

FLASH: Flexible Learning of Adaptive Sampling from History in Temporal Graph Neural Networks

Or Feldman¹, Krishna Sri Ipsit Mantri², Carola-Bibiane Schönlieb³
Chaim Baskin¹, Moshe Eliasof^{1,3}

¹Ben-Gurion University of the Negev, Israel

²Purdue University, USA

³University of Cambridge, UK

orfel@post.bgu.ac.il, kmantri@uni-bonn.de, cbs31@cam.ac.uk, chaimbaskin@bgu.ac.il
eliasof@bgu.ac.il

Abstract

Aggregating temporal signals from historic interactions is a key step in future link prediction on dynamic graphs. However, incorporating long histories is resource-intensive. Hence, temporal graph neural networks (TGNNs) often rely on historical neighbors sampling heuristics such as uniform sampling or recent neighbors selection. These heuristics are static and fail to adapt to the underlying graph structure. We introduce FLASH, a learnable and graph-adaptive neighborhood selection mechanism that generalizes existing heuristics. FLASH integrates seamlessly into TGNNs and is trained end-to-end using a self-supervised ranking loss. We provide theoretical evidence that commonly used heuristics hinder TGNNs performance, motivating our design. Extensive experiments across multiple benchmarks demonstrate consistent and significant performance improvements for TGNNs equipped with FLASH.

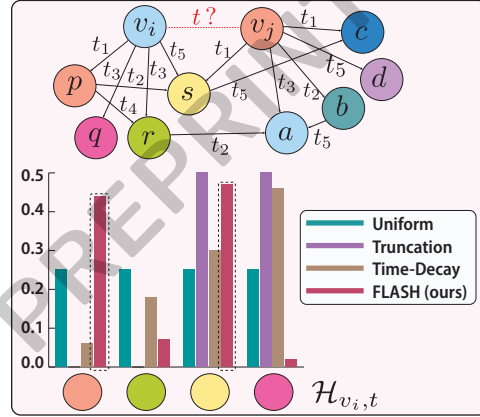


Figure 1: Illustration of different neighborhood selection strategies for predicting a link between v_i and v_j . Circles represent nodes and their colors indicate each node’s feature. The bar chart on the right shows how each strategy scores the neighbors. Static heuristics (truncation or uniform sampling) either discard them or fail to prioritize important neighbors, while FLASH assigns higher scores to these key neighbors.

1 Introduction

Dynamic graphs provide a natural framework for modeling real-world systems where entities and their interactions evolve over time. They underpin a wide range of applications, including social and communication networks [Shetty and Adibi, 2004; Kumar *et al.*, 2019], user-item recommendation systems [Kumar *et al.*, 2019], and financial or knowledge-intensive platforms [Ni *et al.*, 2019; Shamsi *et al.*, 2022]. Predicting future interactions in these settings has emerged as a central learning task, leading to the development of Temporal Graph Neural Networks (TGNNs) – models specifically designed to learn from sequences of timestamped events represented by dynamic graphs.

TGNNs process dynamic graphs by encoding temporal interaction patterns into node representations, enabling them to predict future links. A common challenge in these models is how to efficiently aggregate information from a node’s history of interactions, which can grow unbounded over time. Processing complete histories quickly becomes computationally prohibitive, especially in high-frequency interaction settings. To address this, existing models such as TGN [Rossi

et al., 2020], TGAT [Xu *et al.*, 2020], DyGFormer [Yu *et al.*, 2023], GraphMixer [Cong *et al.*, 2023], FreeDyG [Tian *et al.*, 2023], and others [Zhou *et al.*, 2022; Zou *et al.*, 2024; Gravina *et al.*, 2024; Xu *et al.*, 2024] adopt memory-efficient heuristics. These typically include strategies like uniform sampling, time-decay weighting, or truncating to the k most recent interactions. While effective at reducing computational overhead, these approaches are static and fail to adapt to the local graph structure or task-specific temporal signals. As shown in Figure 1, such heuristics apply uniform or truncated selection schemes that overlook potentially informative neighbors, whereas adaptive strategies can prioritize structurally meaningful interactions.

Static heuristics like uniform sampling [Rossi *et al.*, 2020], truncation [Cong *et al.*, 2023] or hybrid approaches [Luo and Li, 2024] are appealing due to their simplicity, but they treat all interactions as equally informative or rely solely on recency. This ignores the fact that some neighbors may be more relevant than others due to their position in the graph or their interaction patterns. Moreover, the optimal sampled

neighborhood may vary across time, nodes, and tasks, making fixed strategies fundamentally limited. These shortcomings are further amplified in heterogeneous or rapidly evolving graphs, where structural context can shift dramatically over time. This motivates the need for a learnable, structure-aware neighborhood selection mechanism that can adaptively prioritize informative past interactions.

To address these limitations, we propose FLASH – **Flexible Learning of Adaptive Selection from History for Temporal Graph Neural Networks**, a learnable and graph-adaptive neighborhood selection mechanism for TGNNs. FLASH replaces static heuristics with a data-driven approach that learns to prioritize historically informative neighbors based on their structural and temporal context. Crucially, because the true importance of neighbors is not known a priori, our method is trained using a self-supervised ranking objective that encourages selecting neighbors most predictive of future interactions. FLASH is lightweight, general-purpose, and integrates seamlessly into a wide range of existing TGNN architectures, including TGNNs with non-differentiable feature extractors [Yu *et al.*, 2023; Tian *et al.*, 2023]. This allows it to improve predictive performance without requiring architectural changes. Our key contributions are as follows:

- We propose FLASH, a novel graph-adaptive, learnable neighborhood selection mechanism that seamlessly integrates with any TGNN.
- We design a self-supervised training objective based on ranking loss, enabling our method to learn informative neighbor selection without access to ground-truth labels.
- We provide a theoretical analysis showing that FLASH is provably more expressive than the common heuristics of recent neighbors selection and uniform sampling.
- We conduct extensive experiments across multiple dynamic graph benchmarks, demonstrating consistent performance gains across diverse TGNN backbones compared to common neighbor sampling baselines.

2 Related Work

Neighborhood selection in static graphs Existing sampling techniques for large static graphs often rely on substructure sampling (e.g., nodes or edges), as employed by GraphSAGE [Hamilton *et al.*, 2017] and FastGCN [Chen *et al.*, 2018], or utilize random walks, as in PinSage [Ying *et al.*, 2018]. Other methods, such as GraphSAINT [Zeng *et al.*, 2020] and Cluster-GCN [Chiang *et al.*, 2019], are specifically designed to facilitate efficient training on large graphs. However, these approaches typically do not address inference or the temporal nature of dynamic graphs. In many TGNNs, uniform sampling can be viewed as dynamic extension of GraphSAGE-like methods, in which each neighbor in the historical neighborhood is sampled with equal probability. On the other hand, truncation sampling retains only the most recent neighbors, and its stochastic variant of time-weighting, reduces the sampling probability of older neighbors over time, effectively treating recency as a measure of importance. This parallels importance sampling in static

graphs (e.g., FastGCN), where recent interactions are explicitly prioritized. However, our experiments show that recency alone is insufficient to capture the complexities of temporal interactions. Instead, incorporating node-specific contextual information at each interaction point, as done by FLASH, proves crucial for achieving robust and accurate performance in dynamic graph settings.

