

Fixed Aggregation Features Can Rival GNNs

Graph neural networks have become the standard approach for learning from relational data, with node classification as a central benchmark. Most models follow the message-passing paradigm, alternating neighborhood aggregation with learned linear combinations across multiple hops, achieving strong performance across domains ranging from social networks to biology. Yet, it comes at the cost of high model complexity that poses challenges for interpretation. We ask the question whether this high complexity is really necessary. Recent evidence [1] shows that classic models (GCN, GATv2, and GraphSAGE) remain surprisingly competitive when tuned with standard optimization techniques: learning rate, weight decay, width, depth, normalization, dropout, and residuals. When carefully tuned, they can rival graph transformers and heterophily-aware models.

We take this architectural simplification further: we argue that learning the aggregation itself is theoretically unnecessary, and we empirically verify our hypothesis for node-classification tasks, where neighborhood distributions and feature diversity are often more informative for a task rather than the exact graph structure. Our approach, Fixed Aggregation Features (FAF), creates a tabular dataset comprising multi-hop summaries with fixed, label-agnostic reducers over neighborhoods up to depth K —chosen to match the best validation depth of competing models. Only a classifier, e.g. a MLP, is learned on top. Reducers $\mathcal{R} = \{\text{mean, sum, max, min}\}$ recursively construct and concatenate ($\oplus$) features via $h_v^{(0)} = x_v$ and $h_v^{(k)} = \bigoplus_{r \in \mathcal{R}} r(\{h_u^{(k-1)} : u \in N(v)\})$. MLP is trained as classifier of $z_v = [h_v^{(0)} \oplus \dots \oplus h_v^{(K)}]$ with input dimensionality $|x_v| \cdot (1 + |\mathcal{R}| \cdot K)$ per node v . Note that this mimics the layerwise computation in GNNs. In general, it is not equivalent to taking powers of the adjacency matrix.

We motivate FAF via the Kolmogorov-Arnold (KA) representation theorem [2], which states that any multivariate neighborhood function can be written as a univariate function applied to a fixed aggregation. For $d \geq 2$ and scalar neighbor features $x_1, \dots, x_d \in [0, 1]$ there exists a monotone $\phi : [0, 1] \rightarrow \mathcal{C}$ (the Cantor set) such that any function $f : [0, 1]^d \rightarrow \mathbb{R}$ can be written as a univariate g applied to a fixed aggregation independent of f , i.e., $f(x_1, \dots, x_d) = g\left(3 \sum_{p=1}^d 3^{-p} \phi(x_p)\right)$. g even inherits the continuity properties of f . Accordingly, successive aggregation of learned embeddings is not required, as neighborhood information (f) can be compressed without learning (ϕ) and processed by a feed-forward model (g)—lossless in principle, with the usual caveat of finite precision in practice. The GIN architecture [3] has been derived on the basis of a similar result for (permutation-invariant) multiset functions on countable feature spaces, yet, its layers still have to be learned.

Although theory proposes information-preserving aggregators, good learnability and generalization are not guaranteed. We study feature embeddings that are practically feasible and performative. Empirically, FAF matches or surpasses classic GNNs across 12 out of 14 diverse node-classification benchmarks spanning citation, coauthor, and Amazon graphs, as well as heterophilous datasets. It trails only on Minesweeper and Roman-Empire, datasets whose best performing GNNs require linear residual connections; the remaining performance gap matches the benefit conferred by such residuals reported by [1]. This suggests they may need different aggregations per hop or interactions between consecutive layers; on these datasets, GNNs also benefit from deep message passing (10-15 layers).

Because FAF produces tabular features, standard interpretability tools (e.g., feature importance) apply and reveal which hops and reducers matter, helping to distinguish homophily from heterophily, assessing the utility of individual reducers, and discovering long-range interactions. The tabular view also makes it natural to enrich the feature set with common network science descriptors (degree, centrality, communities), and neighborhood-masking features inspired by graph rewiring, all of which can improve robustness and reduce variance.

In summary, we (i) provide a KA-based explanation that fixed, label-agnostic aggregation suffices in principle, removing the need to learn neighborhood embeddings; (ii) present a pipeline to turn graph learning into tabular learning via recursive fixed reducers and MLPs; (iii) demonstrate competitive performance across standard datasets, showing their limitations as benchmarks to meaningfully advance graph learning; and (iv) leverage the tabular formulation for straightforward interpretability of hop depth, aggregator choice, and the interplay of features, graph structure, and labels, thereby establishing a transparent, strong baseline for future graph learning research.

- [1] Y. Luo et al. Classic GNNs are strong baselines: Reassessing GNNs for node classification. NeurIPS, 2024.
- [2] J. Schmidt-Hieber. The Kolmogorov–Arnold representation theorem revisited. Neural Networks, 2021.
- [3] K. Xu et al. How Powerful are Graph Neural Networks? ICLR, 2019.