
GraphFLEx: Structure Learning Framework for Large Expanding Graphs

Anonymous Author(s)

Affiliation

Address

email

Abstract

Graph structure learning is a core problem in graph-based machine learning, essential for uncovering latent relationships and ensuring model interpretability. However, most existing approaches are ill-suited for large-scale and dynamically evolving graphs, as they often require complete re-learning of the structure upon the arrival of new nodes and incur substantial computational and memory costs. In this work, we propose GraphFLEx—a unified and scalable framework for Graph Structure Learning in Large and Expanding Graphs. GraphFLEx mitigates the scalability bottlenecks by restricting edge formation to structurally relevant subsets of nodes identified through a combination of clustering and coarsening techniques. This dramatically reduces the search space and enables efficient, incremental graph updates. The framework supports 48 flexible configurations by integrating diverse choices of learning paradigms, coarsening strategies, and clustering methods, making it adaptable to a wide range of graph settings and learning objectives. Extensive experiments across 26 diverse datasets and Graph Neural Network architectures demonstrate that GraphFLEx achieves state-of-the-art performance with significantly improved scalability. Our implementation is publicly available [here](#).

1 Introduction

Graph representations capture relationships between entities, vital across diverse fields like biology, finance, sociology, engineering, and operations research [1–4]. While some relationships, such as social connections or sensor networks, are directly observable, many, including gene regulatory networks, scene graph generation [5], brain networks, [6] and drug interactions, require inference [7]. Even when available, graph data often contains noise, requiring denoising and recalibration. In such cases, inferring the correct graph structure becomes more crucial than the specific graph model or downstream algorithm.

Graph Structure Learning (GSL) offers a solution, enabling the construction and refinement of graph topologies. GSL has been widely studied in both supervised and unsupervised contexts [8, 9]. In supervised GSL (s-SGL), the adjacency matrix and Graph Neural Networks (GNNs) are jointly optimized for a downstream task, such as node classification. Notable examples of s-GSL include *NodeFormer* [10], *Pro – GNN* [11], *WSGNN* [12], and *SLAPS* [13]. Unsupervised GSL (u-SGL), on the other hand, focuses solely on learning the underlying graph structure, typically through adjacency or Laplacian matrices. Methods in this category include approximate nearest neighbours ($A - NN$) [14, 15], k -nearest neighbours ($k - NN$) [16, 17], covariance estimation (*emp.Cov.*) [18], graphical lasso (*GLasso*) [19], *SUBLIME* [8], and signal processing techniques like *l2-model*, *log-model* and *large-model* [20, 21].

Supervised structure learning (s-SGL) methods have demonstrated effectiveness in specific tasks; however, their reliance on labeled data and optimization for downstream objectives—particularly node classification—significantly constrains their generalizability to settings where annotations are scarce or unavailable [8]. Unsupervised structure learning (u-SGL) methods, which constitute the focus

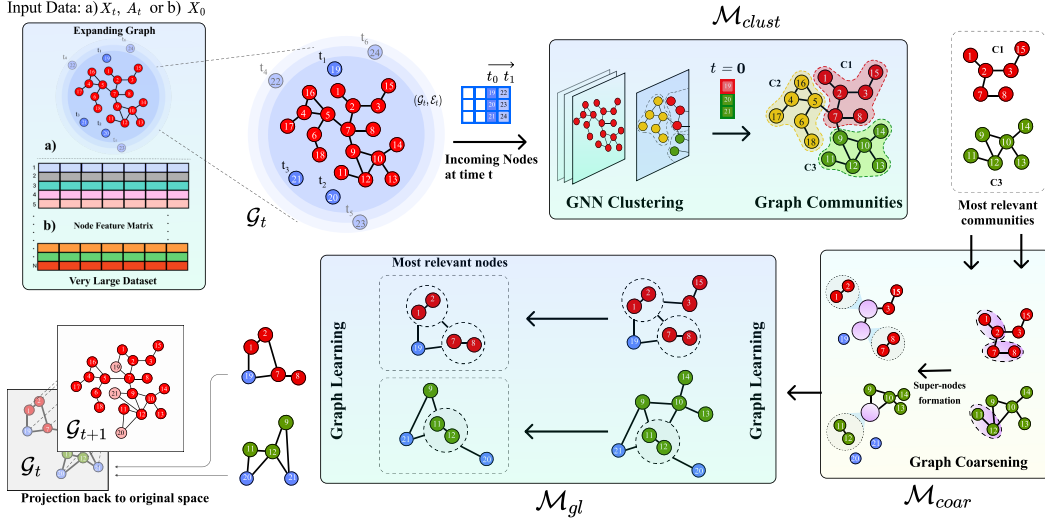


Figure 2: This figure illustrates the general pipeline of GraphFLEX, designed to efficiently handle both a) large datasets with missing structure and b) expanding graphs. Both scenarios can be modeled as expanding graphs (details in Section 3.1). GraphFLEX processes a graph (G_t) and incoming nodes ($\mathcal{E}_{t+1}$) at time t , newly arriving nodes are shown with different timestamps and shades of blue to indicate their arrival time. Our framework comprises of three main components: i) **Clustering**, which infers $\mathcal{E}_{t+1}$ nodes to existing communities using a pre-trained model $\mathcal{M}_{clust}(G_0)$ into smaller, more manageable communities; ii) Since these communities may still be large, a **Coarsening** module is applied to further reduce their size while preserving essential structural information; and iii) Finally, a **Learning** module, where the structure associated with $\mathcal{E}_{t+1}$ nodes are learned using the coarsened graph, followed by projecting this structure onto the G_t graph to create graph G_{t+1} .

of this work, offer broader applicability. Nevertheless, both s-SGL and u-SGL approaches exhibit critical limitations in their ability to scale to large graphs or adapt efficiently to expanding datasets.

To address these challenges, we introduce **GraphFLEX**, a unified and scalable framework for *Graph Structure Learning in Large and Expanding Graphs*. GraphFLEX is built upon the coordinated integration of three foundational paradigms in graph processing: *graph clustering*, *graph coarsening*, and *structure learning*. While each of these methodologies has been studied extensively in isolation, their joint application within a single framework has remained largely unexplored. The novelty of GraphFLEX lies not merely in combining these components, but in the principled manner in which they are algorithmically aligned to reinforce one another—clustering serves to localize the search space, coarsening reduces structural redundancy while preserving global properties, and structure learning operates efficiently within this refined context. This integration enables GraphFLEX to scale effectively to large datasets and accommodate dynamic graphs through *incremental updates*, eliminating the need for expensive re-training. Additionally, the framework supports **48 modular configurations**, enabling broad adaptability across datasets, learning objectives, and deployment constraints. Crucially, we establish *theoretical guarantees* on edge recovery fidelity and computational complexity, offering rigorous foundations for the framework’s efficiency and reliability. As illustrated in Figure 2, GraphFLEX significantly reduces the candidate edge space by operating on structurally relevant node subsets. Empirical evaluations, summarized in Figure 1, demonstrate that GraphFLEX substantially outperforms existing baselines in both runtime and scalability.

Key contributions of this work include:

- GraphFLEX unifies *multiple structure learning strategies* within a single flexible framework.
- GraphFLEX demonstrates effectiveness in *handling growing graphs*.
- GraphFLEX enhances the *scalability* of graph structure learning on large-scale graphs.
- GraphFLEX serves as a *comprehensive framework* applicable individually for clustering, coarsening, and learning tasks.

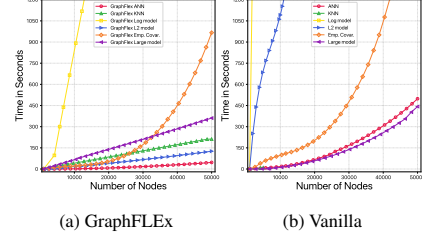


Figure 1: High computational time required to learn graph structures using existing methods, whereas GraphFLEX effectively controls computational growth, achieving near-linear scalability. Notably, Vanilla KNN failed to construct graph structures for more than 10K nodes due to memory limitations.

- We provide both *empirical and theoretical results*, demonstrating the effectiveness of GraphFLEX across a range of datasets.

2 Problem Formulation and Background

A graph $\mathcal{G}$ is represented using $\mathcal{G}(V, A, X)$ where $V = \{v_1, v_2 \dots v_N\}$ is the set of N nodes, each node v_i has a d -dimensional feature vector x_i in $X \in \mathbb{R}^{N \times d}$ and $A \in \mathbb{R}^{N \times N}$ is adjacency matrix representing connection between i^{th} and j^{th} nodes when entry $A_{ij} > 0$. An expanding graph $\mathcal{E}_{\mathcal{G}}$ can be considered a variant of graph $\mathcal{G}$ where nodes v now have an associated timestamp τ_v . We can represent an expanding graph as a sequence of graphs, i.e., $\mathcal{E}_{\mathcal{G}} = \{\mathcal{G}_0, \mathcal{G}_1, \dots, \mathcal{G}_T\}$ where $\{\mathcal{G}_0 \subseteq \mathcal{G}_1 \dots \subseteq \mathcal{G}_T\}$ at $\tau \in \{0, \dots, T\}$ timestamps. New nodes arriving at different timestamps are seamlessly integrating into initial graph $\mathcal{G}_0$.

Problem statement. Given a partially known or missing graph structure, our goal is to incrementally learn the whole graph, i.e., learn adjacency or laplacian matrix. Specifically, we consider two unsupervised GSL tasks:

Goal 1. Large Datasets with Missing Graph Structure: *In this setting, the graph structure is entirely unavailable, and existing methods are computationally infeasible for learning the whole graph in a single step. To address this issue, we first randomly partition the dataset into exclusive subsets. We then learn the initial graph $\mathcal{G}_0(V_0, X_0)$ over a small subset of nodes and incrementally expand it by integrating additional partitions, ultimately reconstructing the full graph $\mathcal{G}_T$.*

Goal 2. Partially Available Graph: *In this case, we only have access to the graph $\mathcal{G}_t$ at timestamp t , with new nodes arriving over time. The goal is to update the graph incrementally to obtain $\mathcal{G}_T$, without re-learning it from scratch at each timestamp.*

GraphFlex addresses these challenges with a unified framework, outlined in Section 3. Before delving into the framework, we review some key concepts.

2.1 Graph Reduction

Graph reduction encompasses sparsification, clustering, coarsening, and condensation [22]. GraphFlex employs clustering and coarsening to refine the set of relevant nodes for potential connections. **Graph Clustering.** Graphs often exhibit global heterogeneity with localized homogeneity, making them well-suited for clustering [23]. Clusters capture higher-order structures, aiding graph learning. Methods like DMon [24] use GNNs for soft cluster assignments, while Spectral Clustering (SC) [25] and K-means [16, 26] efficiently detect communities. DiffPool [27, 28] applies SC for pooling in GNNs.

Graph Coarsening. Graph Coarsening (GC) reduces a graph $\mathcal{G}(V, E, X)$ with N nodes and features $X \in \mathbb{R}^{N \times d}$ into a smaller graph $\mathcal{G}_c(\tilde{V}, \tilde{E}, \tilde{X})$ with $n \ll N$ nodes and $\tilde{X} \in \mathbb{R}^{n \times d}$. This is achieved via learning a coarsening matrix $\mathcal{P} \in \mathbb{R}^{n \times N}$, mapping similar nodes in $\mathcal{G}$ to super-nodes in $\mathcal{G}_c$, ensuring $\tilde{X} = \mathcal{P}X$ while preserving key properties [29–32].

2.2 Unsupervised Graph Structure Learning

Unsupervised graph learning spans from simple k-NN weighting [17, 33] to advanced statistical and graph signal processing (GSP) techniques. Statistical methods, also known as probabilistic graphical models, assume an underlying graph $\mathcal{G}$ governs the joint distribution of data $X \in \mathbb{R}^{N \times d}$ [19, 34, 35]. Some approaches [36] prune elements in the inverse sample covariance matrix $\hat{\Sigma} = \frac{1}{d-1} X X^T$ and sparse inverse covariance estimators, such as Graphical Lasso (GLasso) [19]: $\max_{\Theta} \log \det \Theta - \text{tr}(\hat{\Sigma}\Theta) - \rho \|\Theta\|_1$, where Θ is the inverse covariance matrix. However, these methods struggle with small sample sizes. Graph Signal Processing (GSP) techniques analyze signals on known graphs, ensuring properties like smoothness and sparsity. Signal smoothness on a graph $\mathcal{G}$ is quantified by the Laplacian quadratic form: $Q(\mathbf{L}) = \mathbf{x}^T \mathbf{L} \mathbf{x} = \frac{1}{2} \sum_{i,j} w_{ij} (\mathbf{x}(i) - \mathbf{x}(j))^2$. For a set of vectors X , smoothness is measured using the Dirichlet energy [37]: $\text{tr}(X^T L X)$. State-of-the-art methods [20, 21, 38] optimize Dirichlet energy while enforcing sparsity or specific structural constraints. Table 7 in Appendix D compares various graph learning methods based on their formulations and time complexities. More recently, SUBLIME [8] learns graph structure in an unsupervised manner by leveraging self-supervised contrastive learning to align a learnable graph with a dynamically refined anchor graph derived from the data itself.

Remark 1. Graph Structure Learning (GSL) differs significantly from Continual Learning (CL) [39–41] and Dynamic Graph Learning (DGL) [42–44], as discussed in Appendix C.

3 GraphFLEx

In this section, we introduce GraphFLEx, which has three main modules:

- **Graph Clustering.** Identifies communities and extracts higher-order structural information,
- **Graph Coarsening.** Is used to coarsen down the desired community, if the community itself is large,
- **Graph Learning.** Learns the graph’s structure using a limited subset of nodes from the clustering and coarsening modules, *enabling scalability*.

For pseudocode, see Algorithm 1 in Appendix G.

