

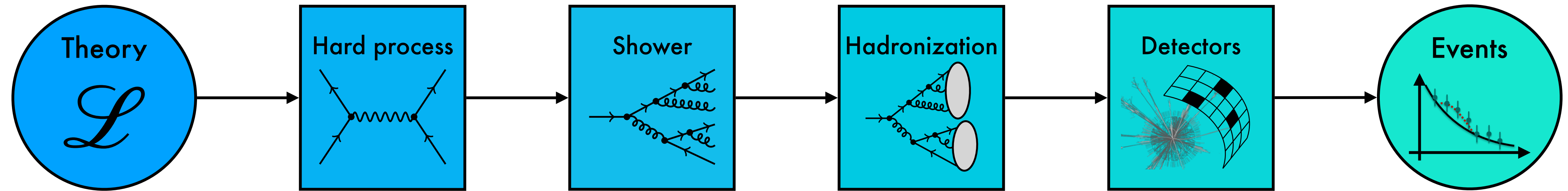
Machine Learning for Event Generation

Making our generators ready for the era of the High-Luminosity LHC

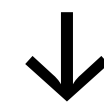
Theo Heimel
CP3, UCLouvain
September 2026



Introduction

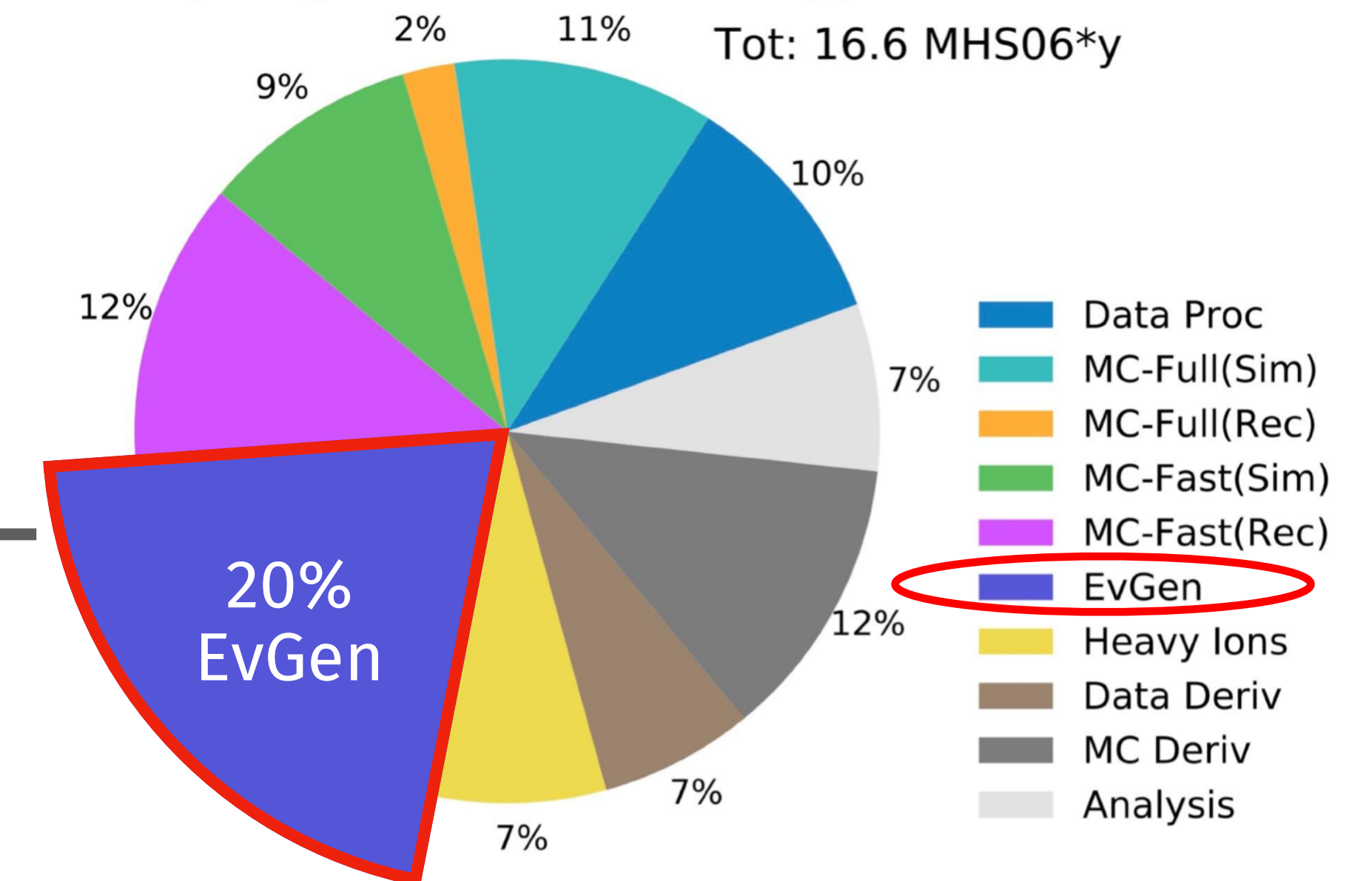


Event generation is a major part of the computational cost



How can we prevent it from becoming a bottleneck at the High-Luminosity LHC?

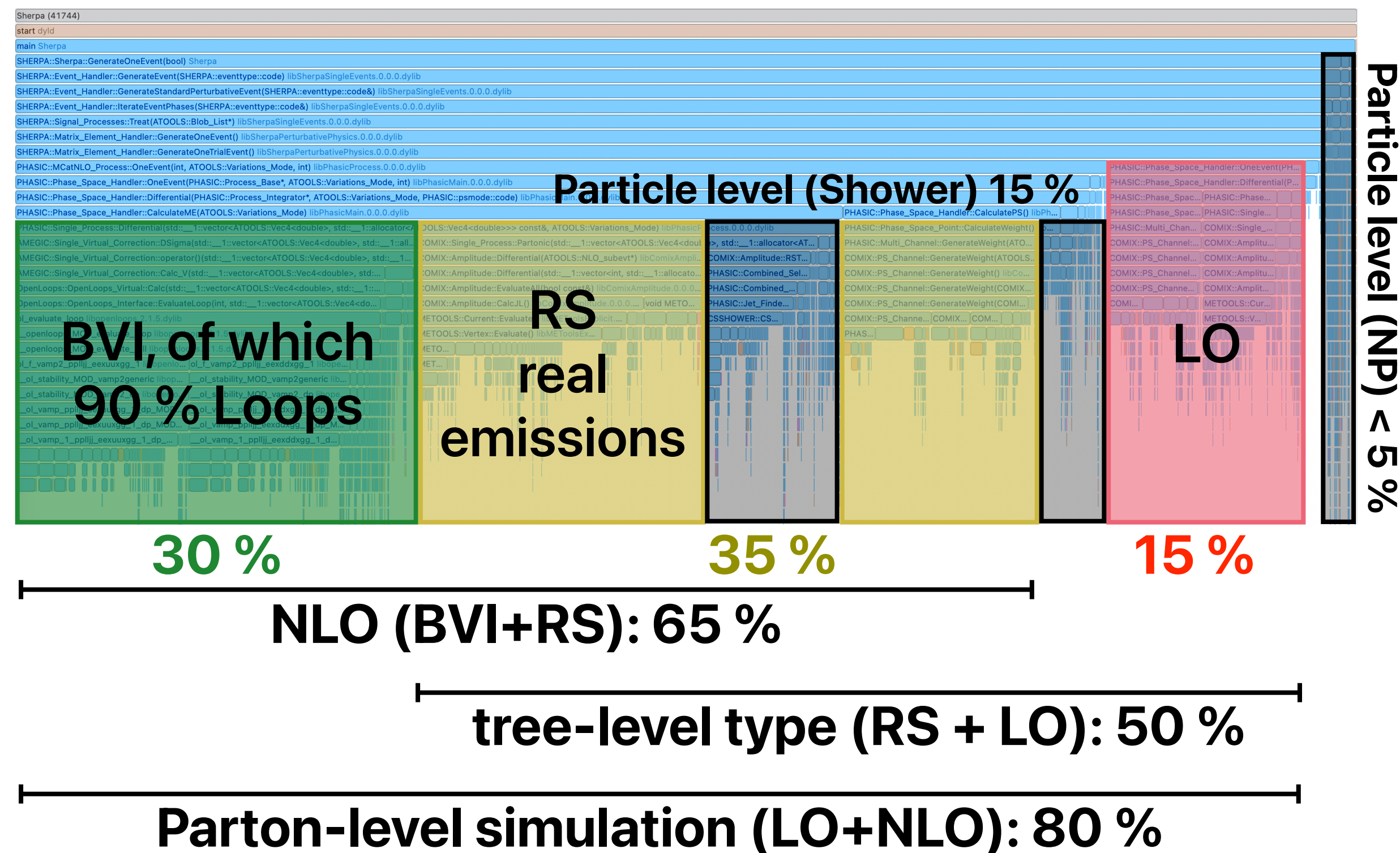
ATLAS Preliminary
2022 Computing Model - CPU: 2031, Aggressive R&D
Tot: 16.6 MHS06*y



[CERN-LHCC-2022-005]

Computational cost

Flame graph of Sherpa in ATLAS Drell-Yan setup
(0-2j @ NLO, >2j @ LO)



- Most expensive: parton-level
→ focus of this talk
- Standard model backgrounds
→ large rates
→ large multiplicities
→ NLO precision necessary
- Signals
→ lower rates
→ LO often sufficient
→ many different processes

