



Contribution ID: 32

Type: **Standard Talk**

Go small then go home - hyperparameter transfer for ML in HEP

Thursday 4 September 2025 12:10 (20 minutes)

Tuning hyperparameters of ML models, especially large ML models, can be time consuming and computationally expensive. As a potential solution, several recent papers have explored hyperparameter transfer. Under certain conditions, the optimal hyperparameters of a small model are also optimal for larger models. One can therefore tune only the small model and transfer the hyperparameters to the larger model, saving a large amount of time and effort. This work explores how well the idea holds up in high-energy physics by applying it to three existing ML pipelines: metric learning for particle tracking, autoencoders for anomaly detection, and particle transformers for jet tagging. These cover several common ML architectures and reflect models currently used or in development at CMS and other experiments. We show that with a few changes to the models, hyperparameters can often be transferred across both neural net depth and width. We focus on learning rate transfer, but also show results on a few other hyperparameters. A few guidelines are introduced, encouraging the use of hyperparameter transfer in future HEP ML models.

Authors: DEZOORT, Gage (Princeton University (US)); VAGE, Liv Helen (Princeton University (US))

Presenter: VAGE, Liv Helen (Princeton University (US))

Session Classification: Contributed talks