3.1 Incremental Graph Learning for Large Datasets

Real-world graph data is continuously expanding. For instance, e-commerce networks accumulate new clicks and purchases daily [45], while academic networks grow with new researchers and publications [46]. To manage such growth, we propose incrementally learning the graph structure over smaller segments.

Given a large dataset $\mathcal{L}(V_{\mathcal{L}}, X_{\mathcal{L}})$, where $V_{\mathcal{L}}$ is the node set and $X_{\mathcal{L}}$ represents node features, we define an *expanding dataset* setting $\mathcal{L}_{\mathcal{E}} = \{\mathcal{E}_{\tau=0}^T\}$. Initially, $\mathcal{L}$ is split into: (i) a *static dataset* $\mathcal{E}_0(V_0, X_0)$ and (ii) an *expanding dataset* $\mathcal{E} = \{\mathcal{E}_{\tau}(V_{\tau}, X_{\tau})\}_{\tau=1}^T$. Both *Goal 1* (large datasets with missing graph structure) and *Goal 2* (partially available graphs with incremental updates), discussed in Section 2, share the common objective of incrementally learning and updating the graph structure as new data arrives. GraphFLEx handles these by decomposing the problem into two key components:

- **Initial Graph $\mathcal{G}_0(V_0, A_0, X_0)$:** For *Goal 1*, where the graph structure is entirely missing, $\mathcal{E}_0(V_0, X_0)$ is used to construct $\mathcal{G}_0$ from scratch using structure learning methods (see Section 2.2). For *Goal 2*, the initial graph $\mathcal{G}_0(V_0, A_0, X_0)$ is already available and serves as the starting point for incremental updates.
- **Expanding Dataset $\mathcal{E} = \{\mathcal{E}_{\tau}(V_{\tau}, X_{\tau})\}_{\tau=1}^T$:** In both cases, $\mathcal{E}$ consists of incoming nodes and features arriving over T timestamps. These nodes are progressively integrated into the existing graph, enabling continuous adaptation and growth.

The partition is controlled by a parameter r , which determines the proportion of static nodes: $r = \frac{\|V_0\|}{\|V_{\mathcal{L}}\|}$. For example, $r = 0.2$ implies that 20% of $V_{\mathcal{L}}$ is treated as static, while the remaining 80% arrives incrementally over T timestamps. In our experiments, we set $r = 0.5$ and $T = 25$.

Remark 2. We can learn $\mathcal{G}_{\tau}(V_{\tau}, A_{\tau}, X_{\tau})$ by aggregating $\mathcal{E}_{\tau}$ nodes in $\mathcal{G}_{\tau-1}$ graph. Our goal is to learn $\mathcal{G}_T(V_T, A_T, X_T)$ after T^{th} -timestamp.

3.2 Detecting Communities

From the static graph $\mathcal{G}_0$, our goal is to learn higher-order structural information, identifying potential communities to which incoming nodes ($V \in V_{\tau}$) may belong. We train the community detection/clustering model $\mathcal{M}_{\text{clust}}$ once using $\mathcal{G}_0$, allowing subsequent inference of clusters for all incoming nodes. While our framework supports spectral and k-means clustering, our primary focus has been on Graph Neural Network (GNN)-based clustering methods. Specifically, we use DMoN [24, 47, 48], which maximizes spectral modularity. Modularity [49] measures the divergence between intra-cluster edges and the expected number. These methods use a GNN layer to compute the partition matrix $C = \text{softmax}(\text{MLP}(\tilde{X}, \theta_{\text{MLP}})) \in \mathbb{R}^{N \times K}$, where K is the number of clusters and $\tilde{X}$ is the updated feature embedding generated by one or more message-passing layers. To optimize the C matrix, we minimize the loss function $\Delta(C; A) = -\frac{1}{2m} \text{Tr}(C^T B C) + \frac{\sqrt{k}}{n} |\sum_i C_i^T|_F - 1$, which combines spectral modularity maximization with regularization to prevent trivial solutions, where B is the modularity matrix [24]. Our static graph $\mathcal{G}_0$ and incoming nodes $\mathcal{E}$ follow Assumption 1.

Assumption 1. *Based on the well-established homophily principle, which forms the basis of most graph coarsening and learning methods. We assume that the generated graphs adhere to the Degree-Corrected Stochastic Block Model (DC-SBM) [50], where intra-class (or intra-community) links are more likely than inter-class links. DC-SBM, an extension of SBM that accounts for degree heterogeneity, making it a more flexible and realistic choice for real-world networks.*

For more details on DC-SBM, see Appendix A.

Lemma 1. $\mathcal{M}_{\text{clust}}$ Consistency. *We adopt the theoretical framework of [50] for a DC-SBM with N nodes and k classes. The edge probability matrix is parameterized as $P_N = \rho_N P$, where $P \in \mathbb{R}^{k \times k}$ is a symmetric matrix containing the between/within community edge probabilities and it is independent of N , $\rho_N = \lambda_N/N$, and λ_N is the average degree of the network. Let*

181 $\hat{y}_N = [\hat{y}_1, \hat{y}_2, \dots, \hat{y}_N]$ denote the predicted class labels, and let $\hat{C}_N$ be the corresponding $N \times k$
 182 one-hot matrix. Let the true class label matrix is C_N , and μ is any $k \times k$ permutation matrix. Under
 183 the adjacency matrix $A^{(N)}$, the global maximum of the objective $\Delta(\cdot; A^{(N)})$ is denoted as $\hat{C}_N^*$. The
 184 consistency of class predictions is defined as:

185 **1. Strong Consistency.**

$$P_N \left[\min_{\mu} \|\hat{C}_N^* \mu - C_N\|_F^2 = 0 \right] \rightarrow 1 \quad \text{as } N \rightarrow \infty,$$

186 **2. Weak Consistency.**

$$\forall \varepsilon > 0, P_N \left[\min_{\mu} \frac{1}{N} \|\hat{C}_N^* \mu - C_N\|_F^2 < \varepsilon \right] \rightarrow 1 \quad \text{as } N \rightarrow \infty.$$

187 where $\|\cdot\|_F$ is the Frobenius norm. Under the conditions of Theorem 3.1 from [50]:

- 188 • The $\mathcal{M}_{\text{clust}}$ objective is strongly consistent if $\lambda_N / \log(N) \rightarrow \infty$, and
- 189 • It is weakly consistent when $\lambda_N \rightarrow \infty$.

190 **Remark 3. Structure Learning within Communities.** In *GraphFLEx*, we focus on learning the
 191 structure within each community rather than the structure of the entire dataset at once. Strong consis-
 192 tency ensures perfect community recovery, meaning no inter-community edges exist representing
 193 the ideal case. Weak consistency, however, allows for a small fraction (ϵ) of inter-community edges,
 194 where ϵ is controlled by ρ_n in $P_n = \rho_n P$, influencing graph sparsity.

195 By Lemma 1 and Assumption 1, stronger consistency leads to more precise structure learning,
 196 whereas weaker consistency permits a limited number of inter-community edges.

197 3.3 Learning Graph Structure on a Coarse Graph

198 After training $\mathcal{M}_{\text{clust}}$, we identify communities for incoming nodes, starting with $\tau = 1$. Once
 199 assigned, we determine significant communities those with at least one incoming node and learn their
 200 connections to the respective community subgraphs. For large datasets, substantial community sizes
 201 may again introduce scalability issues. To mitigate this, we first coarsen the large community graph
 202 into a smaller graph and use it to identify potential connections for incoming nodes. This process
 203 constitutes the second module of *GraphFLEx*, denoted as $\mathcal{M}_{\text{coar}}$, which employs LSH-based hashing
 204 for graph coarsening. The supernode index for i^{th} node is given as:

$$\mathcal{H}_i = \maxOccurance \left\{ \left\lfloor \frac{1}{r} \cdot (\mathcal{W} \cdot X_i + b) \right\rfloor \right\} \quad (1)$$

205 where r (bin width) controls the coarsened graph size, $\mathcal{W}$ represents random projection matrix, X is
 206 the feature matrix, and b is the bias term. For further details, refer to UGC [32]. After coarsening the
 207 i^{th} community (C_i), $\mathcal{M}_{\text{coar}}(C_i) = \{\mathcal{P}_i, S_i\}$ yields a partition matrix $\mathcal{P}_i \in \mathbb{R}^{\|S_i\| \times \|C_i\|}$ and a set of
 208 coarsened supernodes (S_i), as discussed in Section 2.

209 To identify potential connections for incoming nodes, we define their neighborhood as follows:

210 **Definition 1.** The neighborhood of a set of nodes $\mathcal{E}_i$ is defined as the union of the top k most similar
 211 nodes in C_i for each node $v \in \mathcal{E}_i$, where similarity is measured by the distance function $d(v, u)$. A
 212 node $u \in C_i$ is considered part of the neighborhood if its distance $d(v, u)$ is among the k smallest
 213 distances for all $u' \in C_i$.

$$\mathcal{N}_k(\mathcal{E}_i) = \bigcup_{v \in \mathcal{E}_i} \{u \in C_i \mid d(v, u) \leq \text{top-}k[d(v, u') : u' \in C_i]\}$$

214 **Goal 3.** The neighborhood of incoming nodes $\mathcal{N}_k(\mathcal{E}_i)$ represents the ideal set of nodes where the
 215 incoming nodes $\mathcal{E}_i$ are likely to establish connections when the entire community is provided to a
 216 structure learning framework.. A robust coarsening framework must reduce the number of nodes
 217 within each community C_i while ensuring that the neighborhood of the incoming nodes is preserved.

218 3.4 Graph Learning only with Potential Nodes

219 As we now have a smaller representation of the community, we can employ any graph learning
 220 algorithms discussed in Section 2.2 to learn a graph between coarsened supernodes S_i and incoming
 221 nodes ($V_\tau^i \in V_\tau$). This is the third module of *GraphFLEx*, i.e., graph learning; we denote it
 222 as $\mathcal{M}_{gl}$. The number of supernodes in S_i is much smaller compared to the original size of the
 223 community, i.e., $\|S_i\| \ll \|C_i\|$; scalability is not an issue now. We learn a small graph first
 224 using $\mathcal{M}_{gl}(S_i, X_\tau^i) = \tilde{\mathcal{G}}_\tau^i(V_\tau^c, A_\tau^c)$ where X_τ^i represents features of new nodes belonging to i^{th}

Table 1: Time complexity analysis of GraphFLEx. Here, N is the number of nodes in the graph, k is the number of nodes in the static subgraph used for clustering ($k \ll N$), and c represents the number of detected communities. k_τ denotes the number of nodes at timestamp τ . Finally, $\alpha = \|S_\tau^i\| + \|\mathcal{E}_\tau^i\|$ is the sum of coarsened and incoming nodes in the relevant community at τ timestamp.

	$\mathcal{M}_{clust}$	$\mathcal{M}_{coar}$	$\mathcal{M}_{gl}$	GraphFLEx
Best (kNN-UGC-ANN)	$\mathcal{O}(k^2)$	$\mathcal{O}\left(\frac{k_\tau}{c}\right)$	$\mathcal{O}(\alpha \log \alpha)$	$\mathcal{O}(k^2 + \frac{k_\tau}{c} + \alpha \log \alpha)$
Worst (SC-FGC-GLasso)	$\mathcal{O}(k^3)$	$\mathcal{O}\left(\left(\frac{k_\tau}{c}\right)^2 \ S_\tau^i\ \right)$	$\mathcal{O}(\alpha^3)$	$\mathcal{O}(k^3 + \left(\frac{k_\tau}{c}\right)^2 \ S_\tau^i\ + \alpha^3)$

community at time τ , $\tilde{\mathcal{G}}_\tau^i(V_\tau^c, A_\tau^c)$ representing the graph between supernodes and incoming nodes. Utilizing the partition matrix $\mathcal{P}_i$ obtained from $\mathcal{M}_{coar}$, we can precisely determine the set of nodes associated with each supernode. For every new node $V \in V_\tau^i$, we identify the connected supernodes and subsequently select nodes within those supernodes. This subset of nodes is denoted by $\omega_{V_\tau^i}$, the sub-graph associated with $\omega_{V_\tau^i}$ represented by $\mathcal{G}_{\tau-1}^i(\omega_{V_\tau^i})$ then undergoes an additional round of graph learning $\mathcal{M}_{gl}(\mathcal{G}_{\tau-1}^i(\omega_{V_\tau^i}), X_\tau^i)$, ultimately providing a clear and accurate connection of new nodes V_τ^i with nodes of $\mathcal{G}_{\tau-1}$, ultimately updating it to $\mathcal{G}_\tau$. This multi-step approach, characterized by coarsening, learning on coarsened graphs, and translation to the original graph, ensures scalability.

Theorem 1. Neighborhood Preservation. Let $\mathcal{N}_k(\mathcal{E}_i)$ denote the neighborhood of incoming nodes $\mathcal{E}_i$ for the i^{th} community. With partition matrix $\mathcal{P}_i$ and $\mathcal{M}_{gl}(S_i, X_\tau^i) = \mathcal{G}_\tau^c(V_\tau^c, A_\tau^c)$ we identify the supernodes connected to incoming nodes $\mathcal{E}_i$ and subsequently select nodes within those supernodes; this subset of nodes is denoted by $\omega_{V_\tau^i}$. Formally,

$$\omega_{V_\tau^i} = \bigcup_{v \in \mathcal{E}_i} \left\{ \bigcup_{s \in S_i} \{\pi^{-1}(s) | A_\tau^c(v, s) \neq 0\} \right\}$$

Then, with probability $\Pi_{\{c \in \phi\}} p(c)$, it holds that $\mathcal{N}_k(\mathcal{E}_i) \subseteq \omega_{V_\tau^i}$ where

$$p(c) \leq 1 - \frac{2}{\sqrt{2\pi}} \frac{c}{r} \left[1 - e^{-r^2/(2c^2)} \right],$$

and ϕ is a set containing all pairwise distance values ($c = \|v - u\|$) between every node $v \in \mathcal{E}_i$ and the nodes $u \in \omega_{V_\tau^i}$. Here, $\pi^{-1}(s)$ denotes the set of nodes mapped to supernode s , r is the bin-width hyperparameter of $\mathcal{M}_{coar}$.

