

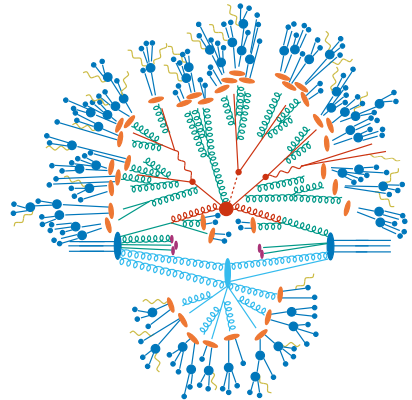
Too good to go: Upcycling Phase-Space Points for Multijet Processes

arXiv:2608.27104

Konrad Helms
with Timo Janßen and Steffen Schumann
15th September 2026

Introduction

- LHC $\xrightarrow{\text{upgrade}}$ HL-LHC finished by ~ 2030
- current MC samples statistically limited compared to HL-LHC data
- $\mathcal{M}_{ab \rightarrow X_N}$ expensive:
 - in particular high-multiplicity final states
 - need stacks: $X + 4j, X + 5j, \dots$
 - dedicated samplers for each process' phase space
- in this talk: reduce $d\Phi_N$ -sampler training costs by reducing No. of $\mathcal{M}_{ab \rightarrow X_N}$ evaluations, while still being unbiased



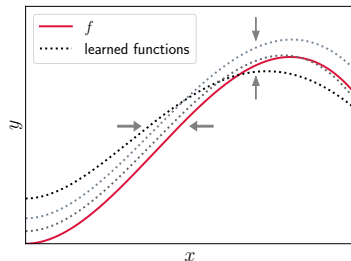
$$\hat{\sigma}_{ab \rightarrow X_N}(\hat{s}) = \frac{1}{2\hat{s}} \int d\Phi_N |\mathcal{M}_{ab \rightarrow X_N}|^2$$

Neural Importance Sampling

- importance sampling:

$$\begin{aligned}
 \int_{[0,1]^d} dx f(x) &= \int_{[0,1]^d} dx \frac{f(x)}{g(x)} g(x) \\
 &\approx \frac{1}{N} \sum_{i=1}^N \underbrace{\frac{f(x_i)}{g(x_i)}}_{\equiv w_i} \text{ with } x_i \sim g
 \end{aligned}$$

- g approximates f well:
 $\iff f/g \simeq \text{const.} \iff \text{spread of weights minimal}$
- neural* importance sampling: [Müller et al., arXiv:1808.03856]
 $g(\cdot) = g(\cdot, \theta)$ with learnable parameters θ
- for us: g is Normalizing Flow



- in HEP, see e.g.:

arXiv:2001.05478,
 arXiv:2604.03511,
 arXiv:2401.09069

Continuous Normalizing Flows & Flow Matching

[Lipman et al. arXiv:2210.02747]

- $$\left. \begin{array}{l} \mathbf{x}_0 \sim \mathcal{N} \\ \mathbf{x}_1 \sim \text{data} \end{array} \right\} \text{linear interpolation: } \mathbf{x}(t) = t\mathbf{x}_1 + (1-t)\mathbf{x}_0 \text{ with } t \in [0, 1]$$
- $\frac{d}{dt}\mathbf{x}(t) = \mathbf{x}_1 - \mathbf{x}_0$
- model target during training: $\mathbf{v}_\theta(\mathbf{x}, t) = \mathbf{x}_1 - \mathbf{x}_0$

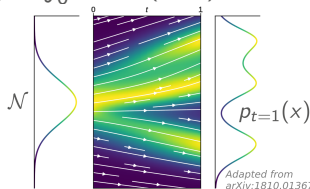
sampling - solving an ODE:

$$\Rightarrow \int_{t=0}^{t=1} d\tau \frac{d}{d\tau} \mathbf{x}(\tau) = \int_0^1 d\tau \mathbf{v}_\theta(\mathbf{x}, \tau) \Rightarrow \mathbf{x}(t=1) = \mathbf{x}(t=0) + \int_0^1 d\tau \mathbf{v}_\theta(\mathbf{x}, \tau)$$

evaluating densities - solving an ODE:

