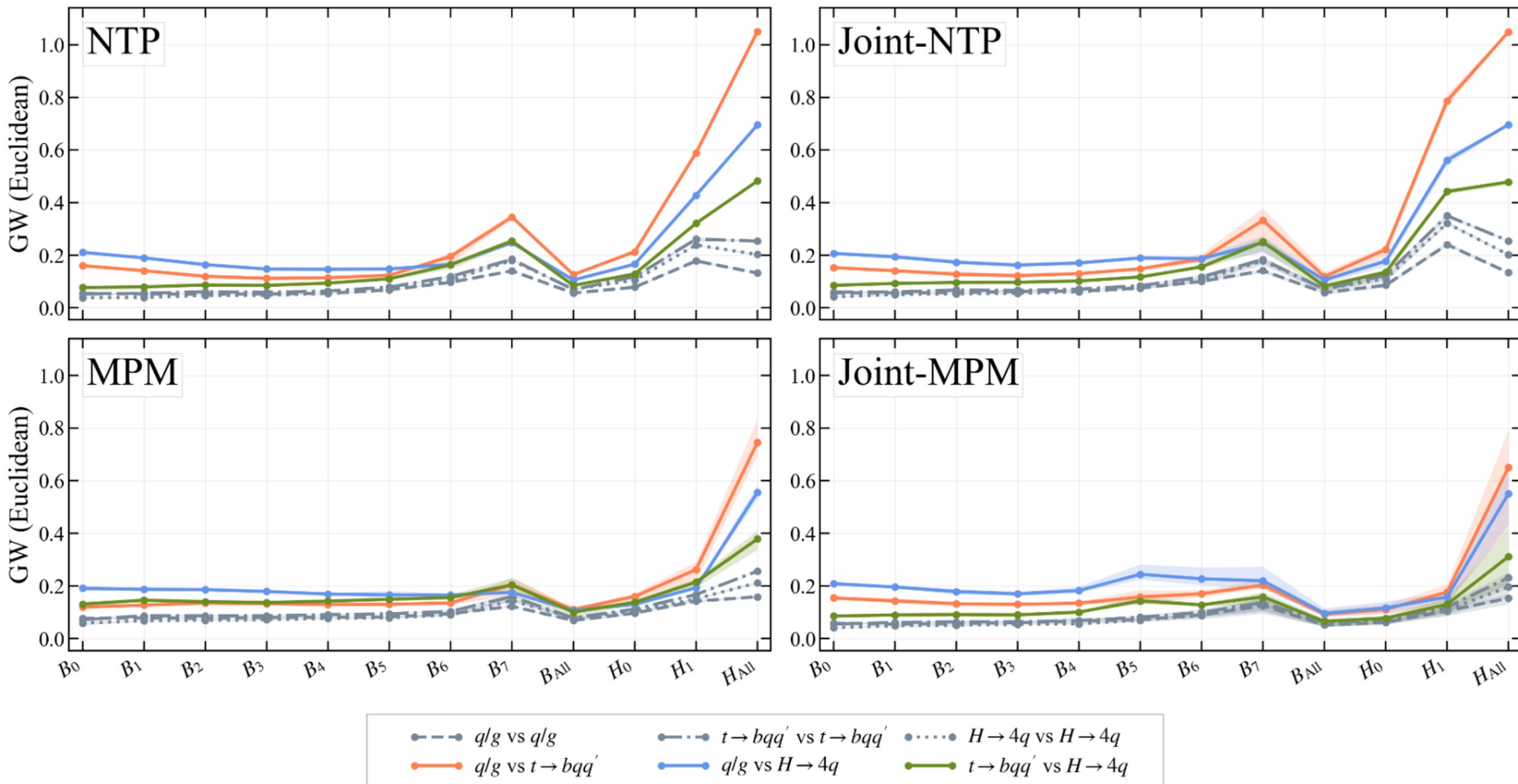
35



... of Differences between Diff Jet Types

$$\text{GW}(c_1, c_2)|_{\mathcal{A}^{b_i}, s=\text{EoT}}$$