Proof. The proof is deferred in Appendix B. □

Remark 4. Theorem 1 establishes that, with a constant probability of success, the neighborhood of incoming nodes $\mathcal{N}_k(\mathcal{E}_i)$ can be effectively recovered using the GraphFLEx multistep approach, which involves coarsening and learning on the coarsened graph, i.e., $\mathcal{N}_k(\mathcal{E}_i) \subseteq \omega_{V_\tau^i}$. The set $\omega_{V_\tau^i}$, estimated by GraphFLEx, identifies potential candidates where incoming nodes are likely to connect. The probability of failure can be reduced by regulating the average degree of connectivity in $\mathcal{M}_{gl}(S_i, X_\tau^i) = \mathcal{G}_\tau^c(V_\tau^c, A_\tau^c)$. While a fully connected network $\mathcal{G}_\tau^c$ ensures all nodes in the community are candidates, it significantly increases computational costs for large communities.

3.5 GraphFLEx: Multiple SGL Frameworks

Each module in Figure 3 controls a distinct aspect of the graph learning process: clustering influences community detection, coarsening reduces graph complexity via supernodes, and the learning module governs structural inference. Altering any of these modules results in a new graph learning method. Currently, we support 48 different graph learning configurations, and this number scales exponentially with the addition of new methods to any module. The number of possible frameworks is given by $\alpha \times \beta \times \gamma$, where α , β , and γ represent the number of clustering, coarsening, and learning methods, respectively.

3.6 Run Time Analysis

GraphFLEx computational time is always bounded by existing approaches, as it operates on a significantly reduced set of nodes. We evaluate the run-time complexity of GraphFLEx in two

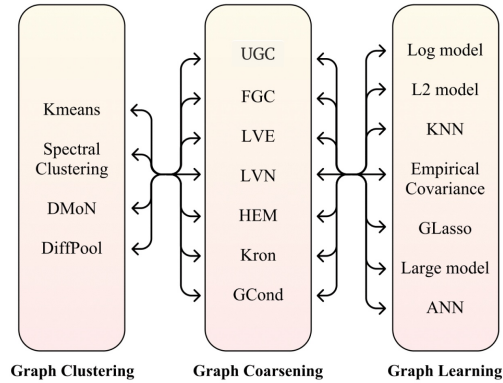


Figure 3: The versatility of GraphFlex in supporting multiple GSL methods.

scenarios: (a) the worst-case scenario, where computationally intensive clustering and coarsening modules are selected, providing an upper bound on time complexity, and (b) the best-case scenario, where the most efficient modules are chosen. Table 1 presents a summary of this analysis for both cases. Due to space limitations, a more comprehensive analysis is provided in Appendix E.

4 Experiments

Tasks and Datasets. To validate GraphFLEX’s utility, we evaluate it across four key dimensions: (i) computational efficiency, (ii) scalability to large graphs, (iii) quality of learned structures, and (iv) adaptability to dynamically growing graphs. To validate the characteristics of GraphFLEX, we conduct extensive experiments on 26 different datasets, including (a) datasets that already have a complete graph structure (allowing comparison between the learned and the original structure), (b) datasets with missing graph structures, (c) synthetic datasets, (d) small datasets for visualizing the graph structure, and (e) large datasets, including datasets with even 2.4M nodes. More details about datasets and *system specifications* are presented in Table 8 in Appendix F.

Table 2: Computational time(in seconds) for learning graph structures using GraphFLEX (GFlex) with existing methods (Vanilla referred to as Van.). The experimental setup involves treating 50% of the data as static, while the remaining 50% of nodes are treated as incoming nodes arriving in 25 different timestamps. The best times are highlighted by color **Green**. OOM and OOT denote out-of-memory and out-of-time, respectively.

Data	ANN		KNN		log-model		l2-model		emp-Covar.		large-model		Sublime	
	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex
Cora	335	100	8.4	36.1	869	81.6	424	55	8.6	30	2115	18.4	7187	493
Citeseer	1535	454	21.9	75	1113	64.5	977	54.0	14.7	59.2	8319	43.9	8750	670
DBLP	2731	988	OOM	270	77000	919	OOT	1470	359	343	OOT	299	OOM	831
CS	22000	12000	OOM	789	OOT	838	32000	809	813	718	OOT	1469	OOM	1049
PubMed	770	227	OOM	164	OOT	176	OOT	165	488	299	OOT	262	OOM	914
Phy.	61000	21000	OOM	903	OOT	959	OOT	908	2152	1182	OOT	2414	OOM	2731
Syn 3	95	37	OOM	30	58000	346	859	53	88	59	5416	42	6893	780
Syn 4	482	71	OOM	73	OOT	555	OOT	145	2072	1043	OOT	392	OOM	1896

Table 3: Node classification accuracies on different GNN models using GraphFLEX (GFlex) with existing Vanilla (Van.) methods. The experimental setup involves treating 70% of the data as static, while the remaining 30% of nodes are treated as new nodes coming in 25 different timestamps. The best and the second-best accuracies in each row are highlighted by dark and lighter shades of **Green**, respectively. GraphFLEX’s structure beats all of the vanilla structures for every dataset. OOM and OOT denotes out-of-memory and out-of-time respectively.

Data	Model	ANN		KNN		log-model		l2-model		COVAR		large-model		Sublime		Base Struct.
		Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	
DBLP	GAT	34.23	67.37	OOM	69.83	OOT	69.83	OOT	68.98	50.48	68.56	OOT	66.38	OOM	68.32	70.84
	SAGE	34.23	69.58	OOM	70.28	OOT	70.28	OOT	70.68	51.47	70.51	OOT	69.32	OOM	70.28	72.57
	GCN	34.12	69.41	OOM	73.39	OOT	73.39	OOT	73.05	51.50	71.75	OOT	68.55	OOM	69.06	74.43
	GIN	34.01	69.69	OOM	68.19	OOT	68.19	OOT	73.08	52.77	72.03	OOT	71.18	OOM	71.87	73.92
CS	GAT	12.47	60.89	OOM	61.09	OOT	60.95	18.64	61.06	58.96	88.06	OOT	86.22	OOM	64.21	60.75
	SAGE	12.70	78.81	OOM	79.43	OOT	79.06	19.24	78.94	56.97	93.30	OOT	92.79	OOM	78.94	80.33
	GCN	12.59	63.81	OOM	67.94	OOT	69.33	19.21	66.01	58.35	91.07	OOT	84.85	OOM	68.92	67.43
	GIN	13.07	77.62	OOM	78.41	OOT	78.55	19.24	77.61	58.26	92.07	OOT	86.03	OOM	77.61	55.65
Pub.	GAT	49.49	83.71	OOM	84.60	OOT	84.60	OOT	84.04	72.63	83.97	OOT	81.15	OOM	82.15	84.04
	SAGE	50.43	87.27	OOM	87.34	OOT	87.34	OOT	87.42	73.57	86.68	OOT	87.34	OOM	83.45	88.88
	GCN	50.45	82.06	OOM	83.56	OOT	83.56	OOT	83.74	73.14	82.39	OOT	78.03	OOM	70.94	85.54
	GIN	51.82	83.13	OOM	84.31	OOT	84.07	OOT	82.93	73.15	83.51	OOT	82.85	OOM	80.72	86.50
Phy.	GAT	29.18	88.06	OOM	88.47	OOT	88.47	OOT	88.68	58.96	88.06	OOT	86.22	OOM	86.12	88.58
	SAGE	29.57	93.47	OOM	93.47	OOT	93.47	OOT	93.78	56.97	93.60	OOT	92.79	OOM	89.58	94.19
	GCN	27.84	91.27	OOM	91.08	OOT	91.08	OOT	91.78	58.35	91.07	OOT	84.85	OOM	88.46	91.48
	GIN	28.38	92.69	OOM	92.04	OOT	92.04	OOT	92.27	58.26	92.07	OOT	86.03	OOM	87.20	88.89

4.1 Computational Efficiency.

Existing methods like k -NN and *log*-model struggle to learn graph structures even for 20k nodes due to out-of-memory (OOM) or out-of-time (OOT) issues, while *l2*-model and *large*-model struggle beyond 50k nodes. Although *A*-NN and *emp*-Covar. are faster, GraphFLEX outperforms them on sufficiently large graphs (Table 2). While traditional methods may be efficient for small graphs, GraphFLEX scales significantly better, excelling on large datasets like *Pubmed* and *Syn 5*, where most methods fail. It accelerates structure learning, making *A*-NN 3× faster and *emp*-Covar. 2× faster.

4.2 Node Classification Accuracy

Experimental Setup. We now evaluate the prediction performance of GNN models when trained on graph structures learned from three distinct scenarios: **1) Original Structure:** GNN models trained

on the original graph structure, which we refer to as the Base Structure, **2) GraphFLEX Structure:** GNN models trained on the graph structure learned from GraphFLEX, and **3) Vanilla Structure:** GNN models trained on the graph structure learned from other existing methods.

For each scenario, a unique graph structure is obtained. We trained GNN models on each of these three structure. For more details on GNN model parameters, see Appendix H.

GNN Models. Graph neural networks (GNNs) such as *GCN* [51], *GraphSage* [52], *GIN* [53], and *GAT* [54] rely on accurate message passing, dictated by the graph structure, for effective embedding. We use these models to evaluate the above-mentioned learned structures. Table 3 reports node classification performance across all methods. Notably, GraphFLEX outperforms vanilla structures by a significant margin across all datasets, achieving accuracies close to those obtained with the original structure. Figure 8 in Appendix H illustrates *GraphSage* classification results, highlighting GraphFLEX’s superior performance. For the *CS* dataset, GraphFLEX (*large-model*) and GraphFLEX (*empCovar.-model*) even surpass the original structure, demonstrating its ability to preserve key structural properties while denoising edges, leading to improved accuracy.

4.3 Scalability of GraphFLEX on Large-Scale Graphs.

To comprehensively evaluate GraphFLEX’s scalability to large-scale graphs, we consider four datasets with a high number of nodes: (a) *Flickr* (89k nodes) [55], (b) *Reddit* (233k nodes) [55], (c) *Ogbn-arxiv* (169k nodes) [46], and (d) *Ogbn-products* (2.4M nodes) [56]. As shown in Table 4, GraphFLEX consistently demonstrates superior scalability across all datasets, outperforming all baseline methods in runtime. In particular, methods such as *log-model*, *l2-model*, and *large-model* fail to run even on *Flickr*, while GraphFLEX successfully scales them on *Flickr*, *Ogbn-arxiv*, and *Reddit*, enabling structure learning where others cannot. For the most computationally demanding dataset, *Ogbn-products*, these methods remain prohibitively expensive even for GraphFLEX. Nonetheless, GraphFLEX efficiently supports scalable structure learning on *Ogbn-products* using the *Covar*, *ANN*, and *KNN* modules. Table 4 also reports node classification accuracy, demonstrating that GraphFLEX maintains performance comparable to the original (base) structure across all datasets. These results confirm that GraphFLEX not only scales effectively, but also preserves the quality of learned structures.

Table 4: Runtime (sec) and Node Classification Accuracy (%) across large datasets. Each cell shows: **Time / Accuracy**. Van = Vanilla, GFlex = GraphFLEX. OOM = Out of Memory, OOT = Out of Time.

Method	ogbn-arxiv (60.13)		ogbn-products (73.72)		Flickr (44.92)		Reddit (94.15)	
	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex
Covar	OOM –	3.7k 60.26	OOM –	83.1k 68.23	2.3k 44.65	682 44.34	OOM –	6.6k 94.13
ANN	7.8k 60.14	4.8k 60.22	OOM –	89.3k 67.91	2.5k 44.09	705 44.92	12.6k 94.14	6.1k 94.18
knn	8.3k 60.09	6.1k 60.23	OOM –	91.8k 68.47	2.7k 43.95	920 44.73	15.6k 94.14	6.9k 94.15
l2	OOT –	9.1k 58.39	OOT –	OOT –	93.3k 44.90	1.2k 44.32	OOT –	5.1 93.47
log	OOT –	45.6k 58.72	OOT –	OOT –	OOT –	18.7k 44.59	OOT –	60.3k 94.13
large	OOT –	5.6k 60.20	OOT –	OOT –	OOT –	2.2k 44.45	OOT –	9.3k 93.71

4.4 GraphFLEX for Link Prediction and Graph Classification.

To further validate the generalization of our framework, we evaluate GraphFLEX on the link prediction task. The results are presented in Table 5, following the same setting as Table 3. The structure learned by GraphFLEX demonstrates strong predictive performance, in some cases even outperforming the base structure. This highlights the effectiveness of GraphFLEX in preserving and even enhancing relational information relevant for link prediction. While our primary focus is on structure learning

Table 5: Link predication accuracy (%) across different datasets. The best and the second-best accuracies in each row are highlighted by dark and lighter shades of Green, respectively.