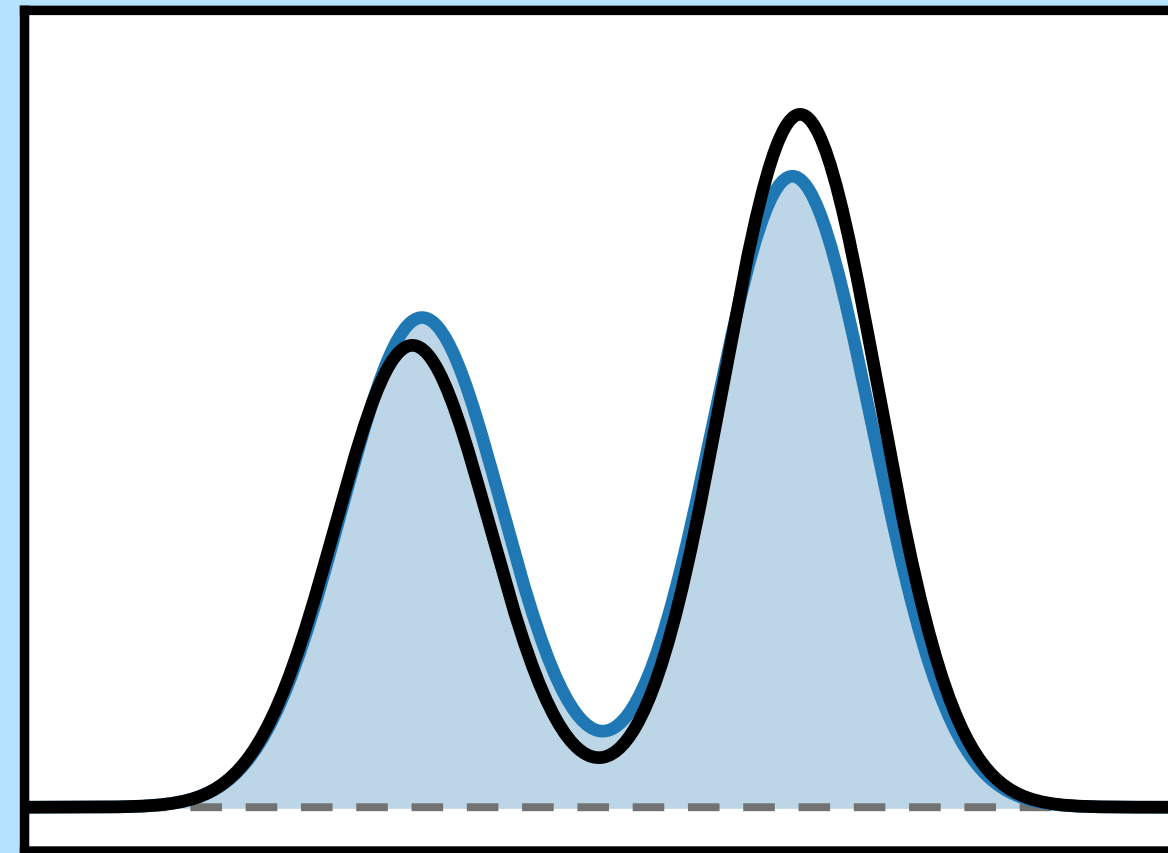
[Figure by Enrico Bothmann]

Leading-order event generation with ML

$$\sigma_{\text{LO}} = \int d\Phi_n B$$

Neural importance sampling

- Use **normalizing flows or CFMs** → density available
- Unweighting step
→ **unbiased**
- Better model
= faster sampling



Tree-level amplitudes

- **Hardware acceleration**
→ move to GPU
→ CPU vector instructions
- Optional:
use **ML surrogates**

Event generators

Two projects for production-ready ML- and hardware-accelerated LHC event generation



MadGraph7

Make end-to-end GPU- and ML-accelerated workflows available for all processes supported by MG5



PEPPER

Complementary to Sherpa with focus on most expensive Standard Model background processes

Both focus on LO for now, NLO will follow

Challenges of ML event generation

Training data

- Don't want to use legacy generators
- Need to “bootstrap” our own training data

Sampling speed

- Matrix element evaluations have become very fast
- Rejecting more samples with cheaper model might be faster

Variety of processes

- Large number of different processes and settings
 - signals and backgrounds
 - model parameters
 - cuts
 - ...

Many trade-offs to consider

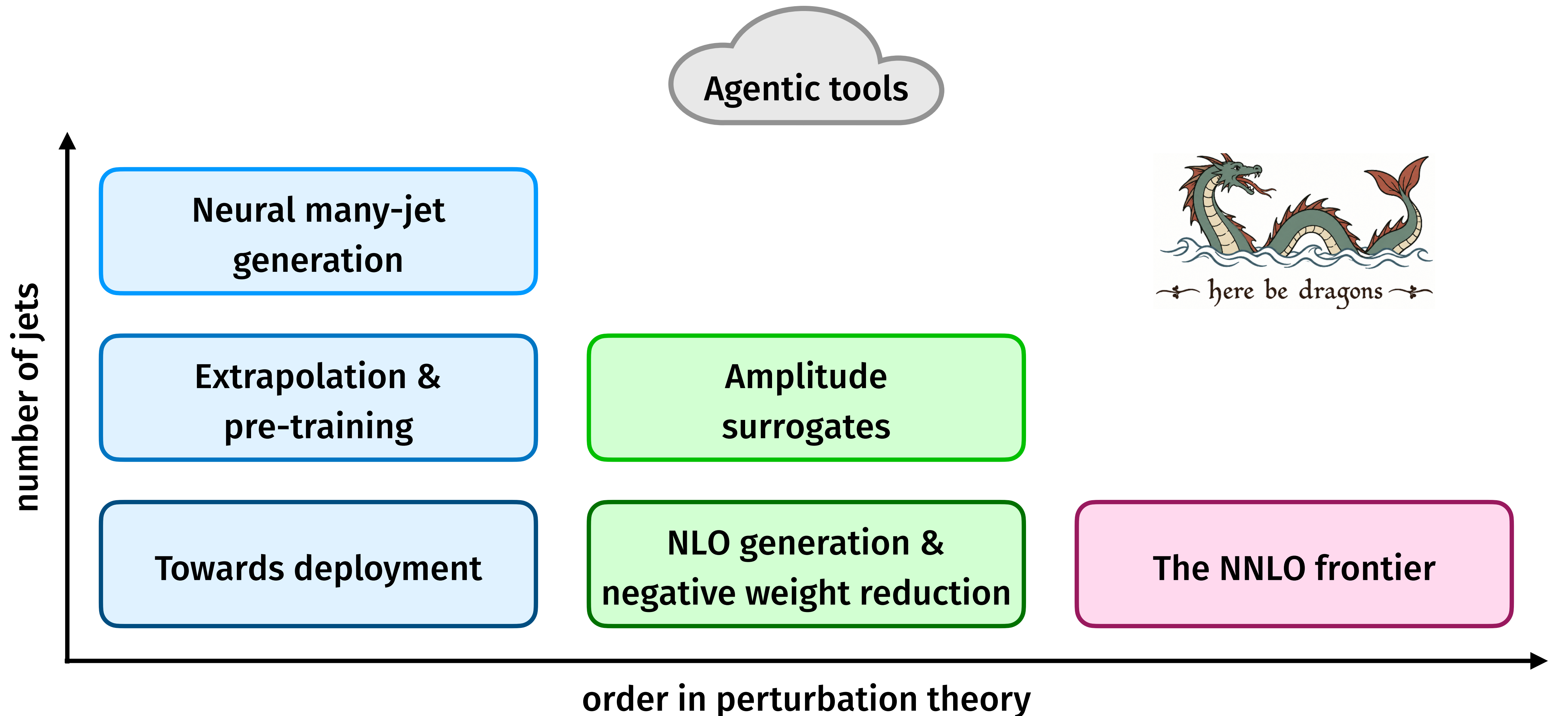
fast training
vs
convergence

expressivity
vs
throughput

surrogates
vs
hardware acceleration

pre-training
vs
training from scratch

Event generation @ ML4Jets 2026



Event generation in MadGraph

Calculate (differential) cross sections

$$d\sigma = \frac{1}{\text{flux}} dx_a dx_b f(x_a) f(x_b) d\Phi_n \langle |M_{\lambda,c,\dots}(p_a, p_b | p_1, \dots, p_n)|^2 \rangle$$

Sum over channels

MadGraph: build channels
from Feynman diagrams

Integrand

MadGraph: $d\sigma/dx$

$$I = \sum_i \left\langle \alpha_i(x) \frac{f(x)}{p_i(x)} \right\rangle_{x \sim p_i(x)}$$

Channel weights

MadGraph: $\alpha_i^{\text{MG}}(x) \sim |M_i|^2$

Channel mappings

MadGraph: use amplitude structure, ...
Analytic mappings + refine with **VEGAS**
(factorized, histogram based
importance sampling)

Event generation in MadNIS

Calculate (differential) cross sections

$$d\sigma = \frac{1}{\text{flux}} dx_a dx_b f(x_a) f(x_b) d\Phi_n \left\langle |M_{\lambda,c,\dots}(p_a, p_b | p_1, \dots, p_n)|^2 \right\rangle$$



Sum over channels

MadGraph: build channels from Feynman diagrams

Integrand

MadGraph: $d\sigma/dx$

$$I = \sum_i \left\langle \alpha_i^\xi(x) \frac{f(x)}{p_i^\omega(x)} \right\rangle_{x \sim p_i^\omega(x)}$$

Learned channel weights

MadGraph: $\alpha_i^{\text{MG}}(x) \sim |M_i|^2$

$$\alpha_i(x) \rightarrow \alpha_i^\xi(x) = \alpha_i^{\text{MG}}(x) \cdot K_i^\xi(x)$$

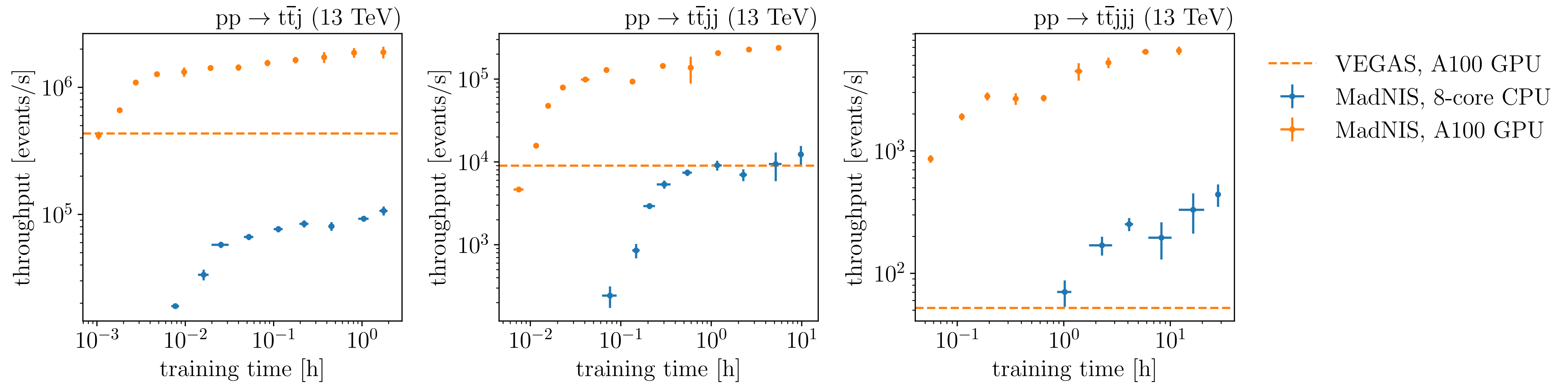
parametrize with NN

Learned channel mappings

MadGraph: use amplitude structure, ...
Analytic mappings + ~~refine with VEGAS~~

refine with
Normalizing Flow

MadNIS performance



- Large speed-up compared to VEGAS baseline already after a few minutes of training, **10x to 100x speed-up** for long training
- Upcoming publication with much larger range of LHC processes

(preliminary)

Upcoming MadGraph7 release

MadMatrix

- Matrix element on GPU
- improved CPU performance from SIMD vectorization

[2507.21039]

AmpliCol

- implement ME using recursion relations to enable higher multiplicities
- made efficient using leading-to full-color reweighting

[2409.12128, 2601.19483]

MadNIS

- neural importance sampling for better unweighting efficiency
- automated defaults: as easy to use as VEGAS
- used automatically for many processes

[2311.01548, 2603.22407]

MadSpace

- new flexible and modular phase space generator
- improved throughput
- GPU- and ML-enabled
- usable beyond MadGraph

[2602.06895]



MadGraph7 alpha release

Alpha release of MadGraph7 out now!

