

Hey there,

Here is your disruptive ML/AI digest for the day. Here's what you folks might find interesting today.

TOP PICKS

Novel Architectures and Objectives

- [Raven: High-Recall Sequence Modeling with Sparse Memory Routing](#) — This paper introduces a genuinely new memory-writing primitive for linear-time sequence models. State-space models (SSMs) update their entire state densely at every step, causing interference that degrades long-range recall; sliding-window attention (SWA) writes sparsely but hard-evicts tokens outside the window. Raven interpolates between these extremes by maintaining a fixed set of memory slots and, at each step, decaying and updating only a learned, input-dependent subset via routing. This is a conceptual departure from both dense SSM updates and position-bounded SWA, and the authors (including Albert Gu, co-creator of Mamba) show it extrapolates to 16x training length while preserving recall where both baselines collapse. The routing mechanism is the key novelty: it decouples memory write density from sequence position, which is something neither SSMs nor SWA achieve.
- [Generator-Aligned Representation Interfaces for Diagnostic Soft Equivariance](#) — Rather than baking prescribed group actions into specialized operators (the standard approach for exact-equivariant networks), this work exposes transformation generators to a *generic* sequence backbone through aligned canonical and generator-induced views. The authors formalize "soft equivariance" as a probe-specific residual over declared data and transformation distributions, explicitly distinguishing representation consistency from both task robustness and exact equivariance. The Direct Equivariance Error (DEE) metric provides a frozen-checkpoint diagnostic of whether the prescribed representation relation holds under known token/voxel actions. This matters because it makes equivariant structure portable across modalities and backbones without group-specific redesign, and it honestly separates finite-probe evidence from certification over a continuous group.

Theory and Generalization

- [Algorithmic Separation between Constant-Depth and Logarithmic-Depth Neural Networks](#) — This is the first algorithmic (not merely approximation-theoretic) separation between constant-depth and logarithmic-depth networks. The authors identify a class of Boolean functions with hierarchically structured Fourier spectra that log-depth networks can learn efficiently via layerwise coordinate descent, reconstructing the spectra hierarchically and adaptively. Critically, they also prove that every constant-depth, polynomial-width network with regular activations and controlled spectral norms must incur constant L2 approximation error on a subclass of these functions. Jason D. Lee is a co-author. The result is significant because most depth-separation work stops at approximation power; this paper shows that depth matters for *learnability* with a concrete, efficient algorithm on one side and a hardness-of-approximation result on the other.

Quantum ML and Hybrid Paradigms

- [Quantum Speedups for Stochastic Optimization with Heavy-Tailed Noise](#) — This paper proposes a novel quantum mean estimator for multivariate heavy-tailed random variables that beats optimal classical estimators in the low-dimensional regime, and proves matching lower bounds (up to log factors) for the tail index range $p > 4/3$. The authors build an unbiased estimator via a generalized multi-level Monte Carlo technique and use it to construct quantum normalized and projected stochastic gradient descent methods with provably sharper query complexity than classical lower bounds in low-dimensional regimes. Shengyu Zhang is among the authors. What makes this more than a routine

quantum-classical comparison is the dimension-dependent lower bound showing that nontrivial dimension dependence is unavoidable, which sharpens our understanding of where quantum advantage actually lives in stochastic optimization.

Back tomorrow with more!

P.S. In the most recent submission window: cs.LG: 115 papers stat.ML: 16 papers