Data	ANN		KNN		log-model		l2-model		COVAR		large-model		Base Struct.
	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	
DBLP	96.57	96.61	OOM	94.23	OOT	97.59	OOT	97.59	97.22	97.59	OOT	96.24	95.13
Citeseer	80.12	96.32	85.17	96.24	80.48	96.24	80.48	96.48	82.05	96.24	84.50	94.38	90.78
Cora	84.47	95.30	79.23	95.14	90.63	95.45	90.81	95.14	86.05	95.30	90.63	94.67	89.53
Pubmed	94.24	96.91	OOM	97.42	OOT	97.42	OOT	97.37	94.89	94.64	OOT	94.41	94.64
CS	94.21	95.73	OOM	96.02	OOT	93.17	OOT	93.17	93.52	92.31	OOT	95.73	95.00
Physics	95.77	91.34	OOM	94.63	OOT	90.79	OOT	94.63	92.03	90.79	OOT	92.97	93.96

for node-level tasks, we briefly discuss the applicability of GraphFLEX to graph classification. In such tasks, especially in domains like molecule or drug discovery, each data point often corresponds to a small individual subgraph. For these cases, applying clustering and coarsening is typically redundant and may introduce unnecessary computational overhead. Nevertheless, GraphFLEX

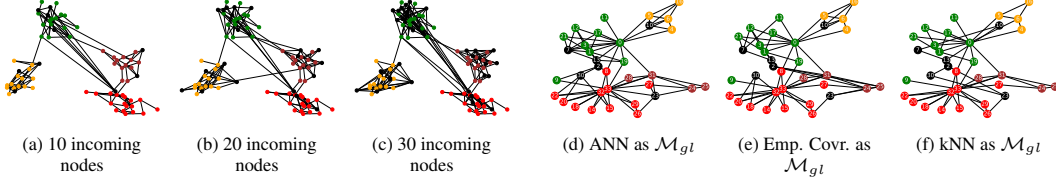


Figure 4: Figures (a), (b), and (c) illustrate the growing structure learned using GraphFLEX for *HE* synthetic dataset. Figures (d), (e), and (f) illustrate the learned structure on Zachary’s karate dataset when existing methods are employed with GraphFLEX. New nodes are denoted using black color.

remains flexible—its learning module can be directly used without the clustering or coarsening steps, making it suitable for graph classification as well. This adaptability reinforces GraphFLEX’s utility across a broad range of graph learning tasks.

4.5 Clustering Quality

We measure three metrics to evaluate the resulting clusters or community assignments: a) Normalized Mutual Information (NMI) [24] between the cluster assignments and original labels; b) Conductance ($\mathcal{C}$) [57] which measures the fraction of total edge volume that points outside the cluster; and c) Modularity ($\mathcal{Q}$) [49] which measures the divergence between the intra-community edges and the expected one. Table 6 illustrates these metrics for single-cell RNA and the MNIST dataset (where the whole structure is missing), and Figure 12 in Appendix K shows the PHATE [58] visualization of clusters learned using GraphFLEX’s clustering module $\mathcal{M}_{clust}$. We also train the aforementioned GNN models for the node classification task in order to illustrate the efficacy of the learned structures; the accuracy values presented in Table 6, clearly highlight the significance of the learned structures, as reflected by the high accuracy values.

4.6 Structure Visualization

We evaluate the structures generated by GraphFLEX through visualizations on four small datasets: (i) MNIST [59], consisting of hand-written digit images, where Figure 5(a) shows that images of the same digit are mostly connected; (ii) Pre-trained GloVe embeddings [60] of English words, with Figure 5(b) revealing that frequently used words are closely connected; (iii) A synthetic *HE* dataset (see Appendix F), demonstrating GraphFLEX’s ability to handle expanding networks without requiring full relearning. Figure 4(a-c) shows the graph structure evolving as 30 new nodes are added over three timestamps; and (iv) Zachary’s karate club network [61], which highlights GraphFLEX’s multi-framework capability. Figure 4(d-f) shows three distinct graph structures after altering the learning module. For a comprehensive ablation study, refer to Appendix L.

Table 6: Clustering (NMI, $\mathcal{C}$, $\mathcal{Q}$) and node classification accuracy using GCN, GraphSAGE, GIN, and GAT.

Data	NMI	$\mathcal{C}$	$\mathcal{Q}$	GCN	SAGE	GIN	GAT
Bar. M.	0.716	0.057	0.741	91.2	96.2	95.1	94.9
Seger.	0.678	0.102	0.694	91.0	93.9	94.2	92.3
Mura.	0.843	0.046	0.706	96.9	97.4	97.5	96.4
Bar. H.	0.674	0.078	0.749	95.3	96.4	97.2	95.8
Xin	0.741	0.045	0.544	98.6	99.3	98.9	99.8
MNIST	0.677	0.082	0.712	92.9	94.5	94.9	82.6

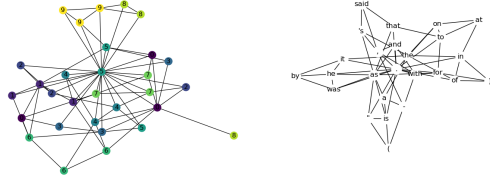


Figure 5: Effectiveness of our framework in learning structure between similar MNIST digits and GloVe embeddings.

5 Conclusion

Large or expanding graphs challenge the best of graph learning approaches. GraphFLEX, introduced in this paper, seamlessly adds new nodes into an existing graph structure. It offers diverse methods for acquiring the graph’s structure. GraphFLEX consists of three key modules: Clustering, Coarsening, and Learning which empowers GraphFLEX to serve as a comprehensive framework applicable individually for clustering, coarsening, and learning tasks. Empirically, GraphFLEX outperforms state-of-the-art baselines, achieving up to 3× speedup while preserving structural quality. It achieves accuracies close to training on the original graph, in most instances. The performance across multiple real and synthetic datasets affirms the utility and efficacy of GraphFLEX for graph structure learning. **Limitations and Future Work.** GraphFLEX is designed assuming minimal inter-community connectivity, which aligns well with many real-world scenarios. However, its applicability to heterophilic graphs may require further adaptation. Future work will focus on extending the framework to supervised GSL methods and heterophilic graphs, broadening its scalability and versatility.

References

- [1] J. Zhou, G. Cui, S. Hu, Z. Zhang, C. Yang, Z. Liu, L. Wang, C. Li, and M. Sun, “Graph neural networks: A review of methods and applications,” *AI open*, vol. 1, pp. 57–81, 2020. **(Cited at p. 1.)**
- [2] A. Fout, J. Byrd, B. Shariat, and A. Ben-Hur, “Protein interface prediction using graph convolutional networks,” *Advances in neural information processing systems*, vol. 30, 2017. **(Not cited.)**
- [3] Y. Wu, D. Lian, Y. Xu, L. Wu, and E. Chen, “Graph convolutional networks with markov random field reasoning for social spammer detection,” in *Proceedings of the AAAI conference on artificial intelligence*, vol. 34, pp. 1054–1061, 2020. **(Not cited.)**
- [4] N. Malik, R. Gupta, and S. Kumar, “Hyperdefender: A robust framework for hyperbolic gnns,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 39, pp. 19396–19404, Apr. 2025. **(Cited at p. 1.)**
- [5] J. Gu, H. Zhao, Z. Lin, S. Li, J. Cai, and M. Ling, “Scene graph generation with external knowledge and image reconstruction,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pp. 1969–1978, 2019. **(Cited at p. 1.)**
- [6] Y. Zhu, W. Xu, J. Zhang, Y. Du, J. Zhang, Q. Liu, C. Yang, and S. Wu, “A survey on graph structure learning: Progress and opportunities,” *arXiv preprint arXiv:2103.03036*, 2021. **(Cited at p. 1.)**
- [7] J. D. Allen, Y. Xie, M. Chen, L. Girard, and G. Xiao, “Comparing statistical methods for constructing large scale gene networks,” *PloS one*, vol. 7, no. 1, p. e29348, 2012. **(Cited at p. 1.)**
- [8] Y. Liu, Y. Zheng, D. Zhang, H. Chen, H. Peng, and S. Pan, “Towards unsupervised deep graph structure learning,” in *Proceedings of the ACM Web Conference 2022*, pp. 1392–1403, 2022. **(Cited at pp. 1 and 3.)**
- [9] Y. Chen and L. Wu, “Graph neural networks: Graph structure learning,” *Graph Neural Networks: Foundations, Frontiers, and Applications*, pp. 297–321, 2022. **(Cited at p. 1.)**
- [10] Q. Wu, W. Zhao, Z. Li, D. P. Wipf, and J. Yan, “Nodeformer: A scalable graph structure learning transformer for node classification,” *Advances in Neural Information Processing Systems*, vol. 35, pp. 27387–27401, 2022. **(Cited at p. 1.)**
- [11] W. Jin, Y. Ma, X. Liu, X. Tang, S. Wang, and J. Tang, “Graph structure learning for robust graph neural networks,” in *Proceedings of the 26th ACM SIGKDD international conference on knowledge discovery & data mining*, pp. 66–74, 2020. **(Cited at p. 1.)**
- [12] D. Lao, X. Yang, Q. Wu, and J. Yan, “Variational inference for training graph neural networks in low-data regime through joint structure-label estimation,” in *Proceedings of the 28th ACM SIGKDD conference on knowledge discovery and data mining*, pp. 824–834, 2022. **(Cited at p. 1.)**
- [13] B. Fatemi, L. El Asri, and S. M. Kazemi, “Slaps: Self-supervision improves structure learning for graph neural networks,” *Advances in Neural Information Processing Systems*, vol. 34, pp. 22667–22681, 2021. **(Cited at p. 1.)**
- [14] W. Dong, C. Moses, and K. Li, “Efficient k-nearest neighbor graph construction for generic similarity measures,” in *Proceedings of the 20th international conference on World wide web*, pp. 577–586, 2011. **(Cited at p. 1.)**
- [15] M. Muja and D. G. Lowe, “Scalable nearest neighbor algorithms for high dimensional data,” *IEEE transactions on pattern analysis and machine intelligence*, vol. 36, no. 11, pp. 2227–2240, 2014. **(Cited at p. 1.)**
- [16] J. MacQueen *et al.*, “Some methods for classification and analysis of multivariate observations,” in *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*, vol. 1, pp. 281–297, Oakland, CA, USA, 1967. **(Cited at pp. 1 and 3.)**

- 425 [17] F. Wang and C. Zhang, “Label propagation through linear neighborhoods,” in *Proceedings of*
 426 *the 23rd international conference on Machine learning*, pp. 985–992, 2006. **(Cited at pp. 1**
 427 **and 3.)**
- 428 [18] C.-J. Hsieh, I. Dhillon, P. Ravikumar, and M. Sustik, “Sparse inverse covariance matrix esti-
 429 mation using quadratic approximation,” *Advances in neural information processing systems*,
 430 vol. 24, 2011. **(Cited at p. 1.)**
- 431 [19] J. Friedman, T. Hastie, and R. Tibshirani, “Sparse inverse covariance estimation with the
 432 graphical lasso,” *Biostatistics*, vol. 9, no. 3, pp. 432–441, 2008. **(Cited at pp. 1 and 3.)**
- 433 [20] X. Dong, D. Thanou, P. Frossard, and P. Vandergheynst, “Learning laplacian matrix in smooth
 434 graph signal representations,” *IEEE Transactions on Signal Processing*, vol. 64, no. 23, pp. 6160–
 435 6173, 2016. **(Cited at pp. 1 and 3.)**
- 436 [21] V. Kalofolias, “How to learn a graph from smooth signals,” in *Artificial intelligence and*
 437 *statistics*, pp. 920–929, PMLR, 2016. **(Cited at pp. 1 and 3.)**
- 438 [22] M. Hashemi, S. Gong, J. Ni, W. Fan, B. A. Prakash, and W. Jin, “A comprehensive
 439 survey on graph reduction: Sparsification, coarsening, and condensation,” *arXiv preprint*
 440 *arXiv:2402.03358*, 2024. **(Cited at p. 3.)**
- 441 [23] S. Fortunato, “Community detection in graphs,” *Physics reports*, vol. 486, no. 3-5, pp. 75–174,
 442 2010. **(Cited at p. 3.)**
- 443 [24] A. Tsitsulin, J. Palowitch, B. Perozzi, and E. Müller, “Graph clustering with graph neural
 444 networks,” *Journal of Machine Learning Research*, vol. 24, no. 127, pp. 1–21, 2023. **(Cited at**
 445 **pp. 3, 4, and 9.)**
- 446 [25] S. D. Kamvar, D. Klein, and C. D. Manning, “Spectral learning,” in *IJCAI*, vol. 3, pp. 561–566,
 447 2003. **(Cited at p. 3.)**
- 448 [26] K. Wagstaff, C. Cardie, S. Rogers, S. Schrödl, *et al.*, “Constrained k-means clustering with
 449 background knowledge,” in *Icml*, vol. 1, pp. 577–584, 2001. **(Cited at p. 3.)**
- 450 [27] J. Bruna, W. Zaremba, A. Szlam, and Y. LeCun, “Spectral networks and deep locally connected
 451 networks on graphs. arxiv,” *arXiv preprint arXiv:1312.6203*, 2014. **(Cited at p. 3.)**
- 452 [28] M. Defferrard, X. Bresson, and P. Vandergheynst, “Convolutional neural networks on graphs
 453 with fast localized spectral filtering,” *Advances in neural information processing systems*, vol. 29,
 454 2016. **(Cited at p. 3.)**
- 455 [29] A. Loukas, “Graph reduction with spectral and cut guarantees,” *J. Mach. Learn. Res.*, vol. 20,
 456 no. 116, pp. 1–42, 2019. **(Cited at p. 3.)**
- 457 [30] M. Kataria, A. Khandelwal, R. Das, S. Kumar, and J. Jayadeva, “Linear complexity framework
 458 for feature-aware graph coarsening via hashing,” in *NeurIPS 2023 Workshop: New Frontiers in*
 459 *Graph Learning*, 2023. **(Cited at pp. 15 and 16.)**
- 460 [31] M. Kumar, A. Sharma, and S. Kumar, “A unified framework for optimization-based graph
 461 coarsening,” *Journal of Machine Learning Research*, vol. 24, no. 118, pp. 1–50, 2023. **(Not**
 462 **cited.)**
- 463 [32] M. Kataria, S. Kumar, and J. Jayadeva, “UGC: Universal graph coarsening,” in *The Thirty-*
 464 *eighth Annual Conference on Neural Information Processing Systems*, 2024. **(Cited at pp. 3, 5,**
 465 **and 14.)**
- 466 [33] X. Zhu, Z. Ghahramani, and J. D. Lafferty, “Semi-supervised learning using gaussian fields and
 467 harmonic functions,” in *Proceedings of the 20th International conference on Machine learning*
 468 *(ICML-03)*, pp. 912–919, 2003. **(Cited at p. 3.)**
- 469 [34] D. Koller and N. Friedman, *Probabilistic graphical models: principles and techniques*. MIT
 470 press, 2009. **(Cited at p. 3.)**