<https://github.com/MadGraphTeam/MadGraph7>



```
cpu_mode 'auto' resolved as 'simd_128'
Start compilation of SubProcesses for device 'simd_128' (1 subprocess(es)
ubprocesses.log
Compilation of SubProcesses done in 5.1 s
Survey
Iteration:      3
Result:         132.7(7.3)
Number of events: 42.0k before cuts, 35.9k after
Run time:       00:00:00.62 wall, 00:00:03.73 cpu

MadNIS training
Batch:          1000 / 1000
Loss:           0.2849
Channels:       3
Run time:       00:00:43.35 wall, 00:03:06.63 cpu

Integration and unweighting
Result:         137.29(20)
Rel. error:     0.1473 %
Rel. stddev:    0.698
Number of events: 225k before cuts, 186k after
Unweighting eff.: 0.16069 before cuts, 0.19422 after
Unweighted events: 36.2k / 100k
Run time:       00:00:04

Individual channels
#   integral ↓   RSD   uweff   N     opt   unweighted
0.5  79.18(16)    0.643  0.23050 78.9k  0     18.2k / 57.6k
0.4  30.105(85)   0.716  0.18156 53.3k  0     9.68k / 22.0k
0.1  28.007(81)   0.738  0.15358 53.9k  0     8.27k / 20.4k
```

Processes + ↻ 🔍

Filter processes

PROCMG7_sm_0

PROCMG7_sm_1

PROCMG7_sm_10

PROCMG7_sm_11

PROCMG7_sm_12

PROCMG7_sm_13

PROCMG7_sm_14

PROCMG7_sm_15

PROCMG7_sm_16

PROCMG7_sm_17

PROCMG7_sm_18

PROCMG7_sm_19

MadGraph7
MadGraph found

MadBoard

PROCESS CARDS MADNIS DIAGRAMS

PROCMG7_sm_10 sm 1 run

$p p > t \bar{t}, (t > w^+ b, w^+ > l^+ \nu_l), (t \sim > w^- b \sim, w^- > l^- \nu_l \sim)$

Run	Cross section (pb)	Unweighted events	Status	Actions
run_01	24.983(25)	100k	done	👁️ ⬇️ 🗑️

Rows per page: 100 1-1 of 1 < >

▶ START RUN

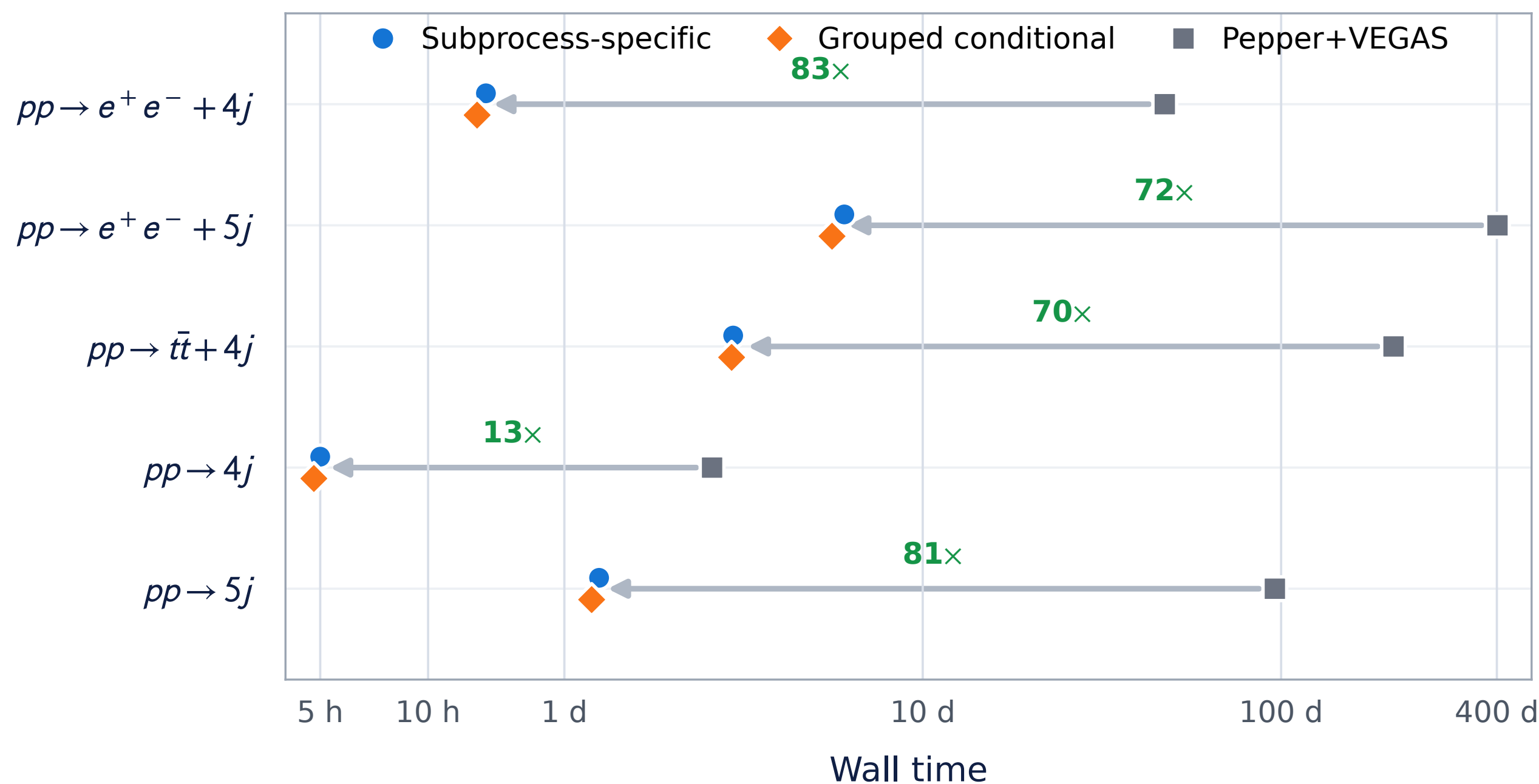
↻ RELOAD

DELETE ALL RUNS

DELETE PROCESS

NIS @ high multiplicity

End-to-end wall time for 1B unweighted events

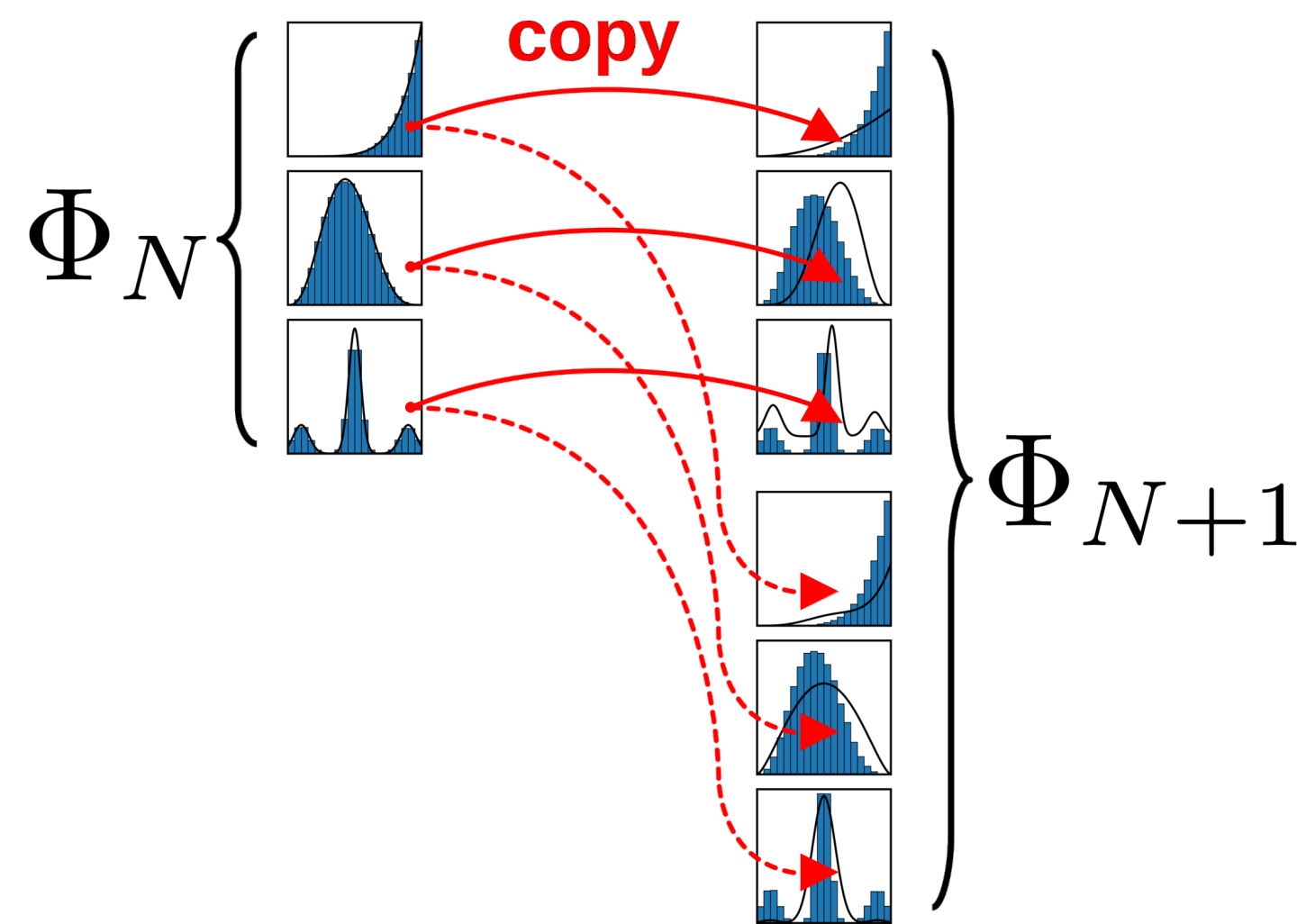


- Neural importance sampling with normalizing flows in PEPPER
- **speed-ups of 10x to 100x**
- End-to-end GPU setup
- Earlier work also with flow matching: high expressivity, but expensive sampling [2506.06203] → for expensive integrands