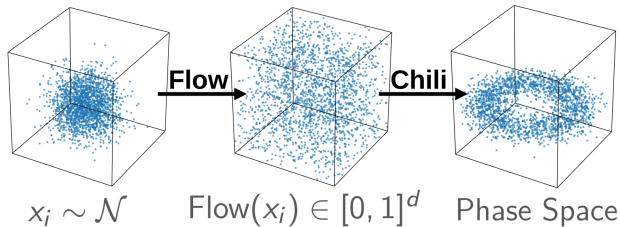
- $$\frac{d \log(p_t(\mathbf{x}))}{dt} = -\nabla \cdot \mathbf{v}_\theta(\mathbf{x}, t)$$

$$\Rightarrow \log(p_{t=1}(\mathbf{x}_1)) = \log(\mathcal{N}(\mathbf{x}_0)) - \int_0^1 d\tau \nabla \cdot \mathbf{v}_\theta(\mathbf{x}, \tau)$$



Workflow: A Chained Mapping

- expert knowledge about integrand, implemented through phase-space mapping
- dimensionality: $d = 3N - 4 + 2$
- Flow: $[0, 1]^d \rightarrow [0, 1]^d$
- CHILI [Bothmann et al., arXiv:2302.10449]: $[0, 1]^d \rightarrow dx_a dx_b d\Phi_N$



CHILI Phase-Space Parameterization

- want to integrate
 $\sigma_{pp \rightarrow X_N} = \sum_{a,b} \int_0^1 dx_a \int_0^1 dx_b \times \text{PDFs} \times \hat{\sigma}_{ab \rightarrow X_N}(x_a, x_b, \mu_R, \mu_F)$ with
 $\hat{\sigma}_{ab \rightarrow X_N}(\hat{s}) = \frac{1}{2\hat{s}} \int d\Phi_N |\mathcal{M}_{ab \rightarrow X_N}|^2$
- with Lorentz-invariant phase-space element $d\Phi_N$, parameterized in $dp_{\perp}^2, dy, d\varphi$
- 4-momentum conservation implemented through combination with light-cone momentum fractions dx_a and dx_b
- coordinate nesting, not literal inclusion:

$$dx_a dx_b d\Phi_{N+1} = [dx_a dx_b d\Phi_N] \times d\Phi_1$$

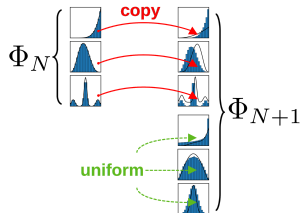
with

$$z_{N+1} = (z_N, p_{\perp,1}^2, y_1, \varphi_1)$$

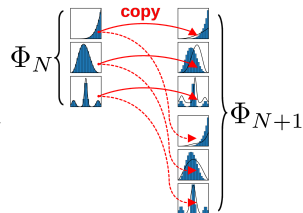
Phase-Space Augmentation: Dimensions

- assuming efficient $d\Phi_N$ sampler exists:

$$\mathbf{z}_N = \begin{matrix} \Phi_N \\ \left[\begin{array}{c} z_N^{(1)} \\ z_N^{(2)} \\ \vdots \\ z_N^{(d_N)} \end{array} \right] \end{matrix} \xrightarrow{\text{copy}} \begin{matrix} \Phi_{N+1} \\ \left[\begin{array}{c} z_N^{(1)} \\ z_N^{(2)} \\ \vdots \\ z_N^{(d_N)} \\ * \\ \vdots \\ * \end{array} \right] \end{matrix}$$