- [35] O. Banerjee, L. El Ghaoui, and A. d’Aspremont, “Model selection through sparse maximum likelihood estimation for multivariate gaussian or binary data,” *The Journal of Machine Learning Research*, vol. 9, pp. 485–516, 2008. **(Cited at p. 3.)**
- [36] A. P. Dempster, “Covariance selection,” *Biometrics*, pp. 157–175, 1972. **(Cited at p. 3.)**
- [37] M. Belkin, P. Niyogi, and V. Sindhvani, “Manifold regularization: A geometric framework for learning from labeled and unlabeled examples,” *Journal of machine learning research*, vol. 7, no. 11, 2006. **(Cited at p. 3.)**
- [38] C. Hu, L. Cheng, J. Sepulcre, G. El Fakhri, Y. M. Lu, and Q. Li, “A graph theoretical regression model for brain connectivity learning of alzheimer’s disease,” in *2013 IEEE 10th International Symposium on Biomedical Imaging*, pp. 616–619, IEEE, 2013. **(Cited at p. 3.)**
- [39] G. M. Van de Ven and A. S. Tolias, “Three scenarios for continual learning,” *arXiv preprint arXiv:1904.07734*, 2019. **(Cited at pp. 3 and 15.)**
- [40] X. Zhang, D. Song, and D. Tao, “Cglib: Benchmark tasks for continual graph learning,” *Advances in Neural Information Processing Systems*, vol. 35, pp. 13006–13021, 2022. **(Not cited.)**
- [41] G. I. Parisi, R. Kemker, J. L. Part, C. Kanan, and S. Wermter, “Continual lifelong learning with neural networks: A review,” *Neural networks*, vol. 113, pp. 54–71, 2019. **(Cited at pp. 3 and 15.)**
- [42] S. Kim, S. Yun, and J. Kang, “Dygrain: An incremental learning framework for dynamic graphs,” in *IJCAI*, pp. 3157–3163, 2022. **(Cited at pp. 3 and 15.)**
- [43] T. Wu, Q. Liu, Y. Cao, Y. Huang, X.-M. Wu, and J. Ding, “Continual graph convolutional network for text classification,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, pp. 13754–13762, 2023. **(Not cited.)**
- [44] J. You, T. Du, and J. Leskovec, “Roland: graph learning framework for dynamic graphs,” in *Proceedings of the 28th ACM SIGKDD conference on knowledge discovery and data mining*, pp. 2358–2366, 2022. **(Cited at pp. 3 and 15.)**
- [45] L. Xiang, Q. Yuan, S. Zhao, L. Chen, X. Zhang, Q. Yang, and J. Sun, “Temporal recommendation on graphs via long-and short-term preference fusion,” in *Proceedings of the 16th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 723–732, 2010. **(Cited at p. 4.)**
- [46] K. Wang, Z. Shen, C. Huang, C.-H. Wu, Y. Dong, and A. Kanakia, “Microsoft academic graph: When experts are not enough,” *Quantitative Science Studies*, vol. 1, no. 1, pp. 396–413, 2020. **(Cited at pp. 4, 8, and 16.)**
- [47] F. M. Bianchi, D. Grattarola, and C. Alippi, “Spectral clustering with graph neural networks for graph pooling,” in *International conference on machine learning*, pp. 874–883, PMLR, 2020. **(Cited at p. 4.)**
- [48] F. M. Bianchi, “Simplifying clustering with graph neural networks,” *arXiv preprint arXiv:2207.08779*, 2022. **(Cited at p. 4.)**
- [49] M. E. Newman, “Modularity and community structure in networks,” *Proceedings of the national academy of sciences*, vol. 103, no. 23, pp. 8577–8582, 2006. **(Cited at pp. 4 and 9.)**
- [50] Y. Zhao, E. Levina, and J. Zhu, “Consistency of community detection in networks under degree-corrected stochastic block models,” 2012. **(Cited at pp. 4, 5, and 14.)**
- [51] T. N. Kipf and M. Welling, “Semi-supervised classification with graph convolutional networks,” *arXiv preprint arXiv:1609.02907*, 2016. **(Cited at pp. 8 and 15.)**
- [52] W. Hamilton, Z. Ying, and J. Leskovec, “Inductive representation learning on large graphs,” *Advances in neural information processing systems*, vol. 30, 2017. **(Cited at p. 8.)**
- [53] K. Xu, W. Hu, J. Leskovec, and S. Jegelka, “How powerful are graph neural networks?,” *arXiv preprint arXiv:1810.00826*, 2018. **(Cited at p. 8.)**

- [54] P. Velickovic, G. Cucurull, A. Casanova, A. Romero, P. Lio, Y. Bengio, *et al.*, “Graph attention networks,” *stat*, vol. 1050, no. 20, pp. 10–48550, 2017. **(Cited at p. 8.)**
- [55] H. Zeng, H. Zhou, A. Srivastava, R. Kannan, and V. Prasanna, “Graphsaint: Graph sampling based inductive learning method,” *arXiv preprint arXiv:1907.04931*, 2019. **(Cited at pp. 8 and 16.)**
- [56] K. Bhatia, K. Dahiya, H. Jain, P. Kar, A. Mittal, Y. Prabhu, and M. Varma, “The extreme classification repository: Multi-label datasets and code,” 2016. **(Cited at pp. 8 and 16.)**
- [57] M. Jerrum and A. Sinclair, “Conductance and the rapid mixing property for markov chains: the approximation of permanent resolved,” in *Proceedings of the twentieth annual ACM symposium on Theory of computing*, pp. 235–244, 1988. **(Cited at p. 9.)**
- [58] K. R. Moon, D. Van Dijk, Z. Wang, S. Gigante, D. B. Burkhardt, W. S. Chen, K. Yim, A. v. d. Elzen, M. J. Hirn, R. R. Coifman, *et al.*, “Visualizing structure and transitions in high-dimensional biological data,” *Nature biotechnology*, vol. 37, no. 12, pp. 1482–1492, 2019. **(Cited at pp. 9 and 19.)**
- [59] Y. LeCun, C. Cortes, and C. Burges, “Mnist handwritten digit database,” *ATT Labs [Online]*. Available: <http://yann.lecun.com/exdb/mnist>, vol. 2, 2010. **(Cited at pp. 9 and 16.)**
- [60] J. Pennington, R. Socher, and C. D. Manning, “Glove: Global vectors for word representation,” in *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, pp. 1532–1543, 2014. **(Cited at pp. 9 and 16.)**
- [61] W. W. Zachary, “An information flow model for conflict and fission in small groups,” *Journal of anthropological research*, vol. 33, no. 4, pp. 452–473, 1977. **(Cited at pp. 9, 16, and 22.)**
- [62] M. Datar, N. Immorlica, P. Indyk, and V. S. Mirrokni, “Locality-sensitive hashing scheme based on p-stable distributions,” in *Proceedings of the twentieth annual symposium on Computational geometry*, pp. 253–262, 2004. **(Cited at p. 14.)**
- [63] L. Lü and T. Zhou, “Link prediction in complex networks: A survey,” *Physica A: statistical mechanics and its applications*, vol. 390, no. 6, pp. 1150–1170, 2011. **(Cited at p. 15.)**
- [64] J. T. Vogelstein, W. G. Roncal, R. J. Vogelstein, and C. E. Priebe, “Graph classification using signal-subgraphs: Applications in statistical connectomics,” *IEEE transactions on pattern analysis and machine intelligence*, vol. 35, no. 7, pp. 1539–1551, 2012. **(Cited at p. 15.)**
- [65] W. Jin, L. Zhao, S. Zhang, Y. Liu, J. Tang, and N. Shah, “Graph condensation for graph neural networks,” *arXiv preprint arXiv:2110.07580*, 2021. **(Cited at p. 16.)**
- [66] Z. Yang, W. Cohen, and R. Salakhudinov, “Revisiting semi-supervised learning with graph embeddings,” in *International conference on machine learning*, pp. 40–48, PMLR, 2016. **(Cited at p. 16.)**
- [67] O. Shchur, M. Mumme, A. Bojchevski, and S. Günnemann, “Pitfalls of graph neural network evaluation,” *arXiv preprint arXiv:1811.05868*, 2018. **(Cited at p. 16.)**
- [68] X. Fu, J. Zhang, Z. Meng, and I. King, “Magnn: Metapath aggregated graph neural network for heterogeneous graph embedding,” in *Proceedings of The Web Conference 2020*, pp. 2331–2341, 2020. **(Cited at p. 16.)**
- [69] F. Yang, W. Wang, F. Wang, Y. Fang, D. Tang, J. Huang, H. Lu, and J. Yao, “scbert as a large-scale pretrained deep language model for cell type annotation of single-cell rna-seq data,” *Nature Machine Intelligence*, vol. 4, no. 10, pp. 852–866, 2022. **(Cited at p. 16.)**
- [70] M. Defferrard, L. Martin, R. Pena, and N. Perraudin, “Pygsp: Graph signal processing in python.” **(Cited at p. 16.)**
- [71] D. J. Watts and S. H. Strogatz, “Collective dynamics of ‘small-world’ networks,” *nature*, vol. 393, no. 6684, pp. 440–442, 1998. **(Cited at p. 16.)**

Appendix

565 A Degree-Corrected Stochastic Block Model(DC-SBM)

The DC-SBM is one of the most commonly used models for networks with communities and postulates that, given node labels $\mathbf{c} = c_1, \dots, c_n$, the edge variables A'_{ij} s are generated via the formula

$$E[A_{ij}] = \theta_i \theta_j P_{c_i} P_{c_j}$$

566 , where θ_i is a "degree parameter" associated with node i , reflecting its individual propensity to
 567 form ties, and P is a $K \times K$ symmetric matrix containing the between/withincommunity edge
 568 probabilities and $P_{c_i} P_{c_j}$ denotes the edge probabilities between community c_i and c_j .
 569 For DC-SBM model [50] assumed P_n on n nodes with k classes, each node v_i is given a label/degree
 570 pair (c_i, θ_i) , drawn from a discrete joint distribution $\Pi_{K \times m}$ which is fixed and does not depend on n .
 571 This implies that each θ_i is one of a fixed set of values $0 \leq x_1 \leq \dots \leq x_m$. To facilitate analysis of
 572 asymptotic graph sparsity, we parameterize the edge probability matrix P as $P_n = \rho_n P$ where P is
 573 independent of n , and $\rho_n = \lambda_n/n$ where λ_n is the average degree of the network.
 574

575 B Neighbourhood Preservation

576 **Theorem 2. Neighborhood Preservation.** Let $\mathcal{N}_k(\mathcal{E}_i)$ denote the neighborhood of incoming nodes
 577 $\mathcal{E}_i$ for the i^{th} community. With partition matrix $\mathcal{P}_i$ and $\mathcal{M}_{gl}(S_i, X_\tau^i) = \mathcal{G}_\tau^c(V_\tau^c, A_\tau^c)$ we identify the
 578 supernodes connected to incoming nodes $\mathcal{E}_i$ and subsequently select nodes within those supernodes;
 579 this subset of nodes is denoted by $\omega_{V_\tau^i}$. Formally,

$$\omega_{V_\tau^i} = \bigcup_{v \in \mathcal{E}_i} \left\{ \bigcup_{s \in S_i} \{\pi^{-1}(s) | A_\tau^c(v, s) \neq 0\} \right\}$$

580 Then, with probability $\Pi_{\{c \in \phi\}} p(c)$, it holds that $\mathcal{N}_k(\mathcal{E}_i) \subseteq \omega_{V_\tau^i}$ where

$$p(c) \leq 1 - \frac{2}{\sqrt{2\pi}} \frac{c}{r} \left[1 - e^{-r^2/(2c^2)} \right],$$

581 and ϕ is a set containing all pairwise distance values ($c = \|v - u\|$) between every node $v \in \mathcal{E}_i$ and the nodes
 582 $u \in \omega_{V_\tau^i}$. Here, $\pi^{-1}(s)$ denotes the set of nodes mapped to supernode s , r is the bin-width hyperparameter of
 583 $\mathcal{M}_{coar}$.

584 **Proof:** The probability that LSH random projection [32, 62] preserves the distance between two
 585 nodes v and u i.e., $d(u, v) = c$, is given by:

$$p(c) = \int_0^r \frac{1}{c} f_2\left(\frac{t}{c}\right) \left(1 - \frac{t}{r}\right) dt,$$

586 where $f_2(x) = \frac{2}{\sqrt{2\pi}} e^{-x^2/2}$ represents the Gaussian kernel when the projection matrix is randomly
 587 sampled from p -stable($p = 2$) distribution [62].