Avoid training from scratch

Too good to go

[2608.27104] Helms, Janßen, Schumann

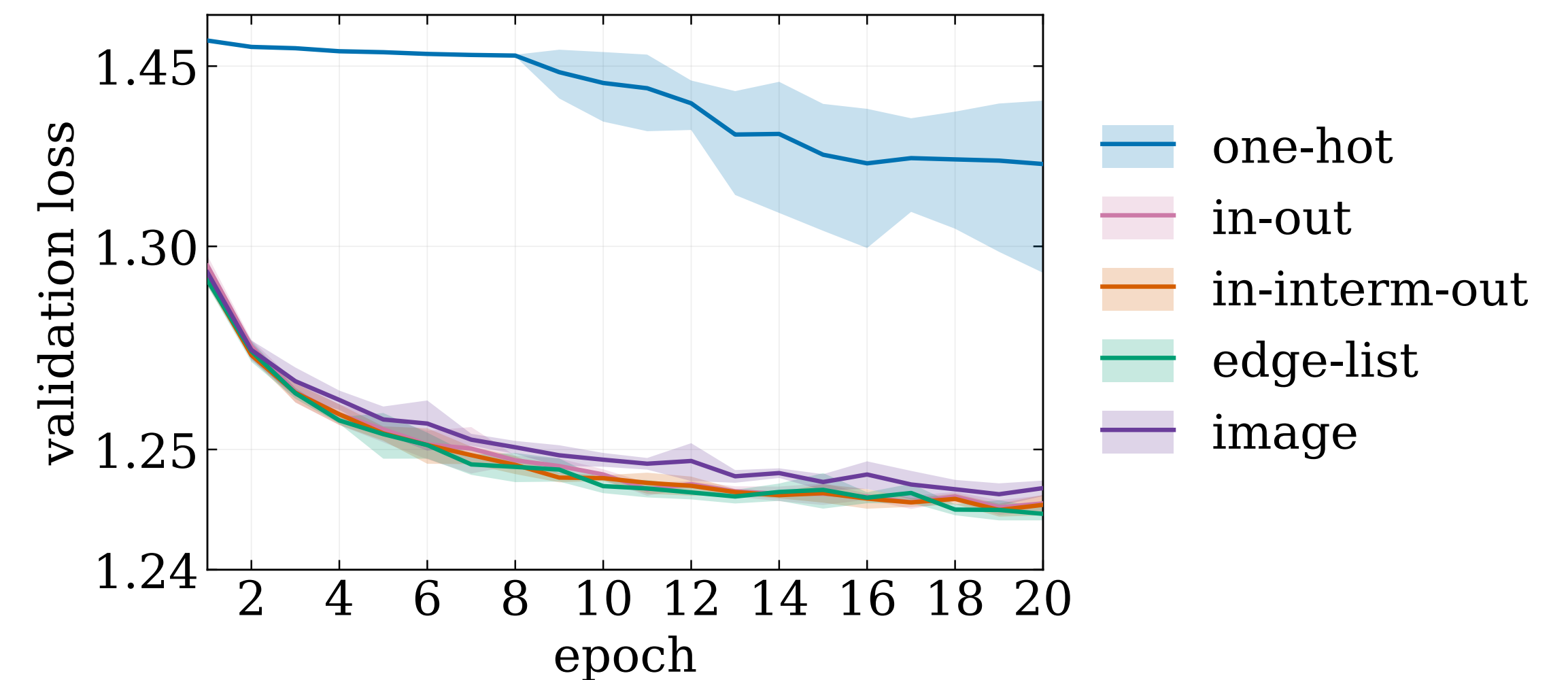


- Cost increases towards more jets
→ reuse results from cheaper trainings
- Initialize (N+1)-particle network based on N-particle network

→ Konrad Helms' talk

LLM conditioning

[2606.23791] Bahl, Plehn, Schiller, Sivagnanalingam



- LLM pre-trained on various SM processes
- Feynman diagrams as input
→ learn physics, not just distribution

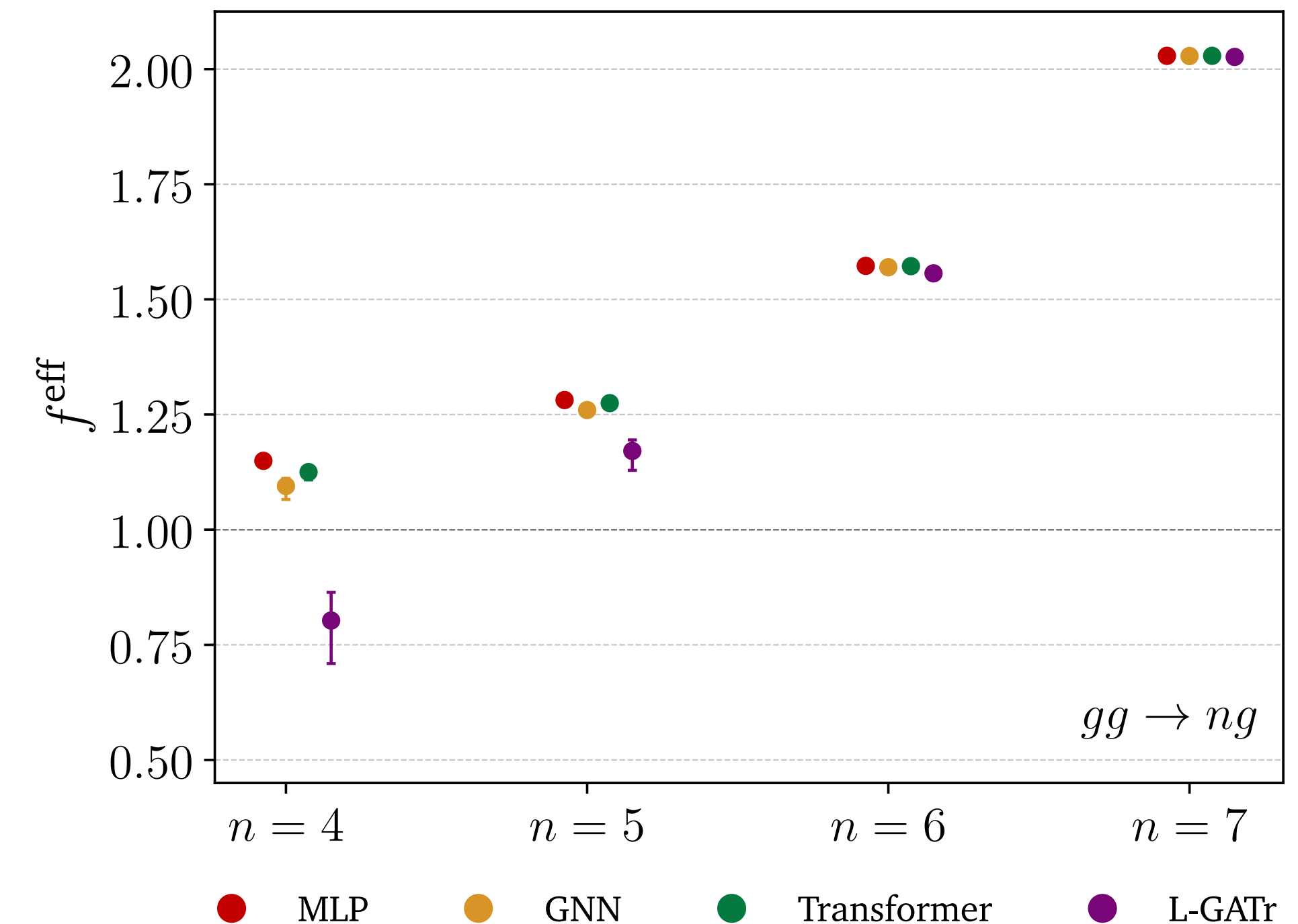
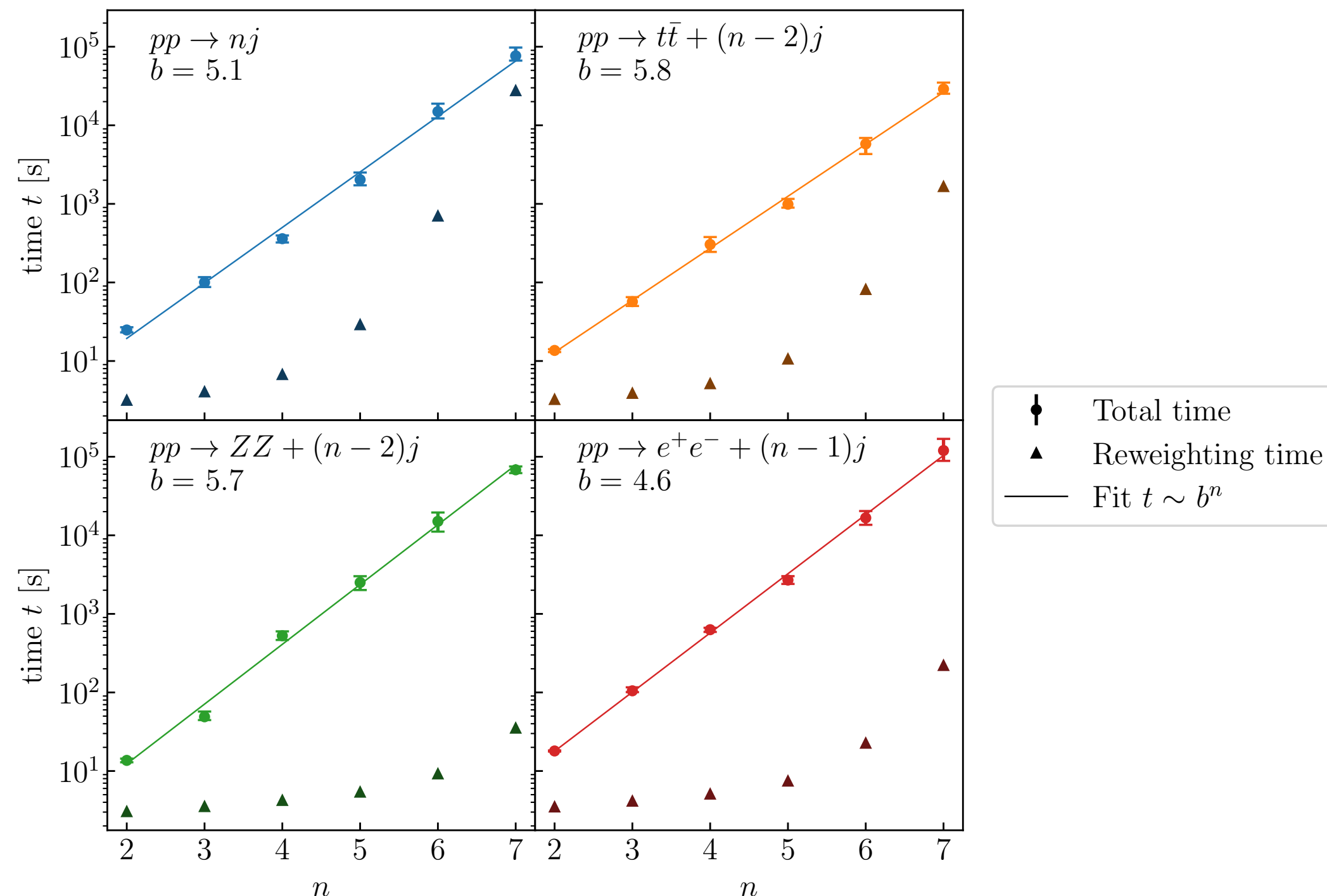
→ Thanush Sivagnanalingam' talk