Learning from Large Historical Neighborhoods Learning from large historical neighborhoods in dynamic graphs poses significant challenges in computational cost and capturing long-term dependencies. To address the latter, the patching technique introduced in [Yu *et al.*, 2023], adopted in other recent studies [Tian *et al.*, 2023; Ding *et al.*, 2025]. The method splits the historical neighborhood into chronological patches, each linearly projected into a single representative vector. Since historical neighbors must be encoded (e.g., via co-occurrence encoding [Yu *et al.*, 2023]), neighborhood selection precedes it to reduce computational overhead. Thus, as proposed by [Yu *et al.*, 2023] for DYGFORMER, patching is a complementary strategy for processing large historical neighborhoods, in future link prediction tasks.

Weighing and Selection Mechanisms. Some approaches, such as TGAT [Xu *et al.*, 2020], employ attention mechanisms to weight neighbors, but only after the neighborhood has been sampled. Other works use reinforcement learning [Wang *et al.*, 2022] to learn selection strategies, but they require task-specific reward design and extensive training. Methods based on hard negative mining or curriculum sampling [Younesian *et al.*, 2024] share our motivation of identifying informative samples, but are not tailored for self-supervised learning and are not readily applicable to TGNNs. In contrast, FLASH learns to sample informative neighbors in a self-supervised, fully differentiable manner.

3 Background

Continuous Time Dynamic Graph (CTDG) is represented as a sequence of time-stamped events $\mathcal{G} = \{(x_1, t_1), (x_2, t_2) \dots\}$ where $t_1 \leq t_2 \leq \dots$. Each event (x_i, t_i) represents an action on the graph that occurred at time t_i . Each such action can be either node addition, node removal, edge addition, or edge removal. We denote $\mathcal{G}_t = (\mathcal{V}_t, \mathcal{E}_t)$ the snapshot of CTDG at time t , which is the graph received by applying all the events in $\mathcal{G}$ that occurred until time t , where $\mathcal{G}_0 = (\emptyset, \emptyset)$. We denote $F_V : \mathcal{V} \times \mathbb{R}^+ \rightarrow \mathbb{R}^{d_V}$ and $F_E : \mathcal{E} \times \mathbb{R}^+ \rightarrow \mathbb{R}^{d_E}$ as the functions that map a node or an edge to their features, at a specific point in time. For a given node v and timestamp t , we denote the historic neighborhood $\mathcal{H}_{v,t}$ as the set of all nodes that interacted with v before time t :

$$\mathcal{H}_{v,t} = \{u | (v, u) \in \mathcal{E}\} \quad (1)$$

Certain TGNNs allow the same node u to appear multiple times in $\mathcal{H}_{v,t}$, where each occurrence of u is associated with a distinct timestamp that records the time v and u interacted.

TGNNs construct node representations by sampling a subset $\mathcal{S}_{v,t}(k) \subseteq \mathcal{H}_{v,t}$ of k historical neighbors. Existing approaches employ various static selection strategies:

$$\mathcal{S}_{v,t}^{tru}(k) = \{u | r_u \leq k\} \quad (2)$$

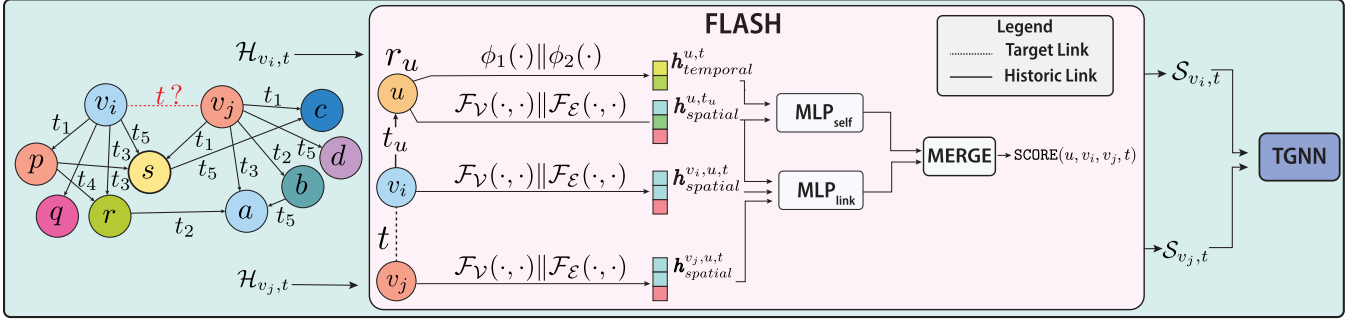


Figure 2: Overview of FLASH. Each historical neighbor u is assigned a score based on its temporal, spatial, and structural relationships with v_i and v_j . The highest-scoring neighbors are selected.

$$S_{v,t}^{\text{uni}}(k) \sim \text{Unif}\{S \subseteq \mathcal{H}_{v,t} \mid |S| = k\}. \quad (3)$$

where r_u is the rank of u with respect to the sequence of sorted neighbors from $\mathcal{H}_{v,t}$ in a decreasing order by their timestamp. Some TGNNs [Rossi *et al.*, 2020] allow sampling beyond the 1-hop historical neighborhood, either by allowing the sampling technique to sample farther nodes in advance or by applying the sampling technique recursively, i.e., sampling from the neighborhood of the sampled neighbors.

Given the sampled neighborhood $\mathcal{S}_{v,t}(k)$, the representation of node v at time t is computed as:

$$\mathbf{z}_v^t = \psi(\mathcal{S}_{v,t}(k)) \quad (4)$$

where $\psi(\cdot)$ is a parametric function of the TGNN that maps node to vector based on its sampled neighborhood.