- uniform sampling of extra dimensions



- fill extra dimensions with previous particle's phase space points

Phase-Space Augmentation: Conditions

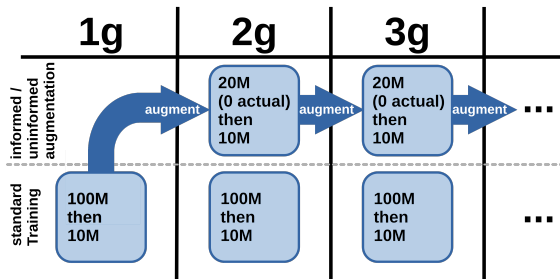
- distributions conditional on the helicity configurations of the external particles
- extra gluon in $X + n \cdot g \rightarrow X + (n + 1)g$ has helicity state $+1$ or -1
 $\implies$ number of helicity combinations on $d\Phi_{N+1}$ doubles w.r.t. $d\Phi_N$
- condition augmentation:

$$[z_N, c] \longrightarrow \begin{cases} [(z_N, p_{\perp, n+1}^2, y_{n+1}, \varphi_{n+1}), c] \\ [(z_N, p_{\perp, n+1}^2, y_{n+1}, \varphi_{n+1}), c + \text{No. helicity cond. on } d\Phi_N] \end{cases}$$

- e.g.: from 32 conditions on $d\Phi_N$ we create 32+32 conditions on $d\Phi_{N+1}$, etc.

Build-up of the Training Stacks

- for multijet processes: can build a stack of consecutive augmented trainings
- training speed-ups accumulate over time

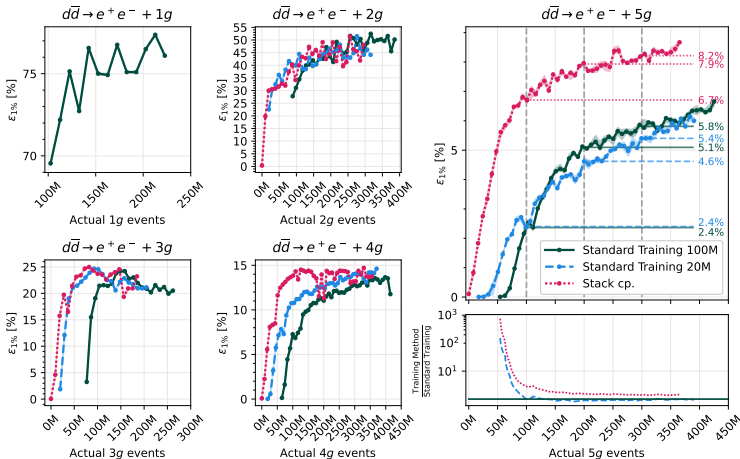


1. for $X + 1g$: sample $\sim 100M$ uniform points $\rightarrow$ evaluate $\mathcal{M}$ to calculate MC weight w
2. (iteratively) train $X + 1g$ phase-space sampler with phase-space points and weights
 $\rightarrow$ iterative training: train model $\rightarrow$ generate phase-space points $\rightarrow$ evaluate $\mathcal{M} \rightarrow$ repeat
3. augment $X + 1g \rightarrow X + 2g$
4. repeat

Results #1

$$d\bar{d} \rightarrow e^+e^- + n \cdot g$$

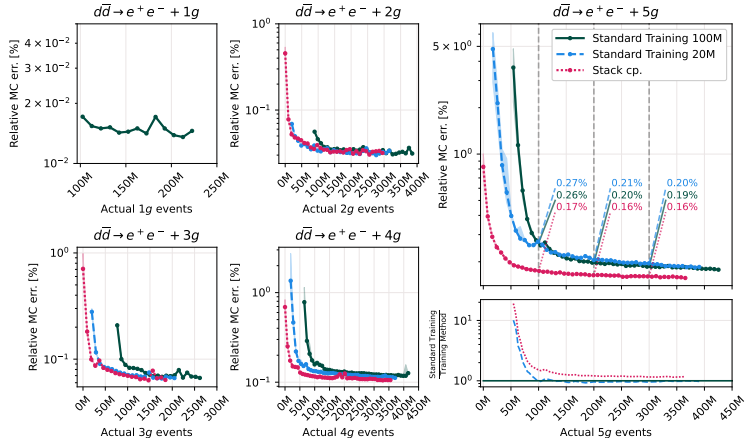
- augmentation vs. standard training
- performance metric: unweighting efficiency $\epsilon = \langle w \rangle / \max(w)$ (larger is better)