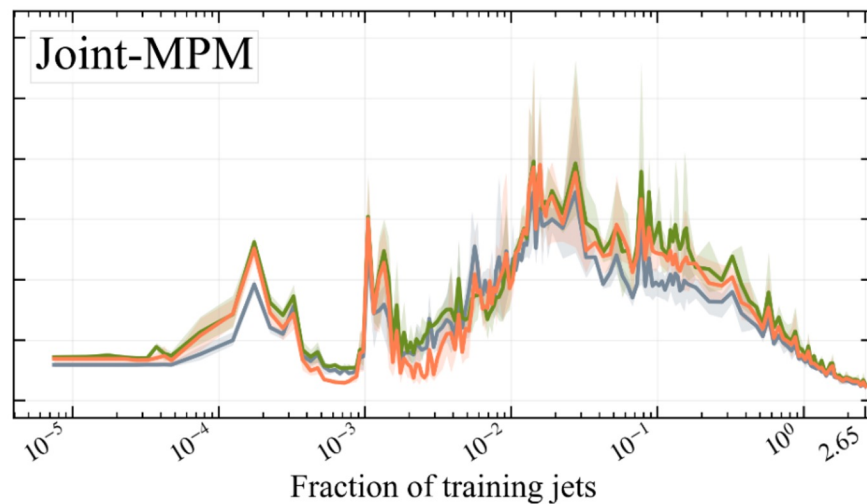
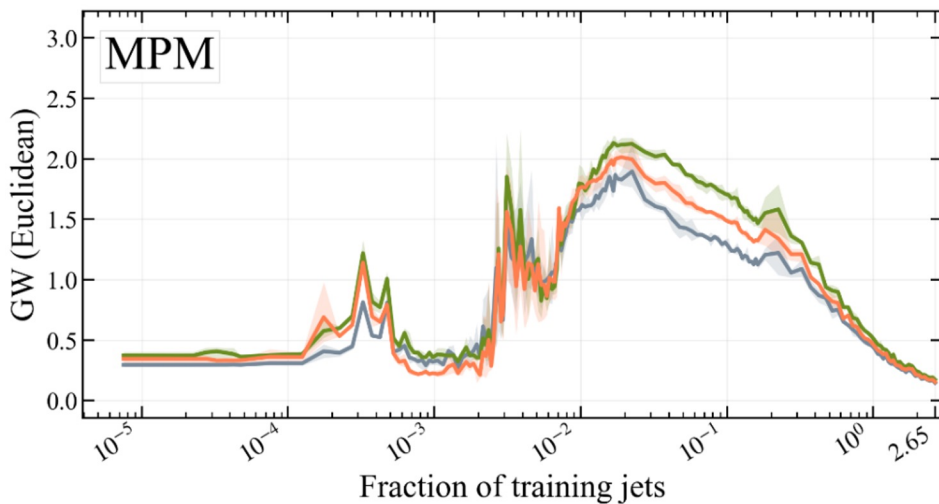
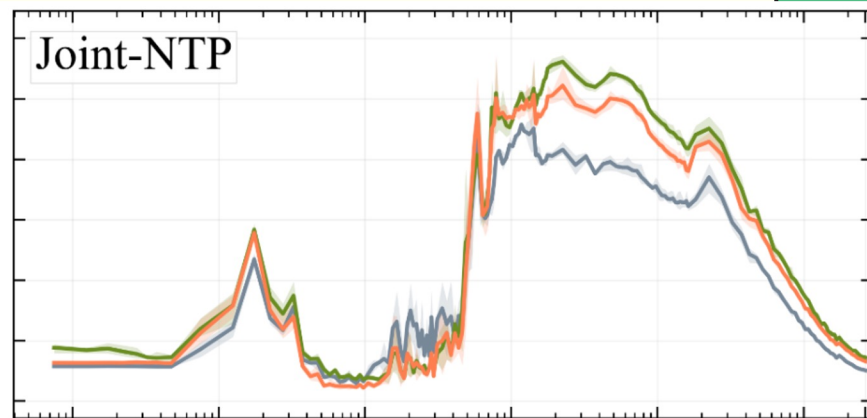
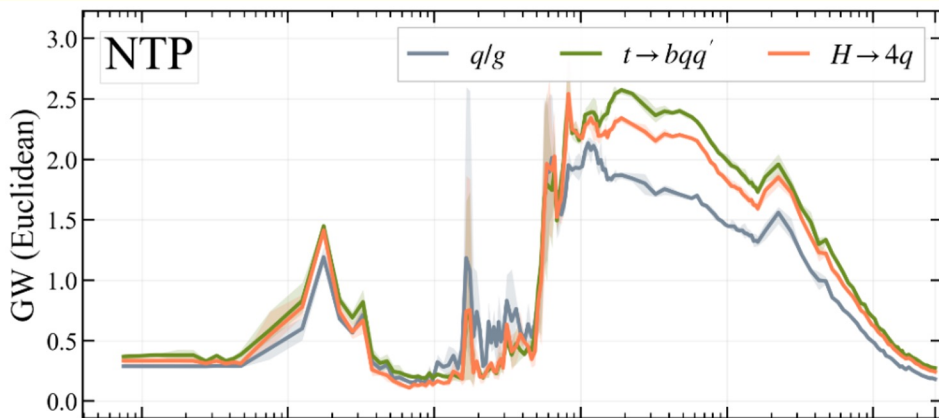
36