For the task of future link prediction, given a pair of nodes (v_i, v_j) and their appropriate representations $\mathbf{z}_{v_i}^t$ and $\mathbf{z}_{v_j}^t$ at time t , a TGNN assigns a probability to the existence of a future edge between them using a learnable MERGE function:

$$p(v_i, v_j \mid t) = \text{MERGE}(\mathbf{z}_{v_i}^t, \mathbf{z}_{v_j}^t) \quad (5)$$

To train and evaluate TGNNs for future link prediction, the common approach is to split the entire sequence of interactions into two consecutive non-overlapping segments: a training prefix and an evaluation suffix. All interactions within the training prefix serve as positive examples, indicating node pairs that do form an edge at a specific time. For negative examples, random node pairs are sampled. The TGNN parameters are then updated by minimizing a binary classification loss (e.g., cross-entropy) that distinguishes positive from negative edges:

$$\mathcal{L}_{\text{task}} = - \sum_{(v_i, v_j, t) \in \text{train}} \left[y_{ij}^t \log p(v_i, v_j \mid t) + (1 - y_{ij}^t) \log (1 - p(v_i, v_j \mid t)) \right] \quad (6)$$

where y_{ij}^t is 1 for observed edges in the training prefix and 0 otherwise. Once trained, the TGNN can be used to predict edges in the evaluation suffix by computing the probability $p(v_i, v_j \mid t)$ for new future node pairs.

3.1 Theoretical analysis of Heuristic Neighborhood Samplings

The lack of flexibility in common neighborhood sampling techniques, such as k recent selection (truncation) or uniform

sampling, hinder the performance of TGNNs. Specifically, we show that there exist dynamic graphs such that any TGNN relying on these common heuristics cannot learn from them (i.e., it cannot overfit on this graph).

Theorem 1. *For any k there exists a dynamic graph on which a TGNN that applies k recent neighbors sampling cannot learn.*

Theorem 2. *For any k , there exists a dynamic graph on which a TGNN that applies uniform sampling of k historical neighbors cannot learn.*

To prove Theorem 1, we find a graph that the $k + 1$ recent neighbor is required to learn the dynamic behavior of the graph. To prove Theorem 2, we find a dynamic graph that requires consistently selecting the same recent neighbor to learn its dynamic behavior. We provide the full proofs and assumptions for Theorem 1 and Theorem 2 in Appendix D.

Implication. These common sampling heuristics discard potentially crucial historical interactions, limiting the performance of TGNNs and reduce their expressive power. We aim to develop a learnable and adaptive sampling technique that not only generalizes these heuristics but also enables TGNNs to become more expressive.

4 Method

In the previous section, we showed that current neighborhood sampling strategies do not fully account for the graph structure and its features, preventing TGNNs from capturing even simple evolving dynamics. Another limitation of these sampling heuristics stems from the fact that different TGNNs learn temporal dynamics in distinct ways due to their diverse designs. For example, TGN [Rossi *et al.*, 2020] uses memory states, whereas DyGFormer [Yu *et al.*, 2023] employs a Transformer-based architecture [Vaswani *et al.*, 2017] and jointly processes $\mathcal{S}_{v_i,t}$ and $\mathcal{S}_{v_j,t}$. Thus, enforcing the same fixed neighborhood sampling scheme across various TGNNs may undermine their performance. A learnable sampling strategy can adapt not only to the structure of the dynamic graph but also to the way each TGNN architecture exploits temporal signals.

4.1 Desiderata for Adaptive Neighborhood Sampling in TGNNs

We propose an adaptive neighborhood sampling mechanism that satisfies the following desiderata:

- (1) **Adapt to CTDG Dynamics:** We seek a mechanism $\text{SAMPLE}(\cdot)$ that accounts for node and edge attributes as well as their timestamps, i.e., the learnable parameters θ should be informed by the CTDG’s interaction patterns. Formally, for a node v with historical neighborhood $\mathcal{H}_{v,t}$ and a candidate interaction partner v' ,

$$\mathcal{S}_{v,t} = \text{SAMPLE}(v, v', \mathcal{H}_{v,t}; \theta, k), \quad (7)$$
 where $|\mathcal{S}_{v,t}| = k$.
- (2) **Generalize Existing Heuristics:** $\text{SAMPLE}(\cdot)$ should subsume commonly used heuristics (e.g., truncation, uniform sampling) as special cases via appropriate choices of θ .
- (3) **Seamless Integration with TGNNs:** Since TGNN backbones encode temporal dynamics differently (e.g., memory modules, attention), the relevance of sampled neighbors may depend on the backbone. Accordingly, the TGNN should inform the sampling procedure and the learning of θ . Moreover, the mechanism should integrate seamlessly with a broad range of TGNNs, including models with non-differentiable feature extractors.
- (4) **Self-Supervised Learning of Neighbor Importance:** Because ground-truth neighbor “importance” is unavailable, the mechanism should be learnable in a self-supervised manner that rewards neighbor subsets conducive to accurate link prediction.

4.2 FLASH

To predict a future link (v_i, v_j) at time t , we learn how informative each historical neighbor $u \in \mathcal{H}_{v_i,t}$ is for this prediction. Concretely, $\text{SAMPLE}(\cdot; \theta)$ assigns a score to each historical neighbor and selects the k neighbors with the highest score.

Learning the scoring function. We first construct an informative representation for each historical neighbor that captures both **structural** (interaction context) and **temporal** (time-dependent) information (**D1**). For a neighbor $u \in \mathcal{H}_{v_i,t}$, let t_u denote the time at which u and v_i interacted, and let r_u denote the rank of u when sorting $\mathcal{H}_{v_i,t}$ in decreasing order of interaction time. We define:

$$\mathbf{h}_{\text{spatial}}^u = [F_{\mathcal{V}}(u, t_u) \parallel F_{\mathcal{E}}(v_i, u, t_u) \parallel \mathcal{M}(u)] \quad (8)$$

$$\mathbf{h}_{\text{spatial}}^{v_i, u, t} = [F_{\mathcal{V}}(v_i, t_u) \parallel F_{\mathcal{V}}(v_i, t) \parallel \mathcal{M}(v_i)] \quad (9)$$

$$\mathbf{h}_{\text{spatial}}^{v_j, u, t} = [F_{\mathcal{V}}(v_j, t_u) \parallel F_{\mathcal{V}}(v_j, t) \parallel \mathcal{M}(v_j)] \quad (10)$$

$$\mathbf{h}_{\text{temporal}}^{u, t} = [\phi_1(t - t_u) \parallel \phi_2(r_u)], \quad (11)$$

where $\phi_1(\cdot)$ and $\phi_2(\cdot)$ are Time2Vec representations kazemi2019time2vec, $F_{\mathcal{V}}$ maps nodes and timestamps to their corresponding features, $F_{\mathcal{E}}$ maps edges and timestamps to their corresponding features (as defined in

Section 3), $\mathcal{M}(\cdot)$ denotes learnable node features, and $\parallel$ denotes concatenation. Following the standard TGNN pipeline and prior work deng2024taser, we project the neighbor and link representations to a common dimensionality using MLPs, and then combine them via a MERGE operator:

$$\text{SCORE}_{u, v_i, v_j, t} = \text{MERGE} \left(\text{MLP}_{\text{self}}(\mathbf{h}_{\text{spatial}}^u, \mathbf{h}_{\text{temporal}}^{u, t}), \text{MLP}_{\text{link}}(\mathbf{h}_{\text{spatial}}^{v_i, u, t}, \mathbf{h}_{\text{spatial}}^{v_j, u, t}) \right) \quad (12)$$

where MERGE is an MLP with a single hidden layer, as suggested for combining source–destination information yu2023towards, tian2023freedyg. $\text{SCORE}_{u, v_i, v_j, t}$ measures the utility of neighbor u for predicting the link (v_i, v_j) at time t , where larger values indicate higher importance. FLASH then selects the k neighbors with the highest scores. Figure 2 provides a detailed schematic of FLASH.