588 The probability $p(c)$ can be decomposed into two terms:

$$p(c) = S_1(c) - S_2(c),$$

589 $S_1(c)$ and $S_2(c)$ are defined as follows:

$$S_1(c) = \frac{2}{\sqrt{2\pi}} \int_0^r e^{-(t/c)^2/2} dt \leq 1,$$

590

$$S_2(c) = \frac{2}{\sqrt{2\pi}} \int_0^r e^{-(t/c)^2/2} \frac{t}{r} dt.$$

591

$$S_2(c) = \frac{2}{\sqrt{2\pi}} \cdot \frac{c}{r} \int_0^r e^{-(t/c)^2/2} \frac{t}{c^2} dt$$

592 Expanding $S_2(c)$:

$$S_2(c) = \frac{2}{\sqrt{2\pi}} \cdot \frac{c}{r} \int_0^{r^2/(2c^2)} e^{-y} dy$$

593

$$S_2(c) = \frac{2}{\sqrt{2\pi}} \cdot \frac{c}{r} \left[1 - e^{-r^2/(2c^2)} \right]$$

594 Thus, the probability $p(c)$ can be bounded as:

$$p(c) \leq 1 - \frac{2}{\sqrt{2\pi}} \frac{c}{r} \left[1 - e^{-r^2/(2c^2)} \right].$$

595 Now, let ϕ be the set of all pairwise distances $d(u, v)$, where $v \in \mathcal{E}_i$ and node $\omega_{V_\tau^i}$. The probability
 596 that all nodes in $\mathcal{N}_k(\mathcal{E}_i)$ are preserved within $\omega_{V_\tau^i}$, requires that all distances $c \in \phi$ are also preserved.
 597 The probability is then given by:

$$\prod_{c \in \phi} p(c).$$

$$\prod_{c \in \phi} p(c) \leq \prod_{c \in \phi} \left(1 - \frac{2}{\sqrt{2\pi}} \frac{c}{r} \left[1 - e^{-r^2/(2c^2)} \right] \right).$$

598 C Continual Learning and Dynamic Graph Learning

599 In this subsection, we highlight the key distinctions between Graph Structure Learning (GSL) and
 600 related fields to justify our specific selection of related works in Section 2.2. GSL is often confused
 601 with topics such as Continual Learning (CL) and Dynamic Graph Learning (DGL).

602 CL [39–41] addresses the issue of catastrophic forgetting, where a model’s performance on previously
 603 learned tasks degrades significantly after training on new tasks. In CL, the model has access only to
 604 the current task’s data and cannot utilize data from prior tasks. Conversely, DGL [42–44] focuses
 605 on capturing the evolving structure of graphs and maintaining updated graph representations, with
 606 access to all prior information.

607 While both CL and DGL aim to *enhance model adaptability* to dynamic data, GSL is primarily
 608 concerned with generating *high-quality graph structures* that can be leveraged for downstream tasks
 609 such as node classification [51], link prediction [63], and graph classification [64]. Moreover, in
 610 CL and DGL, different tasks typically involve distinct data distributions, whereas GSL assumes a
 611 consistent data distribution throughout.

612 D Related Work

613 Table 7 presents the formulations and associated time complexities of various unsupervised Graph
 614 Structure Learning methods.

Table 7: Unsupervised Graph Structure Learning Methods

Method	Time Complexity	Formulation
<i>GLasso</i>	$O(N^3)$	$\max_{\Theta} \log \det \Theta - \text{tr}(\hat{\Sigma}\Theta) - \rho \ \Theta\ _1$
<i>log-model</i>	$O(N^2)$	$\min_{W \in \mathcal{W}} \ W \circ Z\ _{1,1} - \alpha \mathbf{1}^T \log(W\mathbf{1}) + \frac{\beta}{2} \ W\ _F^2$
<i>l2-model</i>	$O(N^2)$	$\min_{W \in \mathcal{W}} \ W \circ Z\ _{1,1} + \alpha \ W\mathbf{1}\ ^2 + \alpha \ W\ _F^2 + \mathbf{1}\{\ W\ _{1,1} = n\}$
<i>large-model</i>	$O(N \log(N))$	$\min_{W \in \tilde{\mathcal{W}}} \ W \circ Z\ _{1,1} - \alpha \mathbf{1}^T \log(W\mathbf{1}) + \frac{\beta}{2} \ W\ _F^2$

615 E Run Time Analysis

616 In the context of clustering module, $k - NN$ is the fastest algorithm, while Spectral Clustering
 617 is the slowest. Suppose we aim to learn the structure of a graph with N nodes. The clustering
 618 module, however, is only applied to a randomly sampled, smaller, static subgraph with k nodes, where
 619 $k \ll N$. In the worst-case scenario, spectral clustering requires $\mathcal{O}(k^3)$ time, whereas in the best case,
 620 $k - NN$ requires $\mathcal{O}(k^2)$ time. For coarsening module, LSH-based coarsening framework [30], has
 621 the best time complexity of $\mathcal{O}(\frac{k_\tau}{c})$ while FGC denotes the worst case with a time-complexity of
 622 $\mathcal{O}((\frac{k_\tau}{c})^2 \|S_\tau^i\|)$ where c is the number of communities detected by clustering module $\mathcal{M}_{\text{clust}}$, $\|S_\tau^i\|$
 623 is the number of coarsened node in the relevant community at τ timestamp and k_τ denotes number
 624 of nodes at τ timestamp. For learning module, $A - NN$ is the most efficient algorithm with time

complexity as $\mathcal{O}(N \log N)$, while *GLasso* has the worst computational cost of $\mathcal{O}(N^3)$. So, the effective time complexity of GraphFLEx is upper bounded by $\mathcal{O}(k^3 + (\frac{k_\tau}{c})^2 \|S_\tau^i\| + \alpha^3)$ and lower bounded by $\mathcal{O}(k^2 + \frac{k_\tau}{c} + \alpha \log \alpha)$ where $\alpha = \|S_\tau^i\| + \|\mathcal{E}_\tau^i\|$. GraphFLEx’s efficiency in term of computational time is evident in Figure 1 and further quantified in Table 2.

Out of the three modules of GraphFLEx first module($\mathcal{M}_{\text{clust}}$) is trained once, and hence its run time is always bounded; computational time for second module($\mathcal{M}_{\text{coar}}$) can also be controlled because some of the methods either needs training once [65] or have linear time complexity [30]. Consequently, both the clustering and coarsening modules contribute linearly to the overall time complexity, denoted as $\mathcal{O}(N)$. Thus, the effective time complexity of GraphFLEx is given by $\mathcal{O}(N + \mathcal{O}(\mathcal{M}_{gl}(\|S_i, X_\tau^i\|)))$. The overall complexity scales either linearly or sub-linearly, depending on α and the $\mathcal{M}_{gl}$ module. For instance, when $\mathcal{M}_{gl}$ is A-NN the complexity remains linear, if $\alpha \log(\alpha) \approx N$, whereas for *GLasso*, a linear behavior is observed when $\alpha^3 \approx N$.

F Datasets

Datasets used in our experiments vary in size, with nodes ranging from 1k to 60k. Table 8 lists all the datasets we used in our work. We evaluate our proposed framework *GraphFlex* on real-world datasets *Cora*, *Citeseer*, *Pubmed* [66], *CS*, *Physics* [67], *DBLP* [68], all of which include graph structures. These datasets allow us to compare the learned structures with the originals. Additionally, we utilize single-cell RNA pancreas datasets [69], including Baron, Muraro, Segerstolpe, and Xin, where the graph structure is missing. The Baron dataset was downloaded from the Gene Expression Omnibus (GEO) (accession no. GSE84133). The Muraro dataset was downloaded from GEO (accession no. GSE85241). The Segerstolpe dataset was accessed from ArrayExpress (accession no. E-MTAB-5061). The Xin dataset was downloaded from GEO (accession no. GSE81608). We simulate the expanding graph scenario by splitting the original dataset across different T timestamps. We assumed 50% of the nodes were static, with the remaining nodes arriving as incoming nodes at different timestamps.

Synthetic datasets: Different data generation techniques validate that our results are generalized to different settings. Please refer to Table 8 for more details about the number of nodes, edges, features, and classes, *Syn* denotes the type of synthetic datasets. Figure 6 shows graphs generated using different methods. We have employed three different ways to generate synthetic datasets which are mentioned below:

- **PyGSP(PyGsp):** We used synthetic graphs created by PyGSP [70] library. PyG-G and PyG-S denotes grid and sensor graphs from PyGSP.
- **Watts–Strogatz’s small world(SW):** [71] proposed a generation model that produces graphs with small-world properties, including short average path lengths and high clustering.
- **Heterophily(HE):** We propose a method for creating synthetic datasets to explore graph behavior across a heterophily spectrum by manipulating heterophilic factor α , and classes. α is determined by dividing the number of edges connecting nodes from different classes by the total number of edges in the graph.

Visualization Datasets: To evaluate, the learned graph structure, we have also included three datasets: (i) MNIST [59], consisting of handwritten digit images; (ii) Pre-trained GloVe embeddings [60] of English words; and (iii) Zachary’s karate club network [61].

Large Datasets: To comprehensively evaluate GraphFLEx’s scalability to large-scale graphs, we consider four datasets with a high number of nodes: (a) *Flickr*(89k nodes) [55], (b) *Reddit* (233k nodes) [55], (c) *Ogbn-arxiv* (169k nodes) [46], and (d) *Ogbn-products* (2.4M nodes) [56].

System Specifications: All the experiments conducted for this work were performed on an Intel Xeon W-295 CPU with 64GB of RAM desktop using the Python environment.

Category	Data	Nodes	Edges	Feat.	Class	Type
Original Structure Known	Cora	2,708	5,429	1,433	7	Citation network
	Citeseer	3,327	9,104	3,703	6	Citation network
	DBLP	17,716	52.8k	1,639	4	Research paper
	CS	18,333	163.7k	6,805	15	Co-authorship network
	PubMed	19,717	44.3k	500	3	Citation network
	Physics	34,493	247.9k	8,415	5	Co-authorship network
Original Structure Not Known	Xin	1,449	NA	33,889	4	Human Pancreas
	Baron Mouse	1,886	NA	14,861	13	Mouse Pancreas
	Muraro	2,122	NA	18,915	9	Human Pancreas
	Segerstolpe	2,133	NA	22,757	13	Human Pancreas
	Baron Human	8,569	NA	17,499	14	Human Pancreas
Synthetic	Syn 1	2,000	8,800	150	4	SW
	Syn 2	5,000	22k	150	4	SW
	Syn 3	10,000	44k	150	7	SW
	Syn 4	50,000	220k	150	7	SW
	Syn 5	400	1,520	100	4	PyG-G
	Syn 6	2,500	9,800	100	4	PyG-S
	Syn 7	1,000	9,990	150	4	HE
	Syn 8	2,000	40k	150	4	HE
Visulization Datasets	MNIST	60,000	NA	784	10	Images
	Zachary’s karate	34	156	34	4	Karate club network
	Glove	2,000	NA	50	NA	GloVe embeddings
Large dataset	Flickr	89,250	899,756	500	7	-
	Reddit	232,965	11.60M	602	41	-
	Ogbn-arxiv	169,343	1.16M	128	40	-
	Ogbn-products	2,449,029	61.85M	100	47	-

Table 8: Summary of the datasets.

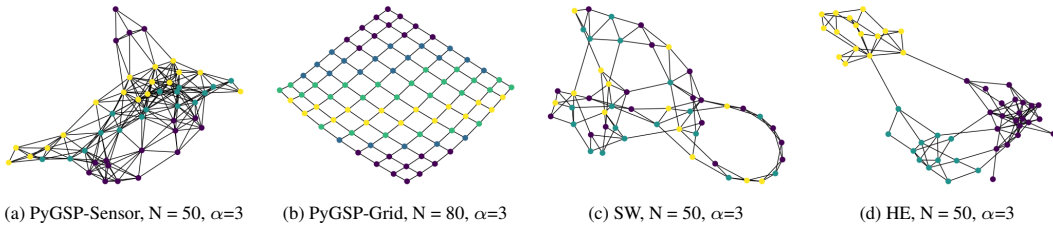


Figure 6: This figure illustrates different types of synthetic graphs generated using i)PyGSP, ii) Watts–Strogatz’s small world(SW), and iii) Heterophily(HE). N denotes the number of nodes, while α denotes the number of classes.

Algorithm 1 GraphFlex: A Unified Structure Learning framework for expanding and Large Scale Graphs

Input: Graph $G_0(X_0, A_0)$, expanding nodes set $\mathcal{E}_1^T = \{\mathcal{E}_\tau(\mathcal{V}_\tau, \mathcal{X}_\tau)\}_{\tau=1}^T$

Parameter: GClust, GCoar, GL $\leftarrow$ Clustering, Coarsening and Learning Module

Output: Graph $G_T(X_T, A_T)$

```

1: Train clustering module  $train(\mathcal{M}_{clust}, \text{GClust}, G_0)$ 
2: for each  $E_t(V_t, X_t)$  in  $\mathcal{E}_1^T$  do
3:    $C_t = infer(\mathcal{M}_{clust}, X_t)$ ,  $C_t \in \mathbb{R}^{N_t}$  denotes the communities of  $N_t$  nodes at time  $t$ .
4:    $I_t = unique(C_t)$ .
5:   for each  $I_t^i$  in  $I_t$  do
6:      $G_{t-1}^i = subgraph(G_{t-1}, I_t^i)$ 
7:      $\{S_{t-1}^i, P_{t-1}^i\} = \mathcal{M}_{coar}(G_{t-1}^i)$ ,  $S_{t-1}^i \in \mathbb{R}^{k \times d}$  are features of  $k$  supernodes,  $P_{t-1}^i \in \mathbb{R}^{k \times N_t^i}$ 
       is the partition matrix.
8:      $G_{t-1}^i(S_{t-1}^i, A_{t-1}^i) = \mathcal{M}_{gl}(S_{t-1}^i, X_t^i)$ ,  $G_{t-1}^i$  is the learned graph on super-nodes  $S_{t-1}^i$ 
       and new node  $X_t^i$ .
9:      $\omega_t^i \leftarrow []$ 
10:    for  $x \in X_t^i$  do
11:       $\omega_t^i.append(x)$ 
12:       $n_p = \{n \mid A_{t-1}^i[n] > 0\}$ 
13:       $\omega_t^i.append(n_p)$ 
14:    end for
15:     $G_{t-1} = update(G_{t-1}, \mathcal{M}_{gl}(\omega_t^i))$ 
16:  end for
17:   $G_t = G_{t-1}$ 
18: end for
19: return  $G_T(X_T, A_T)$ 

```

673 H Other GNN models

674 We used four GNN models, namely GCN, GraphSage, GIN, and GAT. Table 9 contains parameter
675 details we used to train GraphFlex. We have used these parameters across all methods.