Surrogates: double-stage unweighting

- Generate unweighted events using surrogate
→ many rejected samples, but fast
- Evaluate truth and unweight again
→ less rejected samples
→ **unbiased result**
- Largest gains for high multiplicities

Process	$T_{\text{merged process}}^{[\text{approach}]}$ [MHS23y]			Ratio
	Baseline $_{\Sigma}$	Baseline	Surrogate	
Z + ≤ 0 jets	0.04	0.04	-	
Z + ≤ 1 jets	0.11	0.11	-	
Z + ≤ 2 jets	0.32	0.34	-	
Z + ≤ 3 jets	0.67	1.23	-	
Z + ≤ 4 jets	1.60	5.44	-	
Z + ≤ 5 jets	16.09	47.83	4.91	3.3
Z + ≤ 6 jets	297.73	130.51	11.78	11.1

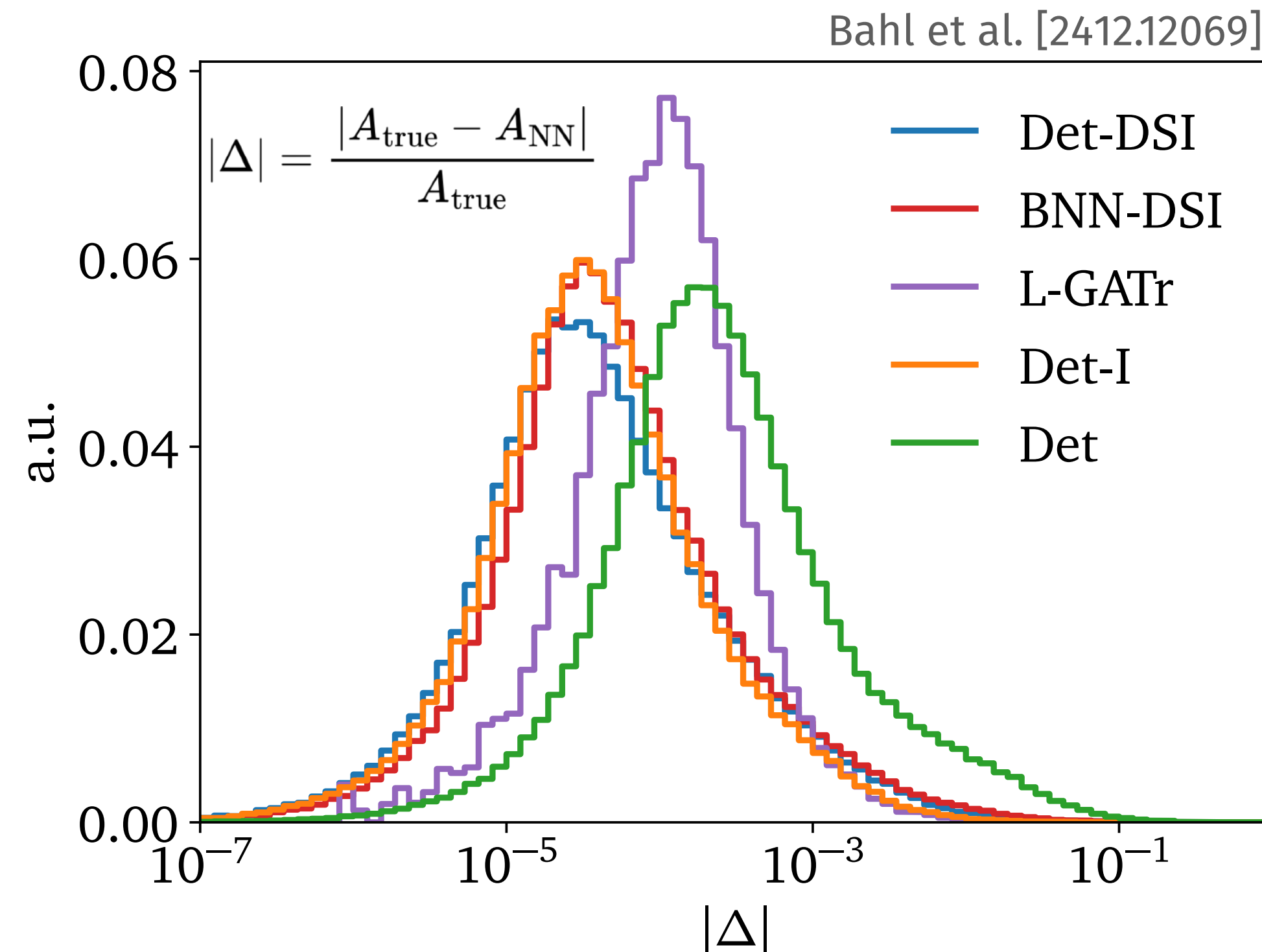
Surrogates: leading-color reweighting



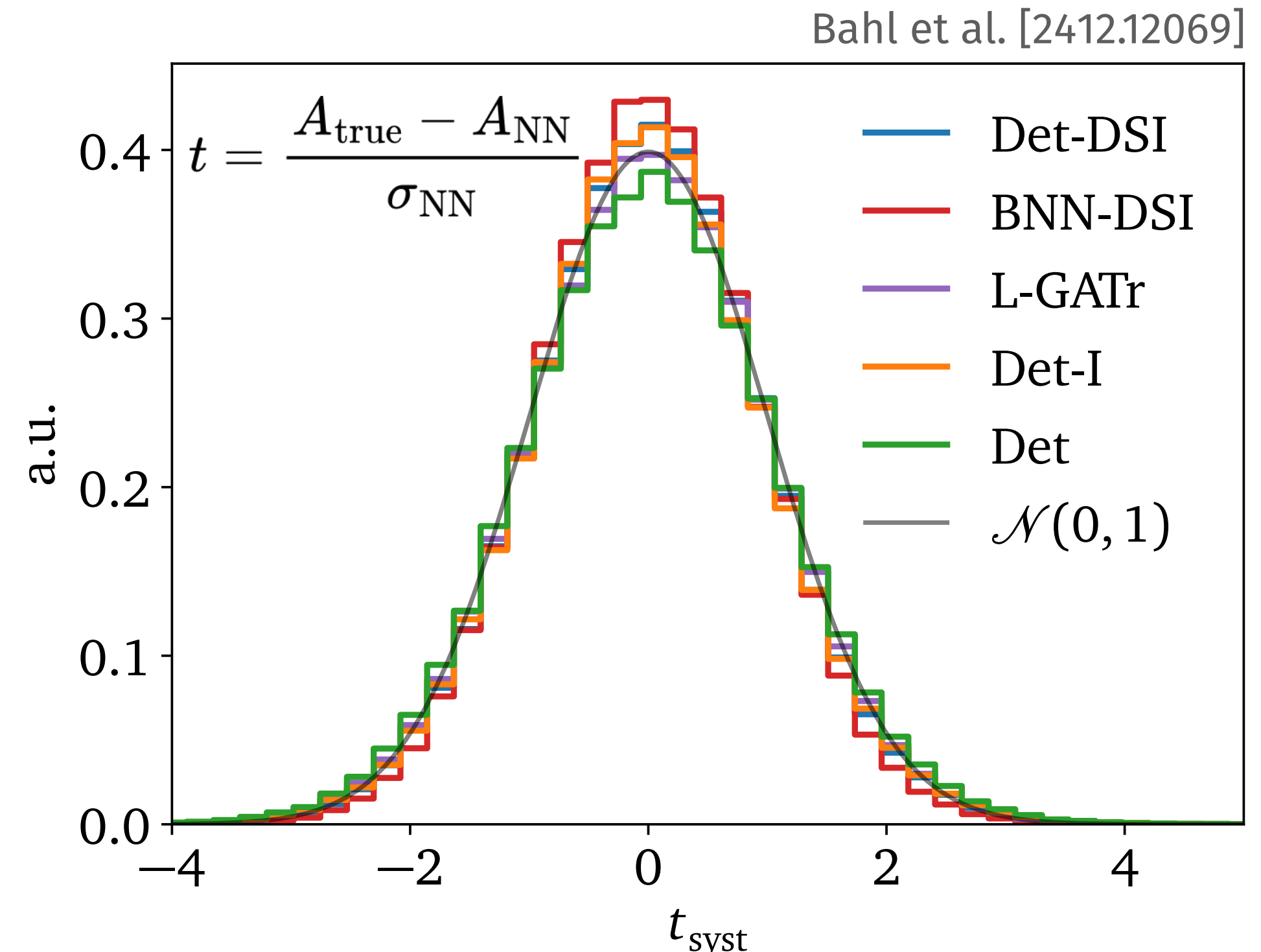
- AmpliCol: leading-color approximation as fast physics-based surrogate [2601.19483]