Training FLASH. For each node pair (v_i, v_j) at time t , let $y_{ij}^t \in \{0, 1\}$ indicate whether (v_i, v_j) is a *positive* (observed) edge ($y_{ij}^t = 1$) or a *negative* (unobserved) edge ($y_{ij}^t = 0$). Let a TGNN compute edge probabilities using the subsets selected by FLASH, $(\mathcal{S}_{v_i, v_j, t}, \mathcal{S}_{v_j, v_i, t})$, and using uniformly sampled subsets $(\mathcal{S}_{v_i, t}^{\text{uni}}, \mathcal{S}_{v_j, t}^{\text{uni}})$. We define the probability difference:

$$\Delta_{ij}^t = p(v_i, v_j | t; \mathcal{S}_{v_i, v_j, t}, \mathcal{S}_{v_j, v_i, t}) - p(v_i, v_j | t; \mathcal{S}_{v_i, t}^{\text{uni}}, \mathcal{S}_{v_j, t}^{\text{uni}}). \quad (13)$$

Let $\overline{s_{v_i}}, \overline{s_{v_j}}$ denote the average FLASH scores over nodes in $\mathcal{S}_{v_i, v_j, t}$ and $\mathcal{S}_{v_j, v_i, t}$, respectively, and let $\overline{s_{v_i}^{\text{uni}}}, \overline{s_{v_j}^{\text{uni}}}$ denote the corresponding averages over $\mathcal{S}_{v_i, t}^{\text{uni}}$ and $\mathcal{S}_{v_j, t}^{\text{uni}}$. For positive edges, we prefer $\Delta_{ij}^t > 0$ (our subsets increase the predicted probability), while for negative edges we prefer $\Delta_{ij}^t < 0$ (our subsets decrease it). We enforce this preference using a pairwise RankNet burges2005learning logistic loss:

$$\mathcal{L}_{ij}^t = \begin{cases} - \left[\log \sigma(\overline{s_{v_j}} - \overline{s_{v_j}^{\text{uni}}}) + \log \sigma(\overline{s_{v_i}} - \overline{s_{v_i}^{\text{uni}}}) \right] & \text{if } (y_{ij}^t - \frac{1}{2}) \Delta_{ij}^t > 0, \\ - \left[\log \sigma(\overline{s_{v_j}^{\text{uni}}} - \overline{s_{v_j}}) + \log \sigma(\overline{s_{v_i}^{\text{uni}}} - \overline{s_{v_i}}) \right] & \text{otherwise.} \end{cases} \quad (14)$$

Minimizing $\mathcal{L}_{ij}^t$ encourages the score averages to move in the direction consistent with y_{ij}^t without requiring any ground-truth importance supervision (**D3**). When the condition in Equation (14) holds, the neighborhoods selected by FLASH improve over uniform sampling, and the loss encourages higher scores for selected neighbors over the uniformly sampled neighbors. Otherwise, uniform sampling performs better and the loss correspondingly increases the scores of uniformly sampled neighbors relative to those selected by

Method ↓ / Dataset →	WIKIPEDIA AP ↑	REDDIT AP ↑	MOOC AP ↑	LASTFM AP ↑	SOCIAL EVO. AP ↑	ENRON AP ↑	UCI AP ↑
TGAT + Trunc.	94.05±0.06	93.63±0.16	79.69±0.24	65.64±0.34	85.36±0.25	72.91±0.58	79.74±0.37
TGAT + Uni.	62.60±0.36	87.94±0.03	60.03±0.11	50.66±0.11	53.22±0.79	51.34±0.32	60.42±0.22
TGAT + NLB	91.49±0.25	92.78±0.11	77.37±0.17	65.64±0.34	83.19±0.07	70.69±0.93	77.80±1.25
TGAT + TASER	67.78±0.99	89.14±0.59	67.10±0.56	52.28±0.52	70.40±0.45	60.55±0.91	70.62±0.33
TGAT + FLASH	94.67 ±0.38	95.92 ±0.13	80.24 ±0.54	74.94 ±1.49	92.17 ±0.25	79.04 ±0.94	87.84 ±0.12
TGN + Trunc.	98.55±0.05	98.61±0.03	90.13±0.64	82.62±2.09	91.63±0.51	86.51±2.29	93.34±0.25
TGN + Uni.	98.49±0.08	98.59±0.01	83.08±1.11	67.60±5.66	65.52±6.06	85.47±2.00	93.34±0.25
TGN + NLB	98.31±0.09	98.61±0.04	89.81±0.60	80.45±1.94	91.19±0.38	86.51±2.29	92.91±0.43
TGN + TASER	98.15±0.07	98.63±0.04	90.57±1.02	75.58±4.07	92.24±0.62	86.87±0.081	89.86±2.61
TGN + FLASH	98.73 ±0.06	99.06 ±0.03	91.19 ±0.51	90.54 ±0.62	93.43 ±0.11	89.90 ±0.76	95.17 ±0.16
GRAPHMIXER + Trunc.	96.23±0.24	95.17±0.03	80.71±1.11	72.98±0.07	87.09±0.12	81.63±0.47	93.14±0.44
GRAPHMIXER + Uni.	77.06±0.16	89.80±0.05	64.75±0.39	63.96±0.09	55.69±0.12	55.65±3.04	71.27±2.79
GRAPHMIXER + NLB	95.09±0.12	95.17±0.03	78.61±0.09	72.98±0.07	85.66±0.09	81.01±0.30	92.51±0.60
GRAPHMIXER + TASER	93.27±1.43	94.78±0.17	78.97±0.66	72.24±0.93	87.57±2.53	77.90±0.48	88.59±1.10
GRAPHMIXER + FLASH	97.51 ±0.25	96.62 ±0.11	80.85 ±0.52	82.68 ±0.74	92.84 ±0.11	85.74 ±0.70	93.21 ±0.61
DYGFORMER + Trunc.	96.91±0.05	95.15±0.07	82.47±0.07	74.81±0.09	85.73±0.06	82.35±0.66	89.61±0.22
DYGFORMER + Uni.	96.87±0.07	95.15±0.07	82.47±0.07	74.81±0.09	85.71±0.06	82.35±0.66	89.61±0.22
DYGFORMER + NLB	96.74±0.07	95.03±0.09	81.09±0.05	74.52±0.18	85.40±0.10	81.86±0.45	89.12±0.19
DYGFORMER + TASER	95.47±0.32	93.83±0.33	81.56±0.67	72.72±1.73	90.09±0.11	81.74±0.64	88.15±0.59
DYGFORMER + FLASH	98.17 ±0.04	98.11 ±0.02	82.96 ±0.42	86.09 ±0.13	92.41 ±0.11	88.70 ±0.17	92.96 ±0.16
FREEDYG + Trunc.	98.35±0.01	97.53±0.02	85.02±0.02	80.19±0.04	90.88±0.05	88.77±0.20	95.21±0.33
FREEDYG + Uni.	98.33±0.02	95.98±0.03	67.52±0.18	67.72±0.10	89.16±0.03	88.46±0.08	95.21±0.33
FREEDYG + NLB	98.33±0.01	97.59±0.02	83.56±0.12	80.02±0.09	90.40±0.05	88.77±0.20	95.17±0.26
FREEDYG + TASER	97.94±0.35	97.11±0.30	76.97±3.87	78.54±1.90	90.84±1.09	87.24±1.47	93.15±1.46
FREEDYG + FLASH	98.94 ±0.04	98.72 ±0.02	85.69 ±0.27	88.46 ±0.07	93.62 ±0.05	91.31 ±0.01	95.86 ±0.16