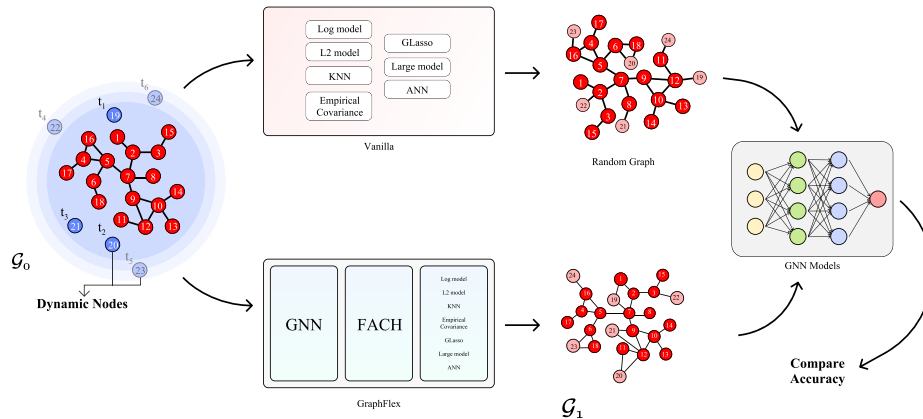


Figure 7: GNN training pipeline.

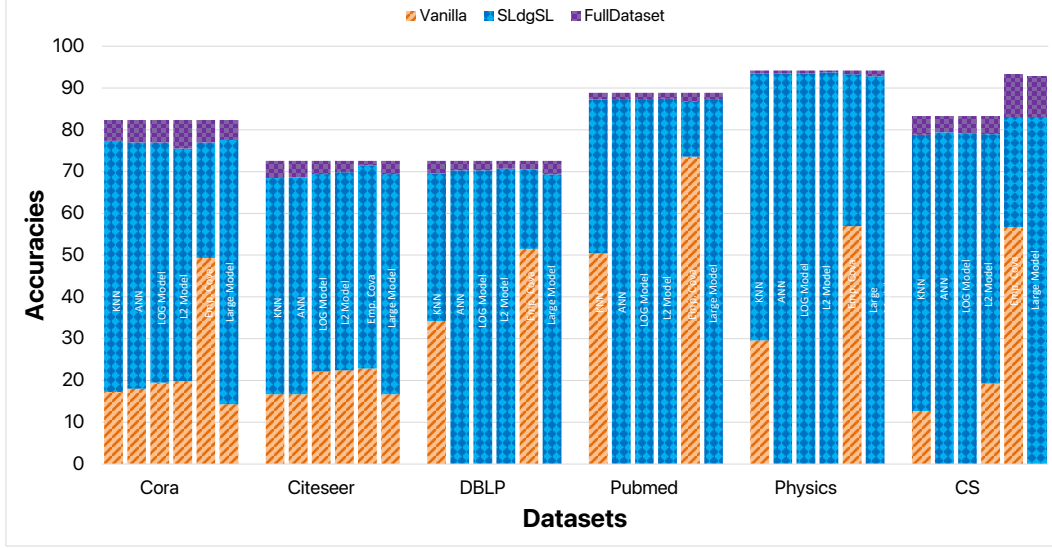


Figure 8: GraphSage accuracies when structure is learned or given with 3 different scenarios (Vanilla, GraphFlex, Original) across different datasets, highlighting performance with 30% node growth over 25 timestamps.

Figure 7 illustrates the pipeline for training our GNN models. Graph structures were learned using both existing methods and GraphFlex, and GNN models were subsequently trained on both structures. Results across all datasets are presented in Table 10 and Table 3.

Table 9: GNN model parameters.

Model	Hidden Layers	L.R	Decay	Epoch
GCN	{64, 64}	0.003	0.0005	500
GraphSage	{64, 64}	0.003	0.0005	500
GIN	{64, 64}	0.003	0.0005	500
GAT	{64, 64}	0.003	0.0005	500

We randomly split data in 60%, 20%, 20% for training-validation-test. The results for these models on synthetic datasets are presented in Table 10.

Figure 7 illustrates the pipeline for training our GNN models. Graph structures were learned using both existing methods and GraphFlex, and GNN models were subsequently trained on both structures.

I Computational Efficiency

Table 11 illustrates the remaining computational time for learning graph structures using GraphFlex with existing Vanilla methods on Synthetic datasets. While traditional methods may be efficient for small graphs, GraphFlex scales significantly better, excelling on large datasets like *Pubmed* and *Syn 5*, where most methods fail.

J Visualization of Growing graphs

This section helps us visualize the phases of our growing graphs. We have generated a synthetic graph of 60 nodes using PyGSP-Sensor and HE methods mentioned in Appendix F. We then added 40 new nodes denoted using black color in these existing graphs at four different timestamps. Figure 9 and Figure 10 shows the learned graph structure after each timestamp for two different Synthetic graphs.

K Clustering Quality

Figure 12 shows the PHATE [58] visualization of clusters learned using GraphFlex’s clustering module $\mathcal{M}_{clust}$ for 6 single-cell RNA datasets, namely *Xin*, *MNIST*, *Baron – Human*, *Muraro*, *BaronMouse*, and *Segerstolpe* datasets.

Table 10: Node classification accuracies on different GNN models using GraphFLEX (GFlex) with existing Vanilla (Van.) methods. The experimental setup involves treating 70% of the data as static, while the remaining 30% of nodes are treated as new nodes coming in 25 different timestamps. The best and the second-best accuracies in each row are highlighted by dark and lighter shades of **Green**, respectively. GraphFLEX’s structure beats all of the vanilla structures for every dataset. OOM and OOT denotes out-of-memory and out-of-time respectively.

Dataset	Model	ANN		KNN		log-model		l2-model		COVAR		large-model		Base Struc.
		Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	
Cora	GAT	18.73	73.84	20.96	73.65	16.14	72.36	18.74	73.10	49.72	77.55	14.28	76.43	79.77
	SAGE	17.25	77.37	18.00	76.99	19.48	77.40	19.85	75.51	49.35	76.99	14.28	77.55	82.37
	GCN	17.99	78.11	17.81	77.92	18.55	77.74	20.41	79.22	47.31	80.52	14.28	79.03	84.60
	GIN	16.69	76.44	18.74	80.52	17.44	76.25	19.29	76.62	48.79	78.85	14.28	76.06	81.63
Citeseer	GAT	16.51	61.82	25.00	62.27	19.24	64.70	18.18	63.48	20.91	62.73	16.67	62.27	66.42
	SAGE	16.66	68.48	16.67	68.64	22.12	69.39	22.42	69.85	22.88	71.52	16.67	69.39	72.57
	GCN	28.18	60.00	16.67	61.97	20.45	65.45	19.70	64.24	21.06	64.70	16.67	63.18	68.03
	GIN	16.66	64.39	16.67	63.94	20.15	59.85	18.64	63.64	22.12	60.30	16.67	61.81	67.38
Syn 4	GAT	29.55	92.07	OOM	90.86	OOT	91.64	OOT	91.64	35.79	92.52	OOT	93.74	89.49
	SAGE	26.75	87.89	OOM	91.05	OOT	86.64	OOT	86.64	32.92	90.44	OOT	86.01	90.03
	GCN	28.85	51.97	OOM	19.58	OOT	18.29	OOT	18.92	33.80	26.60	OOT	36.85	21.43
	GIN	28.50	65.61	OOM	31.06	OOT	26.51	OOT	26.56	34.03	46.40	OOT	47.10	29.35
Syn 6	GAT	44.00	86.80	43.60	86.60	30.00	78.75	55.40	92.80	36.20	93.60	31.80	92.80	97.20
	SAGE	41.00	93.80	41.40	93.60	33.75	88.75	57.60	94.00	35.20	94.80	28.20	95.60	97.40
	GCN	43.60	88.80	42.20	87.40	26.25	81.25	55.60	92.40	31.40	94.40	25.20	94.00	99.40
	GIN	39.60	89.00	40.40	86.60	21.25	82.50	55.20	91.80	30.00	94.60	30.40	92.00	98.80
Syn 8	GAT	29.55	99.75	33.75	88.75	88.25	99.25	88.25	99.25	26.00	85.50	94.00	96.00	98.50
	SAGE	26.75	100.0	32.50	100.0	88.75	99.50	88.75	99.50	26.75	100.0	92.50	100.0	100.0
	GCN	28.85	98.75	31.75	99.75	88.75	99.00	88.75	99.00	28.50	99.25	95.00	100.0	100.0
	GIN	28.50	50.00	30.50	91.00	82.25	91.50	82.25	91.50	27.25	81.75	91.75	92.25	78.25

Table 11: Computational time for learning graph structures using GraphFLEX (GFlex) with existing methods (Vanilla referred to as Van.). The experimental setup involves treating 50% of the data as static, while the remaining 50% of nodes are treated as incoming nodes arriving in 25 different timestamps. The best times are highlighted by color **Green**. OOM and OOT denote out-of-memory and out-of-time, respectively.

Data	ANN		KNN		log-model		l2-model		COVAR		large-model	
	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex	Van.	GFlex
Syn 1	19.4	9.8	2.5	10.5	2418	56.4	37.2	8.8	3.5	8.3	205	9.4
Syn 2	47.3	16.9	6.6	18.3	14000	144	214	22.6	20.3	18.6	1259	16.4
Syn 5	5.1	11.5	0.8	7.3	57.4	28	1.1	5.8	0.2	4.8	3.2	5.3
Syn 6	16.6	9.9	2.8	11.4	1766	96.3	193	101	5.3	8.9	324	9.6
Syn 7	10.6	7.4	1.4	8.9	704	85.2	10.3	7.9	0.9	6.4	36.5	8.2
Syn 8	19.6	11.2	2.5	11.7	2416	457	37.2	17.0	3.4	10.9	204	11.7

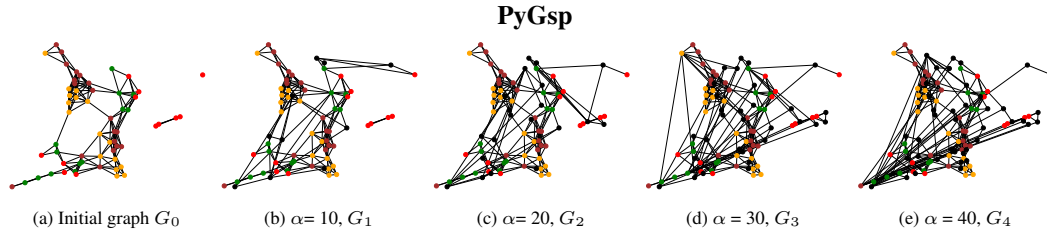


Figure 9: This figure illustrates the growing structure learned using GraphFlex for dynamic nodes. New nodes are denoted using black color, and α denotes number of new nodes. *PyGsp* denotes type synthetic graph.

L Ablation Study

In this section, we present an ablation study to analyze the role of individual modules within GraphFLEX and their influence on the final graph structure. Specifically, we focus on two aspects: (i) the significance of the clustering module, and (ii) the effect of varying module configurations on the learned graph topology.

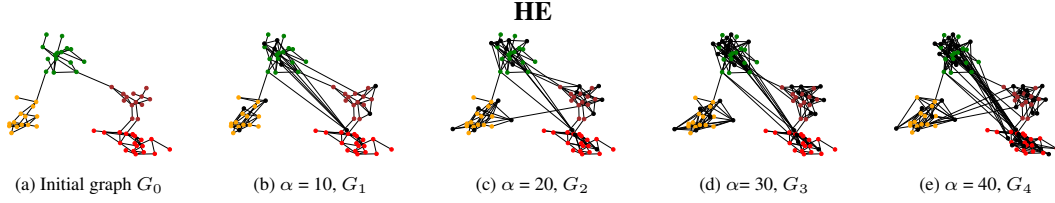


Figure 10: This figure illustrates the growing structure learned using GraphFlex for dynamic nodes. New nodes are denoted using black color, and α denotes the number of new nodes. *HE* denotes the type of synthetic graph.

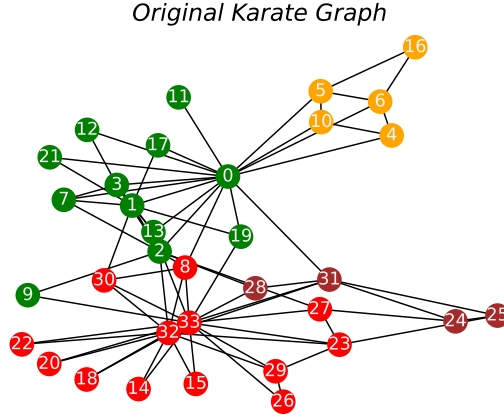


Figure 11: Original Karate Graph

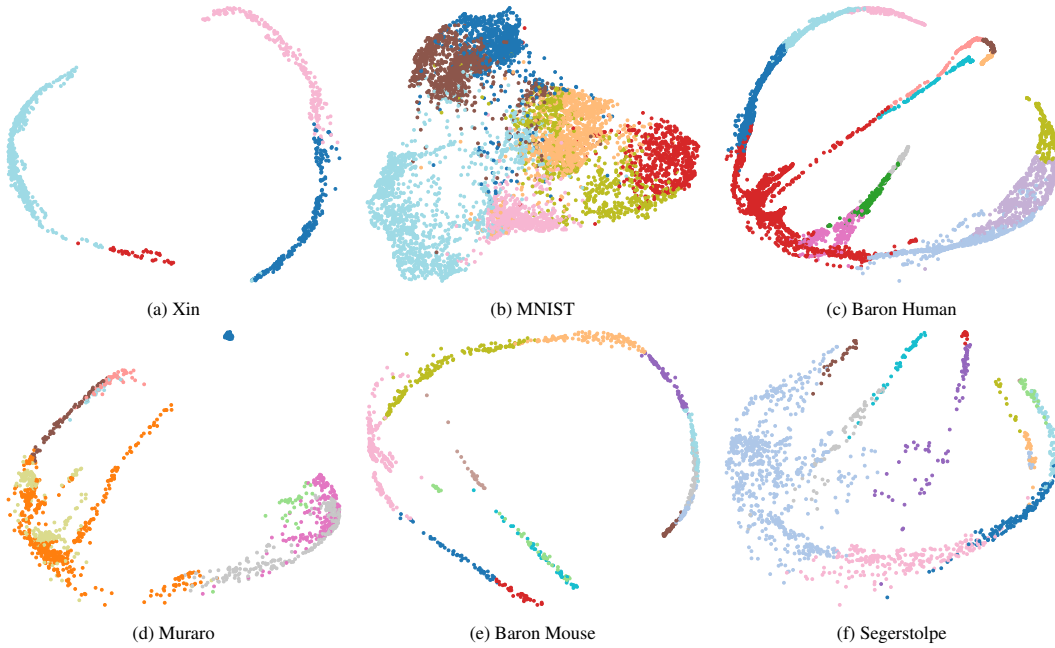


Figure 12: PHATE visualization of clusters learnt using GraphFlex clustering module for scRNA-seq datasets.

