

Hi all,

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

TOP PICKS

Theory and Generalization

[Hilbert Operator for Progressive Encoding \(HOPE\): A Mathematical Framework for Deconstructing Learned Representations in Deep Networks](#) This paper introduces HOPE, a mathematical framework that shifts network compression into a Hilbert space of continuous functions. By modeling neurons as rank-1 Hilbert-Schmidt operators, it unifies pruning and neuron merging as low-rank subspace projections, resolving scale symmetries and architectural biases in standard compression heuristics. Authored by Peter L. Bartlett (UC Berkeley), this work provides a data-free and hyperparameter-free approach to unbiased architectural decisions across different layer types and sizes. It offers a rigorous geometric lens for analyzing the internal knowledge of trained networks, moving beyond empirical compression tricks.

Probabilistic and Statistical Methods

[Expanding Flow Maps](#) This work introduces Expanding Flow Maps (EFMs), a new class of flow maps that define flows between distributions of increasing dimensionality. Unlike standard flow-based generative models constrained to fixed dimensions, EFMs use an expand operator to augment the state space with conditional noise and a transport map to push the expanded state forward. This allows the output size to become a learned, controllable degree of freedom, extending naturally to discrete simplexes for variable-size graph and sequence generation. It represents a principled departure from fixed-canvas generative paradigms, offering a new inductive bias for settings where dimensionality is dynamic.

Quantum ML and Hybrid Paradigms

[Cautious optimism for deep parameterized quantum circuits](#) This paper tackles the scaling behavior of parameterized quantum circuits (PQCs) by rigorously demonstrating the presence of double descent in quantum machine learning. Using add-one-in perturbation techniques and spectral properties of random matrices, the authors provide analytical proof that gradient-based PQCs can improve performance on unseen data as model size increases. This challenges the traditional view that larger quantum models necessarily degrade generalization, providing a theoretical foundation for cautious optimism in deep quantum learning.

Novel Architectures and Objectives

[Scaling Interpretable Transformers with Parity Bottleneck Layers](#) The ParityTransformer introduces a Deep Parity Bottleneck (DPB) to create language models that are interpretable by construction rather than post-hoc. By replacing a learned over-complete basis with a parameter-free algebraic dictionary, the DPB enforces sparsity with a deterministic incoherence guarantee, bypassing the prohibitive memory costs that have limited per-layer interpretable bottlenecks at scale. This architecture allows subsequent computation to act only on features that survive the sparse bottleneck, ensuring that the model's internal representations are native to its forward pass. It offers a compelling alternative to sparse autoencoders by addressing whether the probed features are actually used during computation.

Back tomorrow with more!

P.S. In the most recent submission window: cs.LG: 106 papers stat.ML: 13 papers