Table 1: Comparison of various neighborhood sampling methods on transductive future edge prediction using 2 historical neighbors on different datasets from DyGLib. The best performing method is in bold.

FLASH. We train FLASH jointly with a given TGNN by optimizing the sum of Equation (14) and Equation (6). Gradients from Equation (6) update only the TGNN parameters, while for Equation (14) we stop gradients through Δ_{ij}^t , so that updates flow only to FLASH parameters (D4). To ensure constant runtime, FLASH first applies a recent-neighbors finder `deng2024taser` to restrict the candidate set to a fixed size. It then computes scores only within this reduced neighborhood rather than over the full historical neighborhood. In Appendix G, we provide a throughout analysis of FLASH and TASER.

4.3 Theoretical Analysis of FLASH

The components of FLASH adapt to the evolution of the dynamic graph over time and to the interactions being predicted, by incorporating both spatiotemporal features and potentially interacting nodes when constructing each sampled neighborhood. We show that FLASH can replicate both k -recent neighbor selection and uniform sampling. Furthermore, we show that the graphs used in the proofs of Theorem 1 and Theorem 2 can be learned by a TGNN that utilizes FLASH. Therefore:

Theorem 3. *FLASH is more expressive than k -recent neighbor selection and uniform sampling.*

We provide the full proof of Theorem 3 in Appendix D.

5 Experiments

We evaluate FLASH across multiple dynamic graph benchmarks and compare it with recent state-of-the-art baselines.

Section 5.1 presents our key empirical findings, and Section 5.2 provides additional ablation studies. Additional results are reported in Appendix C. Our experiments aim to answer the following research questions:

- (RQ1) How does our proposed strategy compare to established sampling baselines in predictive accuracy?
- (RQ2) Does our method generalize across varying graph sizes, sparsity levels, and temporal granularities?
- (RQ3) What is the computational overhead of our approach relative to existing methods?

Experimental Setup. We test our method on five common TGNNs: TGAT [Xu *et al.*, 2020], TGN [Rossi *et al.*, 2020], GRAPHMIXER [Cong *et al.*, 2023], DYGFORMER [Yu *et al.*, 2023], and FREEDYG [Tian *et al.*, 2023]. Detailed descriptions of each model are provided in Appendix B. Each model is trained with four neighbor sampling strategies: Truncation [Cong *et al.*, 2023; Yu *et al.*, 2023], Uniform [Rossi *et al.*, 2020], No-Look-Back (NLB) [Luo and Li, 2024], TASER with the recent-neighbors finder and a linear predictor [Deng *et al.*, 2024], and our proposed method. Following standard practice, we use a chronological 70%-15%-15% train-validation-test split. Additional parameters regarding the training scheme and implementation-specific details are provided in Appendix E.

Importance of Small Neighborhoods As discussed in Section 1, processing large neighborhoods is computationally prohibitive in both time and space. Figure 4 reports the throughput of state-of-the-art TGNNs for varying neighborhood sizes. The results show that increasing the neighborhood processed in parallel substantially degrades model