L.1 Clustering Module Evaluation

To evaluate the effectiveness of the clustering module, we compute standard metrics such as Normalized Mutual Information (NMI), Conductance (C), and Modularity (Q) across various datasets (see Table 6 in Section 4.5). These metrics collectively validate the quality of the discovered clusters, thereby justifying the use of a clustering module as a foundational step in GraphFLEX. Since clustering in GraphFLEX is applied only once on a randomly sampled small set of nodes, selecting the right

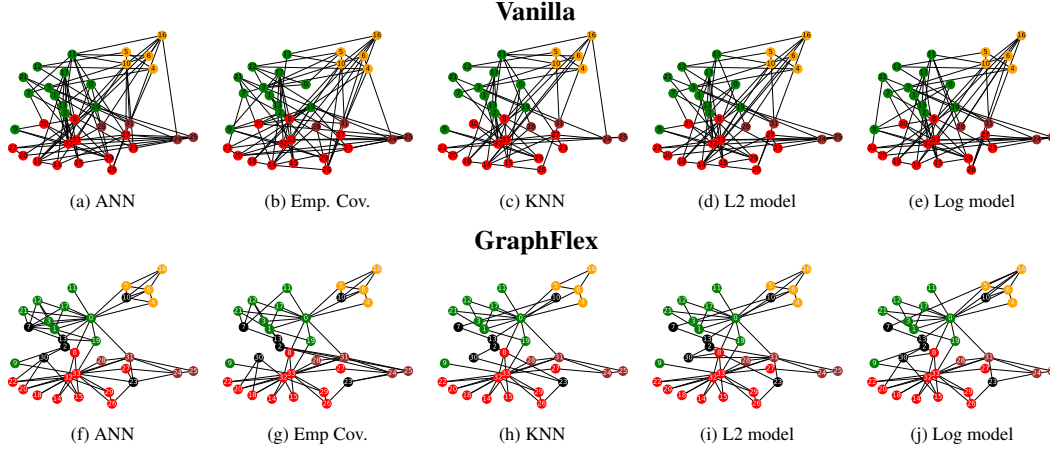


Figure 13: This figure compares the structures learned on Zachary’s karate dataset when existing methods are employed with GraphFlex and when existing methods are used individually. We consider six nodes, denoted in black, as dynamic nodes.

708 method can be considered as part of hyperparameter tuning, where these clustering measures can
 709 guide the optimal choice based on dataset characteristics.

710 L.2 Impact of Module Choices on Learned Graph Structure

711 This section involves a comparison of the graph structure learned from GraphFlex with existing
 712 methods. Six nodes were randomly selected and considered as new nodes. Figure 13 visually depicts
 713 the structures learned using GraphFlex compared to other methods. It is evident from the figure that
 714 the structure known with GraphFlex closely resembles the original graph structure. Figure 11 shows
 715 the original structure of Zachary’s karate club network [61]. We assumed six random nodes to be
 716 dynamic nodes, and the structure learned using GraphFlex compared to existing methods is shown in
 717 Figure 13.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: Yes, all the claims are reflected in paper. See Section 4 and Appendix.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: See Section 5.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: See Appendix.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.

777 • Theorems and Lemmas that the proof relies upon should be properly referenced.

778 **4. Experimental result reproducibility**

779 Question: Does the paper fully disclose all the information needed to reproduce the main experi-

780 mental results of the paper to the extent that it affects the main claims and/or conclusions of the

781 paper (regardless of whether the code and data are provided or not)?

782 Answer: [Yes]

783 Justification: See Section 4 and Appendix.

784 Guidelines:

785 • The answer NA means that the paper does not include experiments.

786 • If the paper includes experiments, a No answer to this question will not be perceived well by the

787 reviewers: Making the paper reproducible is important, regardless of whether the code and data

788 are provided or not.

789 • If the contribution is a dataset and/or model, the authors should describe the steps taken to make

790 their results reproducible or verifiable.

791 • Depending on the contribution, reproducibility can be accomplished in various ways. For

792 example, if the contribution is a novel architecture, describing the architecture fully might

793 suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary

794 to either make it possible for others to replicate the model with the same dataset, or provide

795 access to the model. In general, releasing code and data is often one good way to accomplish

796 this, but reproducibility can also be provided via detailed instructions for how to replicate the

797 results, access to a hosted model (e.g., in the case of a large language model), releasing of a

798 model checkpoint, or other means that are appropriate to the research performed.

799 • While NeurIPS does not require releasing code, the conference does require all submissions

800 to provide some reasonable avenue for reproducibility, which may depend on the nature of the

801 contribution. For example

802 (a) If the contribution is primarily a new algorithm, the paper should make it clear how to

803 reproduce that algorithm.

804 (b) If the contribution is primarily a new model architecture, the paper should describe the

805 architecture clearly and fully.

806 (c) If the contribution is a new model (e.g., a large language model), then there should either

807 be a way to access this model for reproducing the results or a way to reproduce the model

808 (e.g., with an open-source dataset or instructions for how to construct the dataset).

809 (d) We recognize that reproducibility may be tricky in some cases, in which case authors are

810 welcome to describe the particular way they provide for reproducibility. In the case of

811 closed-source models, it may be that access to the model is limited in some way (e.g.,

812 to registered users), but it should be possible for other researchers to have some path to

813 reproducing or verifying the results.

814 **5. Open access to data and code**

815 Question: Does the paper provide open access to the data and code, with sufficient instructions to

816 faithfully reproduce the main experimental results, as described in supplemental material?

817 Answer: [Yes]

818 Justification: All datasets used are publicly available. See Abstract for codebase.

819 Guidelines:

820 • The answer NA means that paper does not include experiments requiring code.

821 • Please see the NeurIPS code and data submission guidelines ([https://nips.cc/public/](https://nips.cc/public/guides/CodeSubmissionPolicy)

822 [guides/CodeSubmissionPolicy](https://nips.cc/public/guides/CodeSubmissionPolicy)) for more details.

823 • While we encourage the release of code and data, we understand that this might not be possible,

824 so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless

825 this is central to the contribution (e.g., for a new open-source benchmark).

826 • The instructions should contain the exact command and environment needed to run to reproduce

827 the results. See the NeurIPS code and data submission guidelines ([https://nips.cc/public/](https://nips.cc/public/guides/CodeSubmissionPolicy)

828 [guides/CodeSubmissionPolicy](https://nips.cc/public/guides/CodeSubmissionPolicy)) for more details.

829 • The authors should provide instructions on data access and preparation, including how to access

830 the raw data, preprocessed data, intermediate data, and generated data, etc.

831 • The authors should provide scripts to reproduce all experimental results for the new proposed

832 method and baselines. If only a subset of experiments are reproducible, they should state which

833 ones are omitted from the script and why.

834 • At submission time, to preserve anonymity, the authors should release anonymized versions (if

835 applicable).

836 • Providing as much information as possible in supplemental material (appended to the paper) is
837 recommended, but including URLs to data and code is permitted.

838 **6. Experimental setting/details**

839 Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters,
840 how they were chosen, type of optimizer, etc.) necessary to understand the results?

841 Answer: [Yes]

842 Justification: See Section 4 and Appendix.

843 Guidelines:

- 844 • The answer NA means that the paper does not include experiments.
- 845 • The experimental setting should be presented in the core of the paper to a level of detail that is
846 necessary to appreciate the results and make sense of them.
- 847 • The full details can be provided either with the code, in appendix, or as supplemental material.

848 **7. Experiment statistical significance**

849 Question: Does the paper report error bars suitably and correctly defined or other appropriate
850 information about the statistical significance of the experiments?

851 Answer: [Yes]

852 Justification: See Section 4 and Appendix.

853 Guidelines:

- 854 • The answer NA means that the paper does not include experiments.
- 855 • The authors should answer "Yes" if the results are accompanied by error bars, confidence
856 intervals, or statistical significance tests, at least for the experiments that support the main claims
857 of the paper.
- 858 • The factors of variability that the error bars are capturing should be clearly stated (for example,
859 train/test split, initialization, random drawing of some parameter, or overall run with given
860 experimental conditions).
- 861 • The method for calculating the error bars should be explained (closed form formula, call to a
862 library function, bootstrap, etc.)
- 863 • The assumptions made should be given (e.g., Normally distributed errors).
- 864 • It should be clear whether the error bar is the standard deviation or the standard error of the
865 mean.
- 866 • It is OK to report 1-sigma error bars, but one should state it. The authors should preferably
867 report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of
868 errors is not verified.
- 869 • For asymmetric distributions, the authors should be careful not to show in tables or figures
870 symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- 871 • If error bars are reported in tables or plots, The authors should explain in the text how they were
872 calculated and reference the corresponding figures or tables in the text.

873 **8. Experiments compute resources**

874 Question: For each experiment, does the paper provide sufficient information on the computer
875 resources (type of compute workers, memory, time of execution) needed to reproduce the experi-
876 ments?

877 Answer: [Yes]

878 Justification: See Appendix.

879 Guidelines:

- 880 • The answer NA means that the paper does not include experiments.
- 881 • The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud
882 provider, including relevant memory and storage.
- 883 • The paper should provide the amount of compute required for each of the individual experimental
884 runs as well as estimate the total compute.
- 885 • The paper should disclose whether the full research project required more compute than the
886 experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it
887 into the paper).

888 **9. Code of ethics**

889 Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS
890 Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

891 Answer: [Yes]

892 Justification: Research conducted in the paper conform, in every respect, with the NeurIPS Code
893 of Ethics.

894 Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. **Broader impacts**

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There is no societal impact of the work performed.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. **Safeguards**

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. **Licenses for existing assets**

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: Assets are properly credited and publicly available.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

953 • For scraped data from a particular source (e.g., website), the copyright and terms of service of
954 that source should be provided.

955 • If assets are released, the license, copyright information, and terms of use in the package should
956 be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for
957 some datasets. Their licensing guide can help determine the license of a dataset.

958 • For existing datasets that are re-packaged, both the original license and the license of the derived
959 asset (if it has changed) should be provided.

960 • If this information is not available online, the authors are encouraged to reach out to the asset’s
961 creators.

962 **13. New assets**

963 Question: Are new assets introduced in the paper well documented and is the documentation
964 provided alongside the assets?

965 Answer: [NA]

966 Justification: The paper does not release new assets.

967 Guidelines:

968 • The answer NA means that the paper does not release new assets.

969 • Researchers should communicate the details of the dataset/code/model as part of their sub-
970 missions via structured templates. This includes details about training, license, limitations,
971 etc.

972 • The paper should discuss whether and how consent was obtained from people whose asset is
973 used.

974 • At submission time, remember to anonymize your assets (if applicable). You can either create
975 an anonymized URL or include an anonymized zip file.

976 **14. Crowdsourcing and research with human subjects**

977 Question: For crowdsourcing experiments and research with human subjects, does the paper
978 include the full text of instructions given to participants and screenshots, if applicable, as well as
979 details about compensation (if any)?

980 Answer: [NA]

981 Justification: The paper does not involve crowdsourcing nor research with human subjects.

982 Guidelines:

983 • The answer NA means that the paper does not involve crowdsourcing nor research with human
984 subjects.

985 • Including this information in the supplemental material is fine, but if the main contribution of
986 the paper involves human subjects, then as much detail as possible should be included in the
987 main paper.

988 • According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other
989 labor should be paid at least the minimum wage in the country of the data collector.

990 **15. Institutional review board (IRB) approvals or equivalent for research with human subjects**

991 Question: Does the paper describe potential risks incurred by study participants, whether such
992 risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals
993 (or an equivalent approval/review based on the requirements of your country or institution) were
994 obtained?

995 Answer: [NA]

996 Justification: The paper does not involve crowdsourcing nor research with human subjects.

997 Guidelines:

998 • The answer NA means that the paper does not involve crowdsourcing nor research with human
999 subjects.

1000 • Depending on the country in which research is conducted, IRB approval (or equivalent) may be
1001 required for any human subjects research. If you obtained IRB approval, you should clearly
1002 state this in the paper.

1003 • We recognize that the procedures for this may vary significantly between institutions and
1004 locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for
1005 their institution.

1006 • For initial submissions, do not include any information that would break anonymity (if applica-
1007 ble), such as the institution conducting the review.

1008 **16. Declaration of LLM usage**

1009 Question: Does the paper describe the usage of LLMs if it is an important, original, or non-
1010 standard component of the core methods in this research? Note that if the LLM is used only for

1011 writing, editing, or formatting purposes and does not impact the core methodology, scientific
1012 rigorousness, or originality of the research, declaration is not required.
1013 Answer: [NA]
1014 Justification: Declaration is not required as LLM is only used for writing, editing, or formatting
1015 purposes.
1016 Guidelines:
1017 • The answer NA means that the core method development in this research does not involve LLMs
1018 as any important, original, or non-standard components.
1019 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what
1020 should or should not be described.