- ML to speed up reweighting to full color
- Gains of up to **factor 2** for large n

Surrogates: learned uncertainties

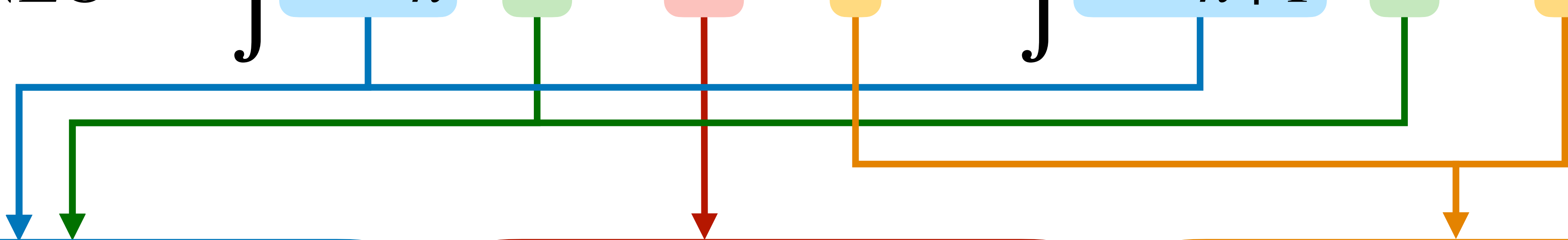


- Symmetry- and physics-aware networks perform better
- Uncertainties small compared to theory uncertainties like scale, PDF, ...



- Well-calibrated uncertainties
- Learned uncertainty also useful to catch low-accuracy events

Next-to-leading order integration

$$\sigma_{\text{NLO}} = \int d\Phi_n (B + V + I) + \int d\Phi_{n+1} (R - S)$$


Importance sampling & tree-level amplitudes

Similar to leading order:

- Neural importance sampling
- Hardware acceleration

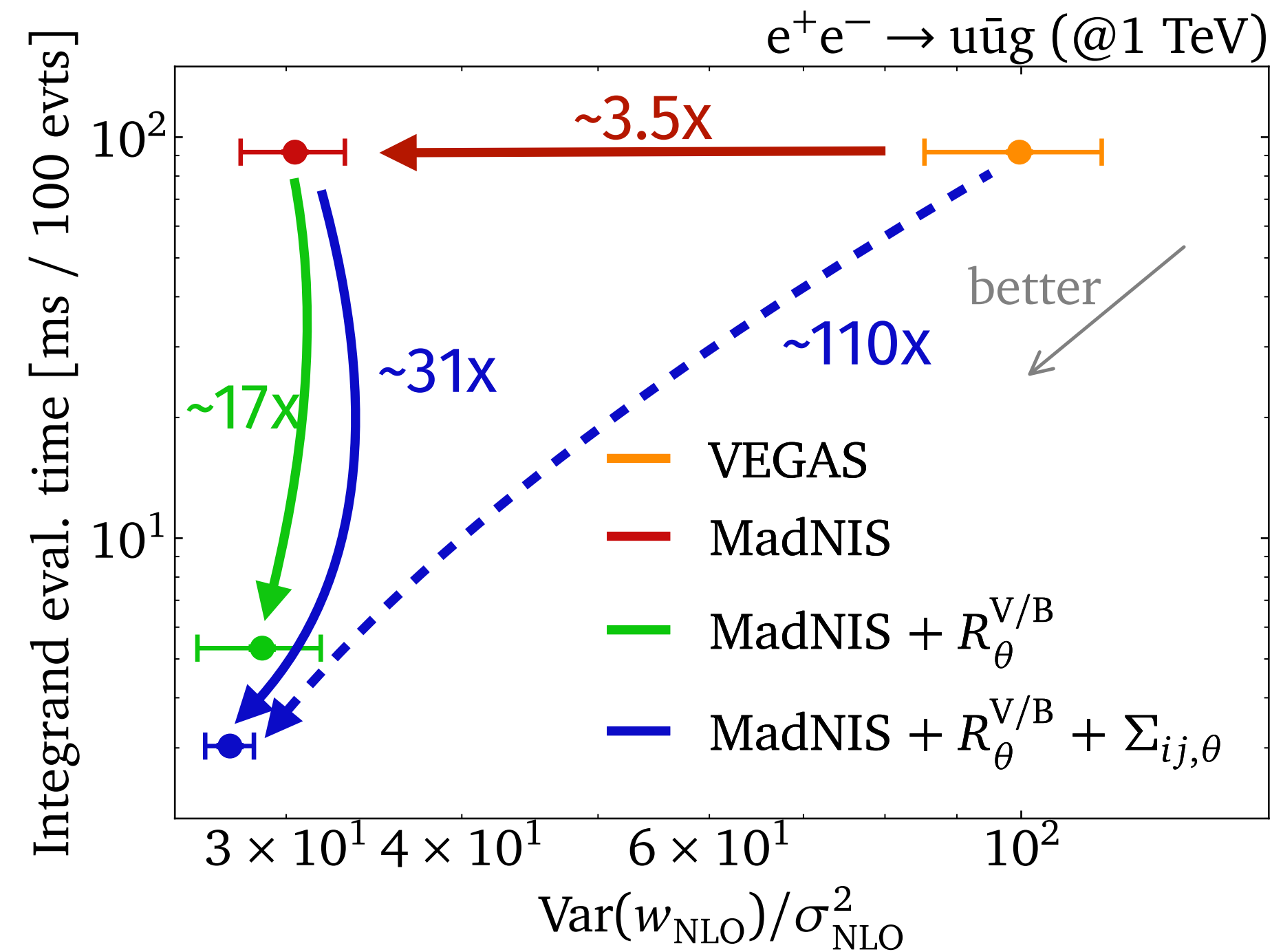
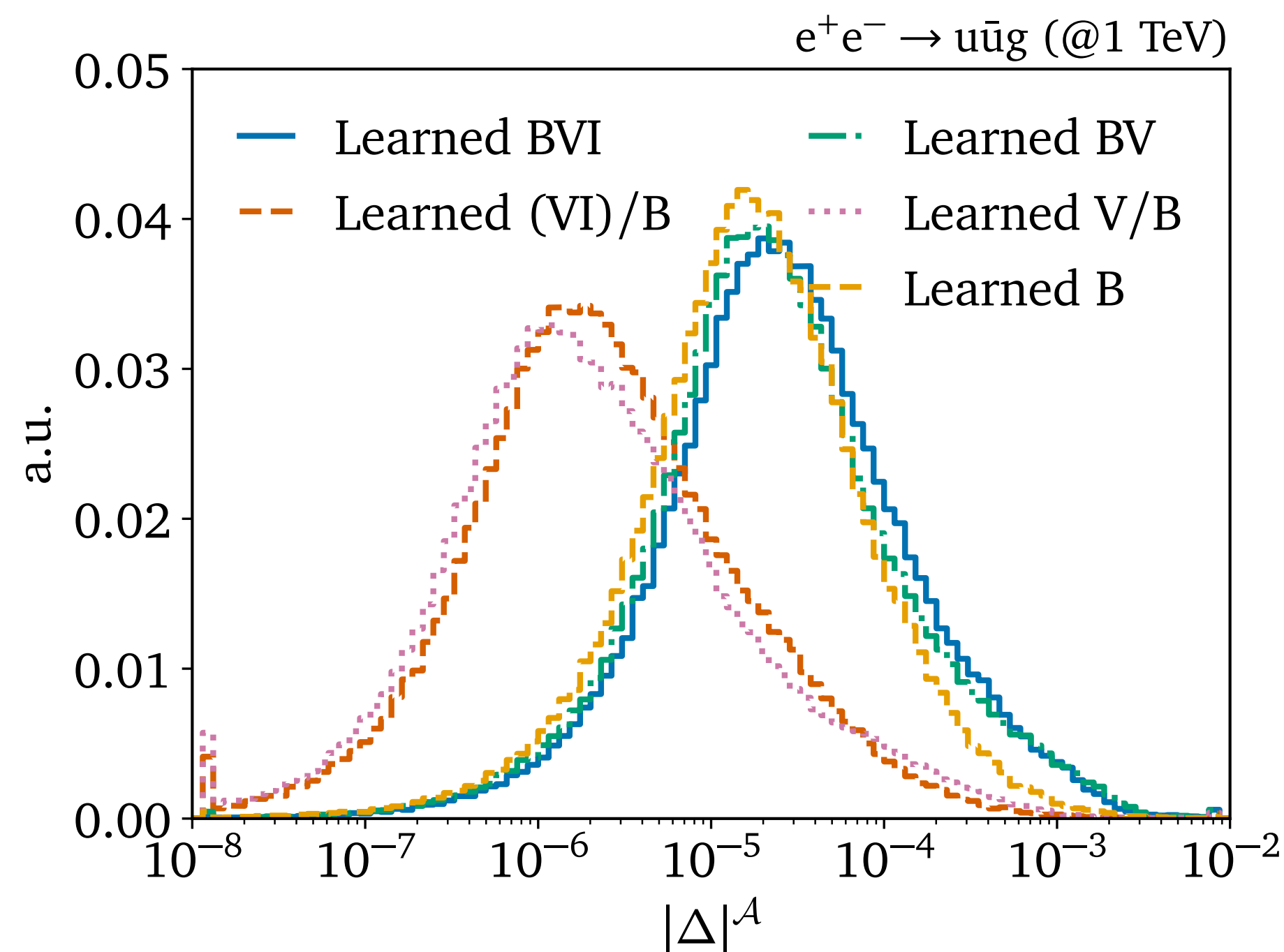
Loop amplitudes

- Most expensive part
- Often need quadruple precision → bad for GPUs
- Speed up with **surrogates**