Method ↓ / Dataset →	WIKIPEDIA AP ↑	REDDIT AP ↑	MOOC AP ↑	LASTFM AP ↑	SOCIAL EVO. AP ↑	ENRON AP ↑	UCI AP ↑
TGAT + Trunc.	94.26 $\pm$ 0.09	90.82 $\pm$ 0.18	78.12 $\pm$ 0.18	72.35 $\pm$ 0.42	82.73 $\pm$ 0.29	70.25 $\pm$ 1.30	80.95 $\pm$ 0.29
TGAT + Uni.	62.77 $\pm$ 0.43	82.85 $\pm$ 0.13	58.38 $\pm$ 0.14	50.75 $\pm$ 0.07	53.06 $\pm$ 0.40	51.63 $\pm$ 0.72	59.65 $\pm$ 0.36
TGAT + NLB.	91.46 $\pm$ 0.24	89.70 $\pm$ 0.15	75.42 $\pm$ 0.22	72.35 $\pm$ 0.42	80.66 $\pm$ 0.12	68.77 $\pm$ 1.03	78.47 $\pm$ 0.68
TGAT + TASER	67.62 $\pm$ 1.08	84.51 $\pm$ 0.75	63.60 $\pm$ 0.52	53.53 $\pm$ 0.85	67.22 $\pm$ 0.59	58.46 $\pm$ 1.38	64.94 $\pm$ 0.90
TGAT+FLASH	94.70 $\pm$ 0.35	94.24 $\pm$ 0.10	78.85 $\pm$ 0.61	79.56 $\pm$ 1.57	90.23 $\pm$ 0.25	76.70 $\pm$ 1.13	87.77 $\pm$ 0.17
TGN + Trunc.	97.84 $\pm$ 0.05	97.28 $\pm$ 0.08	90.02 $\pm$ 1.30	87.23 $\pm$ 1.18	89.08 $\pm$ 1.42	79.66 $\pm$ 2.23	89.36 $\pm$ 0.52
TGN + Uni.	97.82 $\pm$ 0.12	97.13 $\pm$ 0.15	82.47 $\pm$ 0.41	76.10 $\pm$ 6.58	57.15 $\pm$ 0.35	77.91 $\pm$ 2.34	89.36 $\pm$ 0.52
TGN + NLB.	97.61 $\pm$ 0.08	97.23 $\pm$ 0.09	89.19 $\pm$ 0.48	85.79 $\pm$ 2.00	86.51 $\pm$ 2.51	79.66 $\pm$ 2.23	88.34 $\pm$ 0.52
TGN + TASER	97.39 $\pm$ 0.17	97.33 $\pm$ 0.12	89.90 $\pm$ 1.18	81.97 $\pm$ 3.73	90.31 $\pm$ 1.50	78.99 $\pm$ 2.56	85.28 $\pm$ 1.22
TGN+FLASH	98.08 $\pm$ 0.12	98.24 $\pm$ 0.07	90.50 $\pm$ 0.61	92.04 $\pm$ 0.40	92.06 $\pm$ 0.48	82.95 $\pm$ 1.02	92.40 $\pm$ 0.27
GraphMixer + Trunc.	95.80 $\pm$ 0.24	92.21 $\pm$ 0.07	79.38 $\pm$ 0.18	79.52 $\pm$ 0.11	84.83 $\pm$ 0.14	76.17 $\pm$ 0.43	91.60 $\pm$ 0.19
GraphMixer + Uni.	78.72 $\pm$ 0.12	84.29 $\pm$ 0.10	65.91 $\pm$ 0.30	73.74 $\pm$ 0.07	51.13 $\pm$ 0.15	52.46 $\pm$ 3.33	71.69 $\pm$ 1.43
GraphMixer + NLB	94.34 $\pm$ 0.10	92.21 $\pm$ 0.07	76.89 $\pm$ 0.10	79.52 $\pm$ 0.11	83.09 $\pm$ 0.08	74.99 $\pm$ 0.34	90.89 $\pm$ 0.23
GraphMixer + TASER	92.42 $\pm$ 1.54	91.25 $\pm$ 0.30	76.99 $\pm$ 0.59	79.43 $\pm$ 1.02	85.10 $\pm$ 2.73	70.70 $\pm$ 0.44	86.78 $\pm$ 1.76
GraphMixer + FLASH	97.06 $\pm$ 0.29	94.77 $\pm$ 0.21	79.44 $\pm$ 0.55	86.70 $\pm$ 0.74	90.74 $\pm$ 0.11	80.61 $\pm$ 1.41	91.66 $\pm$ 0.24
DyGFormer + Trunc.	97.02 $\pm$ 0.07	93.40 $\pm$ 0.08	81.27 $\pm$ 0.09	79.97 $\pm$ 0.11	82.86 $\pm$ 0.07	80.46 $\pm$ 0.91	89.99 $\pm$ 0.18
DyGFormer + Uni.	97.01 $\pm$ 0.06	93.40 $\pm$ 0.08	81.27 $\pm$ 0.09	79.86 $\pm$ 0.12	82.80 $\pm$ 0.07	80.46 $\pm$ 0.91	89.81 $\pm$ 0.16
DyGFormer + NLB	96.86 $\pm$ 0.08	93.27 $\pm$ 0.09	79.62 $\pm$ 0.11	79.66 $\pm$ 0.23	82.68 $\pm$ 0.10	80.09 $\pm$ 0.38	89.52 $\pm$ 0.22
DyGFormer + TASER	95.70 $\pm$ 0.26	91.05 $\pm$ 0.51	80.37 $\pm$ 0.65	77.93 $\pm$ 1.60	88.20 $\pm$ 0.25	79.28 $\pm$ 1.20	84.98 $\pm$ 0.58
DyGFormer + FLASH	97.91 $\pm$ 0.04	97.36 $\pm$ 0.04	81.42 $\pm$ 0.51	88.55 $\pm$ 0.15	90.56 $\pm$ 0.28	84.88 $\pm$ 0.66	92.73 $\pm$ 0.15
FreeDyG + Trunc.	98.03 $\pm$ 0.03	96.34 $\pm$ 0.04	84.04 $\pm$ 0.06	85.09 $\pm$ 0.07	89.06 $\pm$ 0.09	84.59 $\pm$ 0.33	93.98 $\pm$ 0.15
FreeDyG + Uni.	98.01 $\pm$ 0.02	93.86 $\pm$ 0.04	68.53 $\pm$ 0.24	76.41 $\pm$ 0.09	86.92 $\pm$ 0.09	83.96 $\pm$ 0.21	93.98 $\pm$ 0.15
FreeDyG + NLB.	97.99 $\pm$ 0.01	96.34 $\pm$ 0.04	82.13 $\pm$ 0.10	84.98 $\pm$ 0.13	88.38 $\pm$ 0.17	84.59 $\pm$ 0.33	93.81 $\pm$ 0.16
FreeDyG + TASER	97.67 $\pm$ 0.34	95.63 $\pm$ 0.38	74.93 $\pm$ 4.60	84.07 $\pm$ 1.25	88.87 $\pm$ 1.17	81.31 $\pm$ 2.04	92.02 $\pm$ 1.32
FreeDyG + FLASH	98.59 $\pm$ 0.03	98.21 $\pm$ 0.04	84.38 $\pm$ 0.25	90.73 $\pm$ 0.09	91.53 $\pm$ 0.63	87.13 $\pm$ 0.62	94.60 $\pm$ 0.13

Table 2: Comparison of our suggested sampling strategy with other sampling strategies (Truncation, Uniform and NLB) in the inductive setting. Results are reported in AP for future edge prediction with random negative sampling over five different seeds. The significantly best result for each benchmark appears in bold font. 2 historical neighbors are used by each method.

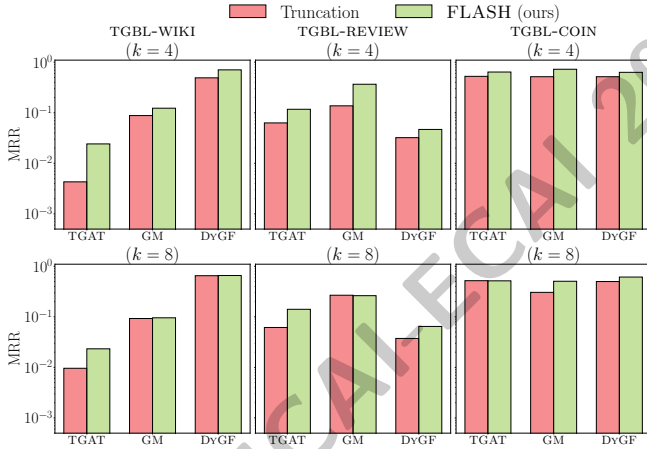


Figure 3: FLASH vs. Truncation on the TGB benchmark. Results are reported in MRR (dynamic link prediction) with random negative sampling over three runs, using $k = 4$ and $k = 8$ historical neighbors. Here GM refers to GRAPHMIXER and DYGF refers to DYGFORMER.

throughput. Consequently, TGNNs should employ small neighborhoods (2–4) to achieve peak runtime performance. Accordingly, we evaluate neighborhood selection strategies under a small neighborhood budget (2–4) to determine which strategy is optimal for future link prediction while sustaining peak throughput in TGNNs.