Results #2

$$d\bar{d} \rightarrow e^+e^- + n \cdot g$$

- augmentation vs. standard training
- performance metric: relative MC integration error $\sigma_{\text{error}}/\sigma$ (smaller is better)



Thank you for your attention

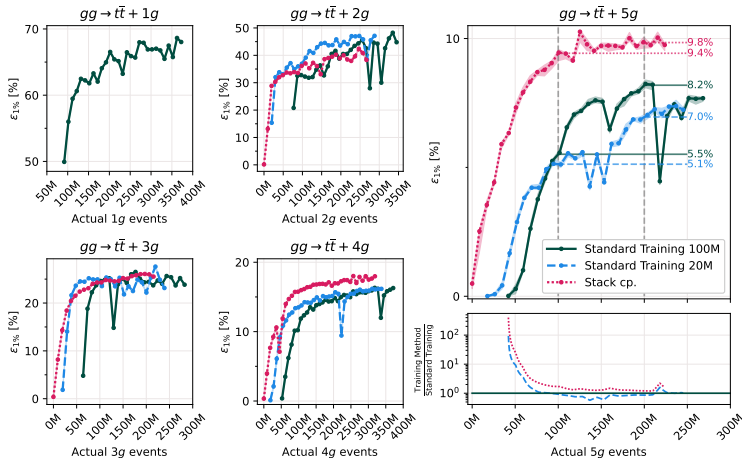
Georg-August-Universität Göttingen

E-Mail: konrad.helms@theorie.physik.uni-goettingen.de

Results #3

$$gg \rightarrow t\bar{t} + n \cdot g$$

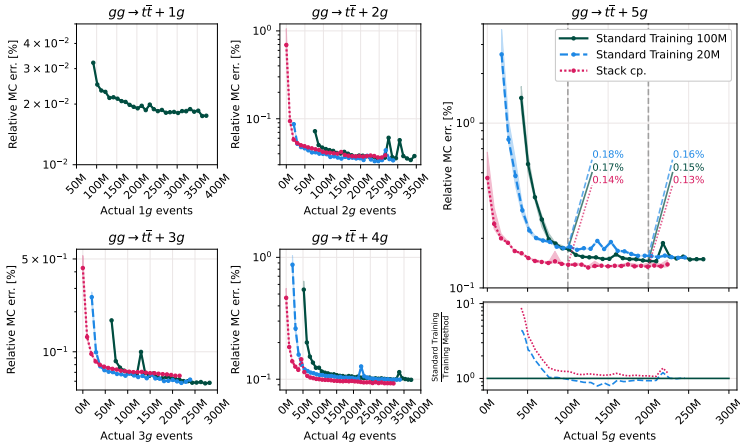
- augmentation vs. standard training
- performance metric: unweighting efficiency
 $\epsilon = \langle w \rangle / \max(w)$
 (larger is better)



Results #4

$$gg \rightarrow t\bar{t} + n \cdot g$$

- augmentation vs. standard training
- performance metric: relative MC integration error
 $\sigma_{\text{error}}/\sigma$
(smaller is better)



Results #5

$$d\bar{d} \rightarrow e^+e^- + n \cdot g \text{ and } gg \rightarrow t\bar{t} + n \cdot g$$

0.1% unweighting efficiency and Kish effective sample size: $\frac{N_{\text{eff.}}}{N_{\text{events}}} = \frac{\langle w \rangle^2}{\langle w^2 \rangle}$
(both: **larger is better**)