Subtraction scheme

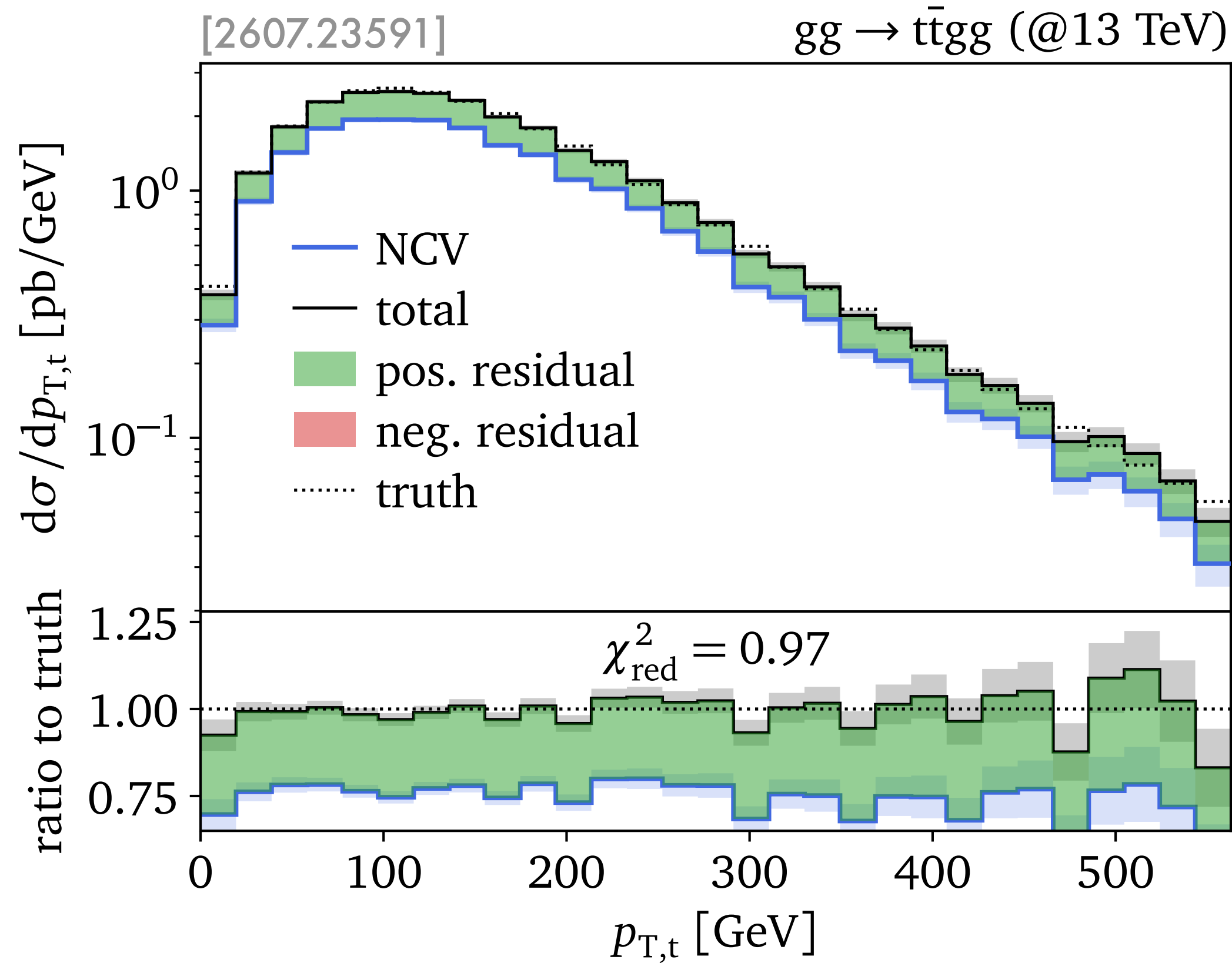
- Causes **negative weights** → large cost downstream
- Some freedom in how to define it → improve with ML

MadNIS @ NLO



- Variance reduction from MadNIS → less samples for given precision
- Replace ME evaluation with surrogate → less time per sample
- Virtual contribution learned accurately, best for ratio virtual / Born

Neural Control Variates @ LO



Network with known integral (NF)
× learned constant

$$I = \langle f(x) - C_\theta p_\theta(x) \rangle_{x \sim \text{unif}} + \langle C_\theta p_\theta(x) \rangle_{x \sim \text{unif}}$$

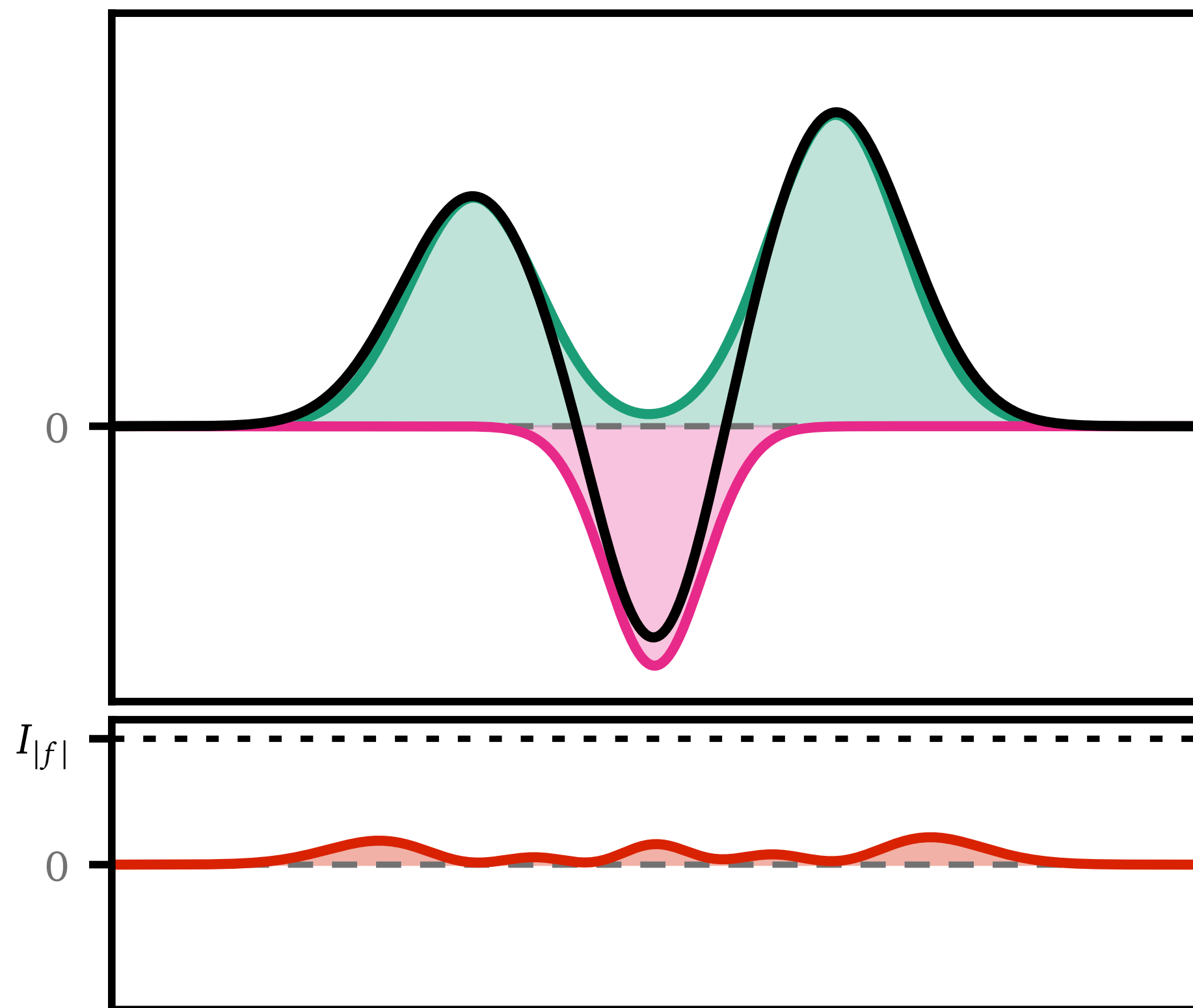
add MadNIS
for res. integral

$$I = \left\langle \frac{f(x) - C_\theta p_\theta(x)}{g_\phi(x)} \right\rangle_{x \sim g_\phi} + \langle C_\theta \rangle_{x \sim p_\theta}$$

**100% unweighting eff.
by construction!**

- Unweighting efficiency: 29% (NIS) to 78% (NIS+NCV)
→ on top of 20× speedup compared to MG5 baseline
- Need penalty loss to avoid negative weights

Neural Control Variates @ NLO



$$I = \langle f(x) - c_\theta(x) \rangle_{x \sim \text{unif}} + \int dx c_\theta(x)$$