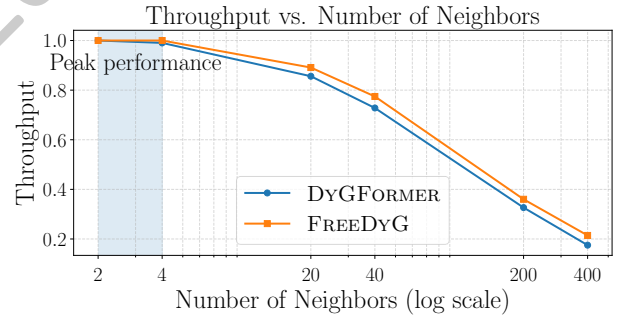


Figure 4: Throughput comparison of TGNNs as a function of selected neighborhood size, reported as the number of edges processed per second. Values are normalized by the throughput at neighborhood size 2 (e.g., if DYGFORMER processes 200 edges/s with neighborhood size 2 and 100 edges/s with neighborhood size 40, the reported value is 0.5).

5.1 Results

Evaluation with DyGLib. We use the WIKIPEDIA, REDDIT, MOOC, LASTFM, ENRON, SOCIAL EVO., and UCI datasets from the DyGLib benchmark [Yu *et al.*, 2023]. DyGLib includes datasets spanning a wide range of graph sizes, with some containing over a million edges (RQ2). Full dataset statistics are provided in Appendix A. We evaluate future edge prediction (Equation (6)) under two settings: (1) *transductive*, where all nodes appear during training, and (2)

ϕ_2	ϕ_1	$\mathcal{M}$	MLP_{link}	WIKIPEDIA	ENRON	UCI
✓	✓	✓	✓	98.73 $\pm$ 0.07	89.63 $\pm$ 0.72	95.08 $\pm$ 0.19
✗	✗	✓	✓	98.20 $\pm$ 0.06	89.50 $\pm$ 0.74	95.08 $\pm$ 0.33
✓	✓	✗	✓	98.58 $\pm$ 0.06	86.45 $\pm$ 1.25	93.27 $\pm$ 0.61
✓	✓	✓	✗	98.54 $\pm$ 0.09	87.38 $\pm$ 0.81	93.06 $\pm$ 0.60
✓	✓	✓	✓	98.73 $\pm$ 0.06	89.90 $\pm$ 0.76	95.17 $\pm$ 0.16

Table 3: Ablation of design choices of FLASH. (✓) denotes inclusion and (✗) denotes exclusion of the component. (*transductive*)

inductive, where test nodes are unseen during training. Tables 1 and 2 report results across models and methods. Overall, FLASH consistently outperforms previously suggested neighborhood selection strategies, both in the transductive setting and the inductive setting (**RQ1**). For example, on SOCIAL EVO., FLASH achieves a $\sim 8\%$ relative improvement over truncation for TGAT, and on LASTFM, a $\sim 15\%$ relative improvement for FREEDYG, demonstrating strong effectiveness even for computationally expensive models.

Evaluation with TGB. We evaluate FLASH on the TGB benchmark [Huang *et al.*, 2023] using TGBL-WIKI, TGBL-REVIEW, and TGBL-COIN. In contrast to DyGLib, TGB samples multiple random negative edges for each positive edge and is evaluated using Mean Reciprocal Rank (MRR). From Figure 3, FLASH yields performance that is consistently on par with or better than baselines across TGNNs and neighborhood sizes. Notably, on TGBL-REVIEW, our method improves performance by $\sim 2.67\times$ for GRAPHMIXER when $k = 4$, highlighting robustness across diverse dynamic graph scenarios (**RQ2**).

5.2 Ablation Study

FLASH Design Choices. We ablate key components of FLASH, namely the historical position embedding ϕ_2 , temporal embedding ϕ_1 , learnable node embedding $\mathcal{M}$, and the link-awareness module MLP_{link} in Equation (12), to assess the contribution of temporal, structural, and interaction context for future edge prediction (**RQ2**). Table 3 shows that the positional embedding ϕ_2 is particularly important on ENRON and UCI, indicating that relying solely on temporal information is suboptimal. The interaction context module, MLP_{link} , is consistently beneficial across all datasets, providing reliable performance improvements.

Computational Overhead. Table 4 reports a throughput comparison for FLASH, measured as the number of processed edges per second, against the baselines (**RQ3**). For a fair comparison, both TASER and FLASH employ the same neighbor finder, which constrains their neighborhoods to an identical size, this benefits both methods relative to uniform sampling. Because all experiments use a unified batch size, model latency is inversely proportional to throughput. Appendix F further provides a time-complexity analysis of FLASH and the baselines.

Model	Uniform	NLB	TASER	FLASH
TGAT	51%	44%	59%	75%
TGN	67%	80%	82%	82%
GRAPHMIXER	42%	40%	59%	64%
DYGFORMER	97%	87%	86%	94%
FREEDYG	84%	85%	91%	92%

Table 4: Throughput comparison of neighbor sampling strategies. Values denote the number of edges processed per second by the TGNN relative to the Truncation baseline. For example, if TGAT with Uniform processes 300 edges/s over a full run while TGAT with Truncation processes 600 edges/s, the reported throughput for TGAT combined with Uniform is 50%.

6 Conclusion

In this work, we introduce FLASH, a flexible, end-to-end learnable neighborhood selection framework for dynamic graph learning. FLASH factorizes the neighbor scoring function into spatial and temporal components, ensuring that neighbor importance reflects both topological structure and temporal context. Our training objective couples the learned scores with downstream predictive performance, enabling effective optimization of the selection policy even when ground-truth importance scores are unavailable and TGNN operations are non-differentiable. This design allows FLASH to integrate seamlessly with a broad range of TGNN architectures, providing an efficient mechanism for selecting the most relevant historical neighbors. We further evaluate FLASH across multiple dynamic graph benchmarks and show that it matches or outperforms the considered baselines.

Contribution Statement

O.F. contributed to the design of the method, implementation, experiments, figure/plot design, and manuscript writing. K.S.I.M. contributed to the design of the method, figure/plot design, and manuscript writing. C.B.S. contributed to manuscript writing. C.B. and M.E. contributed to manuscript writing and equally supervised the work.