Training Evolution of Same Jet Type

$$\text{GW}(\mathcal{A}_{s_i}, \mathcal{A}_{s_{i+1}}) |_{b_i=\text{final}, c=\text{fix}}$$

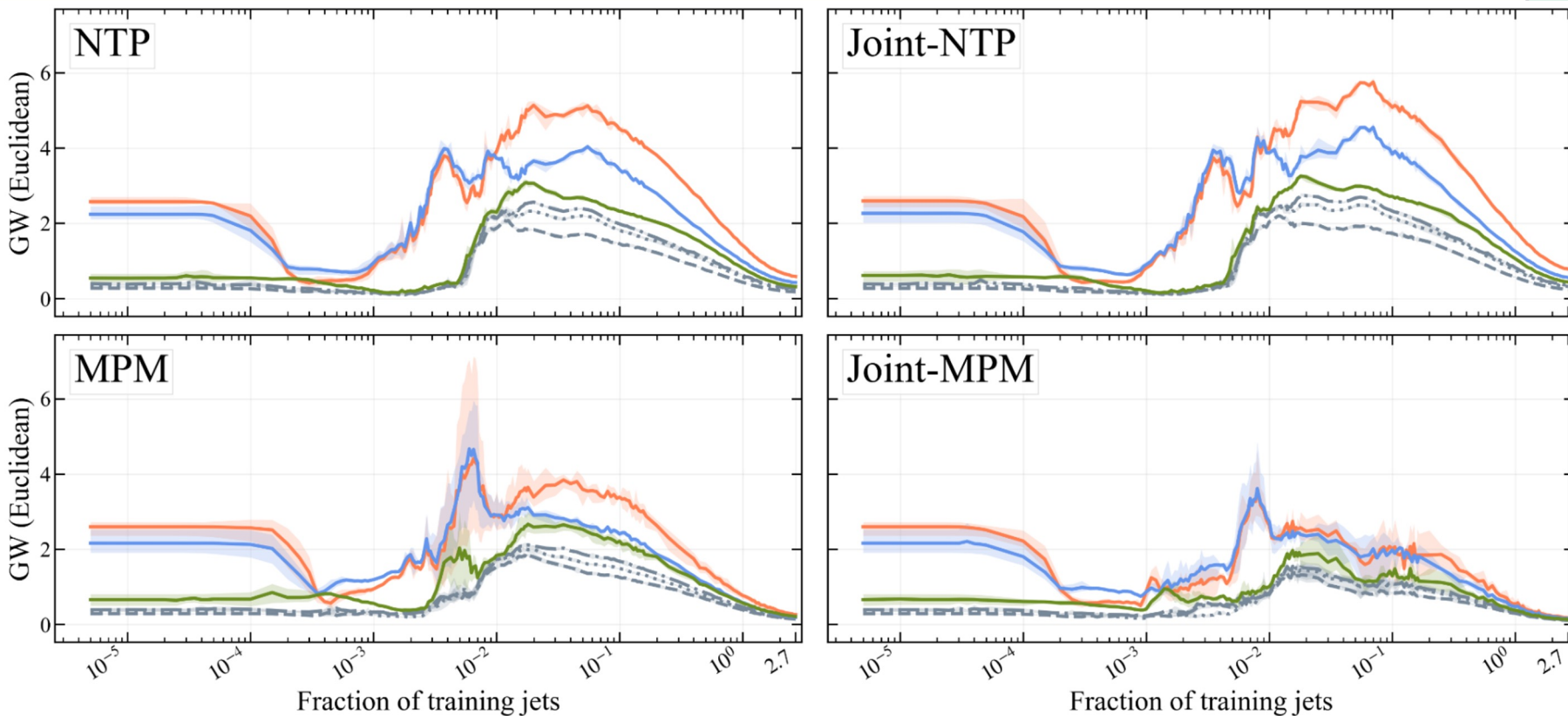
37



... of Differences between Diff Jet Types

$$GW(c_1, c_2)|_{\mathcal{A}^{b_i}, s=\text{EoT}}$$

38



GW Difference between Two Models for Same jet Type