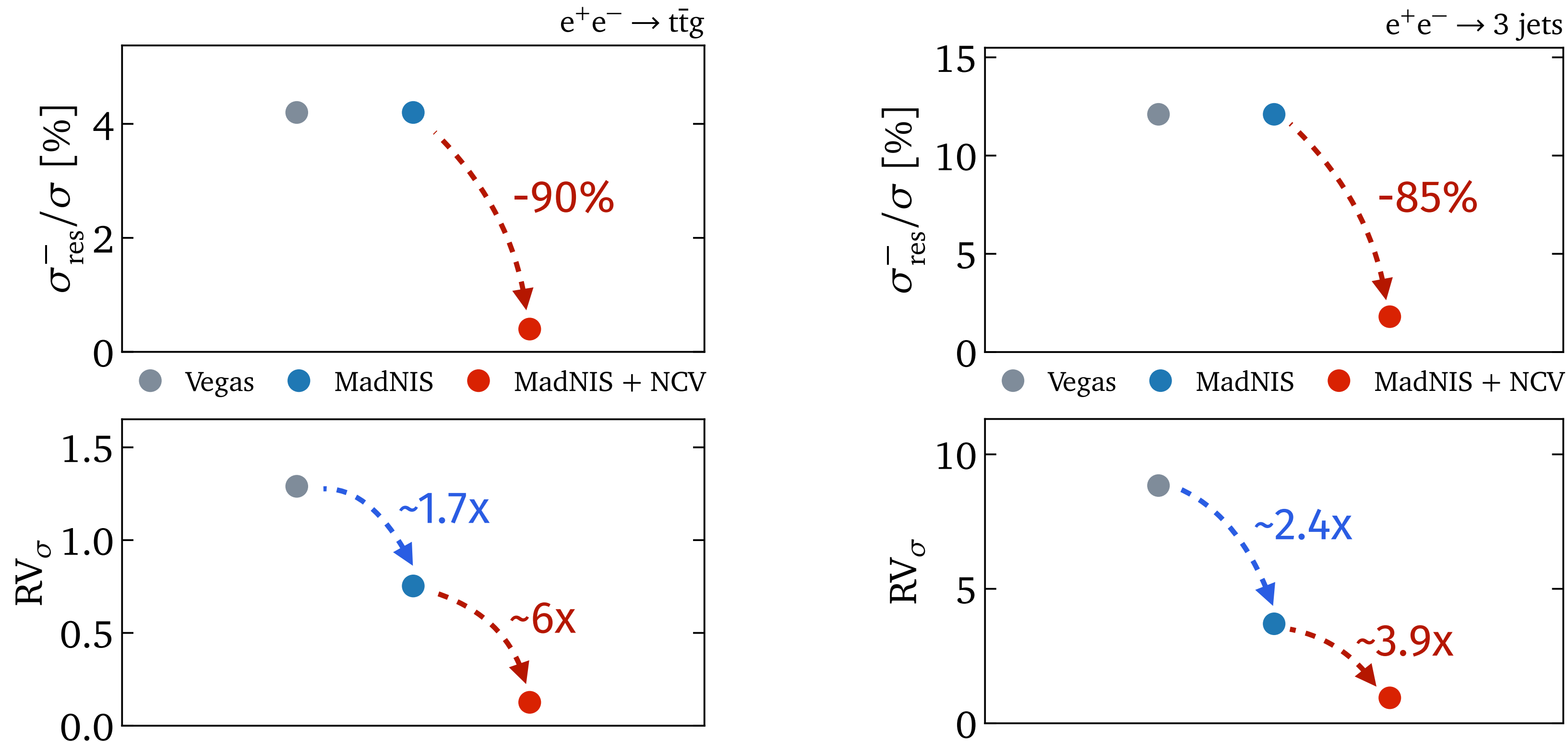
↓ Parametrize *signed*-NCV with
1 positive NCV + 1 negative NCV

$$I = \langle f(x) - C_\theta^+ p_\theta^+(x) + C_\theta^- p_\theta^-(x) \rangle_{x \sim \text{unif}} + \langle C_\theta^+ \rangle_{x \sim p_\theta^+} - \langle C_\theta^- \rangle_{x \sim p_\theta^-}$$

↓ add MadNIS
for res. integral

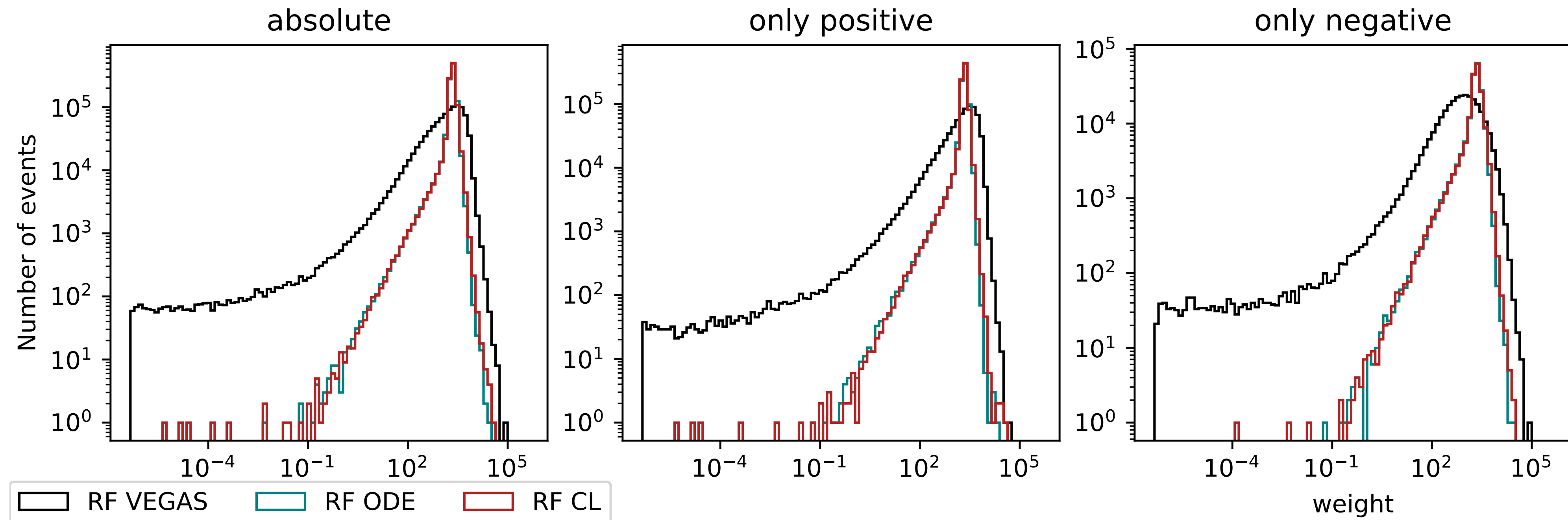
$$I = \left\langle \frac{f(x) - C_\theta^+ p_\theta^+(x) + C_\theta^- p_\theta^-(x)}{g_\phi(x)} \right\rangle_{x \sim g_\phi} + \langle C_\theta^+ \rangle_{x \sim p_\theta^+} - \langle C_\theta^- \rangle_{x \sim p_\theta^-}$$

Neural Control Variates @ NLO



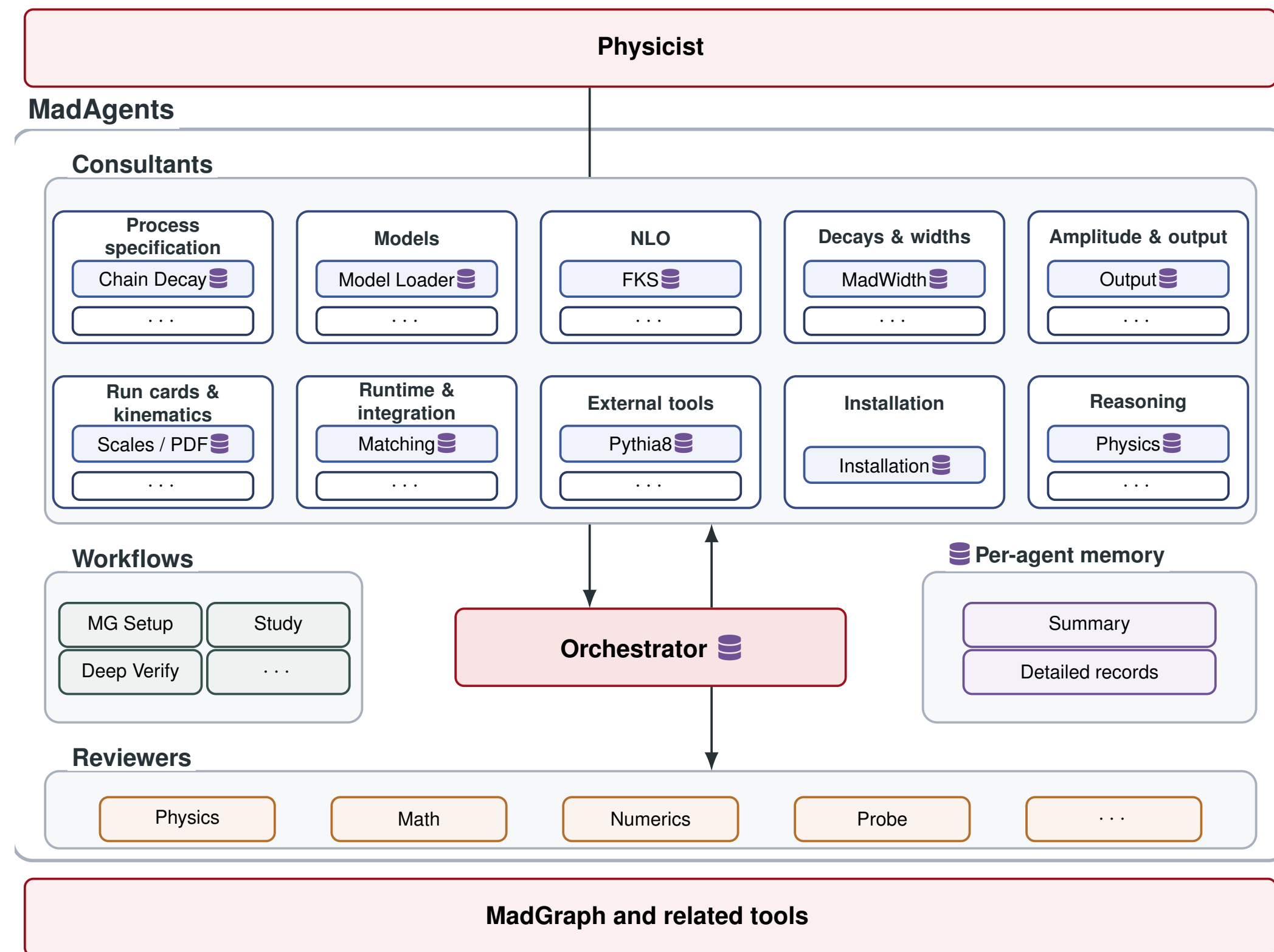
- Better variance reduction than NIS alone
- Loss in effective statistics from sign change: $N_{\text{eff}} = N(1 - 2N_{\text{neg}}/N)^2$
→ save compute in downstream tasks

Neural Importance Sampling @ NNLO



- Neural importance sampling with flow matching
- Tackles negative weights by splitting in positive and negative integrands
→ more efficient than common sampler
- Speed-up by factor 8 compared to STRIPPER baseline

MadAgents: agentic interfaces



- Event generators are complex tools
→ many options unknown even to experts
→ ML+GPU workflows add even more
- Build agentic interfaces specialized for event generation
→ teach beginners
→ help experts
→ reproduce results
→ automate global fits

Outlook

- Neural importance sampling for leading order event-generation is (close to) **production-ready**
→ **speed-ups up to 100x**, on top of hardware acceleration
- Methods developed for LO also **succeed at NLO and NNLO**
- ML solves problems exclusive to higher orders
→ negative weights
→ expensive loop amplitudes
- Agentic interfaces make event generators easier to use



**ML will be an integral part of event generation
at the High-Lumi LHC and beyond!**