References

- [Chen *et al.*, 2018] Jie Chen, Tengfei Ma, and Cao Xiao. Fastgcn: Fast learning with graph convolutional networks via importance sampling. In *International Conference on Learning Representations*. International Conference on Learning Representations, ICLR, 2018.
- [Chiang *et al.*, 2019] Wei-Lin Chiang, Xuanqing Liu, Si Si, Yang Li, Samy Bengio, and Cho-Jui Hsieh. Cluster-gcn: An efficient algorithm for training deep and large graph convolutional networks. In *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, pages 257–266, 2019.
- [Cong *et al.*, 2023] Weilin Cong, Si Zhang, Jian Kang, Baichuan Yuan, Hao Wu, Xin Zhou, Hanghang Tong, and Mehrdad Mahdavi. Do we really need complicated model architectures for temporal networks? *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*.

- [Deng *et al.*, 2024] Gangda Deng, Hongkuan Zhou, Hanqing Zeng, Yinglong Xia, Christopher Leung, Jianbo Li, Rajgopal Kannan, and Viktor Prasanna. Taser: Temporal adaptive sampling for fast and accurate dynamic graph representation learning. In *2024 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*, pages 926–937. IEEE, 2024.
- [Ding *et al.*, 2025] Zifeng Ding, Yifeng Li, Yuan He, Antonio Norelli, Jingcheng Wu, Volker Tresp, Michael M Bronstein, and Yunpu Ma. Dygmamba: Efficiently modeling long-term temporal dependency on continuous-time dynamic graphs with state space models. *Transactions on Machine Learning Research*, 2025.
- [Gravina *et al.*, 2024] Alessio Gravina, Giulio Lovisotto, Claudio Gallicchio, Davide Bacciu, and Claas Grohnfeldt. Long range propagation on continuous-time dynamic graphs. In *International Conference on Machine Learning*, pages 16206–16225. PMLR, 2024.
- [Hamilton *et al.*, 2017] Will Hamilton, Zhitao Ying, and Jure Leskovec. Inductive representation learning on large graphs. *Advances in neural information processing systems*, 30, 2017.
- [Huang *et al.*, 2023] Shenyang Huang, Farimah Poursafaei, Jacob Danovitch, Matthias Fey, Weihua Hu, Emanuele Rossi, Jure Leskovec, Michael Bronstein, Guillaume Rabusseau, and Reihaneh Rabbany. Temporal graph benchmark for machine learning on temporal graphs. *Advances in Neural Information Processing Systems*, 36, 2023.
- [Kumar *et al.*, 2019] Srijan Kumar, Xikun Zhang, and Jure Leskovec. Predicting dynamic embedding trajectory in temporal interaction networks. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1269–1278, 2019.
- [Luo and Li, 2024] Yuhong Luo and Pan Li. Scalable and efficient temporal graph representation learning via forward recent sampling. In *The Third Learning on Graphs Conference*, 2024.
- [Ni *et al.*, 2019] Jianmo Ni, Jiacheng Li, and Julian McAuley. Justifying recommendations using distantly-labeled reviews and fine-grained aspects. In *Proceedings of the 2019 conference on empirical methods in natural language processing and the 9th international joint conference on natural language processing (EMNLP-IJCNLP)*, pages 188–197, 2019.
- [Rossi *et al.*, 2020] Emanuele Rossi, Ben Chamberlain, Fabrizio Frasca, Davide Eynard, Federico Monti, and Michael Bronstein. Temporal graph networks for deep learning on dynamic graphs. *CoRR*, abs/2006.10637, 2020.
- [Shamsi *et al.*, 2022] Kiarash Shamsi, Friedhelm Victor, Murat Kantarcioglu, Yulia Gel, and Cuneyt G Akcora. Chartalist: Labeled graph datasets for utxo and account-based blockchains. *Advances in Neural Information Processing Systems*, 35:34926–34939, 2022.
- [Shetty and Adibi, 2004] Jitesh Shetty and Jafar Adibi. The enron email dataset database schema and brief statistical report. *Information sciences institute technical report, University of Southern California*, 4(1):120–128, 2004.
- [Tian *et al.*, 2023] Yuxing Tian, Yiyan Qi, and Fan Guo. Freedyg: Frequency enhanced continuous-time dynamic graph model for link prediction. In *The Twelfth International Conference on Learning Representations*, 2023.
- [Vaswani *et al.*, 2017] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [Wang *et al.*, 2022] Kaixin Wang, Cheng Long, Da Yan, Jie Zhang, and H. V. Jagadish. Reinforcement learning enhanced weighted sampling for accurate subgraph counting on fully dynamic graph streams, 2022.
- [Xu *et al.*, 2020] Da Xu, Chuanwei Ruan, Evren Korpeoglu, Sushant Kumar, and Kannan Achan. Inductive representation learning on temporal graphs. *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26-30, 2020*.
- [Xu *et al.*, 2024] Yuanyuan Xu, Wenjie Zhang, Ying Zhang, Maria Orlowska, and Xuemin Lin. Timesgn: Scalable and effective temporal graph neural network. In *2024 IEEE 40th International Conference on Data Engineering (ICDE)*, pages 3297–3310. IEEE, 2024.
- [Ying *et al.*, 2018] Rex Ying, Ruining He, Kaifeng Chen, Pong Eksombatchai, William L Hamilton, and Jure Leskovec. Graph convolutional neural networks for web-scale recommender systems. In *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*, pages 974–983, 2018.
- [Younesian *et al.*, 2024] Taraneh Younesian, Daniel Daza, Emile van Krieken, Thiviyan Thanapalasingam, and Peter Bloem. Grapes: Learning to sample graphs for scalable graph neural networks, 2024.
- [Yu *et al.*, 2023] Le Yu, Leilei Sun, Bowen Du, and Weifeng Lv. Towards better dynamic graph learning: New architecture and unified library. *Advances in Neural Information Processing Systems*, 36:67686–67700, 2023.
- [Zeng *et al.*, 2020] Hanqing Zeng, Hongkuan Zhou, Ajitesh Srivastava, Rajgopal Kannan, and Viktor Prasanna. Graphsaint: Graph sampling based inductive learning method. In *International Conference on Learning Representations*, 2020.
- [Zhou *et al.*, 2022] Hongkuan Zhou, Da Zheng, Israt Nisa, Vasileios Ioannidis, Xiang Song, and George Karypis. Tgl: A general framework for temporal gnn training on billion-scale graphs. *arXiv preprint arXiv:2203.14883*, 2022.
- [Zou *et al.*, 2024] Tao Zou, Yuhao Mao, Junchen Ye, and Bowen Du. Repeat-aware neighbor sampling for dynamic graph learning. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 4722–4733, 2024.