

FairGC: Fairness-aware Graph Condensation

Yihan Gao¹, Chenxi Huang¹, Wen Shi¹, Ke Sun², Ziqi Xu³, Xikun Zhang³, Mingliang Hou⁴, Renqiang Luo^{1*}

¹College of Computer Science and Technology, Jilin University, Changchun, China

²Key Laboratory of Advanced Design and Intelligent Computing, Dalian University, Dalian, China

³School of Computing Technologies, RMIT University, Melbourne, Australia

⁴Guangdong Institute of Smart Education, Jinan University, Guangzhou, China

Email: {gaoyh5523, huangcx5523, shiwen24}@mails.jlu.edu.cn, sunke@dlu.edu.cn

{ziqi.xu, xikun.zhang}@rmit.edu.au, teemohold@outlook.com, lrenqiang@jlu.edu.cn

Abstract—Graph condensation (GC) has become a vital strategy for scaling Graph Neural Networks by compressing massive datasets into small, synthetic node sets. While current GC methods effectively maintain predictive accuracy, they are primarily designed for utility and often ignore fairness constraints. Because these techniques are bias-blind, they frequently capture and even amplify demographic disparities found in the original data. This leads to synthetic proxies that are unsuitable for sensitive applications like credit scoring or social recommendations. To solve this problem, we introduce FairGC, a unified framework that embeds fairness directly into the graph distillation process. Our approach consists of three key components. First, a Distribution-Preserving Condensation module synchronizes the joint distributions of labels and sensitive attributes to stop bias from spreading. Second, a Spectral Encoding module uses Laplacian eigen-decomposition to preserve essential global structural patterns. Finally, a Fairness-Enhanced Neural Architecture employs multi-domain fusion and a label-smoothing curriculum to produce equitable predictions. Rigorous evaluations on four real-world datasets, show that FairGC provides a superior balance between accuracy and fairness. Our results confirm that FairGC significantly reduces disparity in Statistical Parity and Equal Opportunity compared to existing state-of-the-art condensation models. The codes are available at <https://github.com/LuoRenqiang/FairGC>.

Index Terms—Fairness, Graph neural networks, Graph condensation

I. INTRODUCTION

Graph learning has become a fundamental framework for modeling complex interactions across various domains [1], [2], yet the rapid expansion of modern datasets has made training standard Graph Neural Networks (GNNs) increasingly difficult due to heavy memory and time costs [3]. To overcome these computational barriers, graph condensation (GC) has emerged as a powerful paradigm that compresses massive, large-scale graphs into a tiny set of synthetic nodes while ensuring that models trained on this compact proxy achieve performance comparable to the original data [4]. This technique is currently finding widespread use in data-heavy applications, such as analyzing billion-scale social networks to map user connections or processing massive financial transaction graphs for real-time risk assessment [5].

Despite the technical maturity of mainstream GC paradigms, these frameworks are primarily engineered to maximize pre-

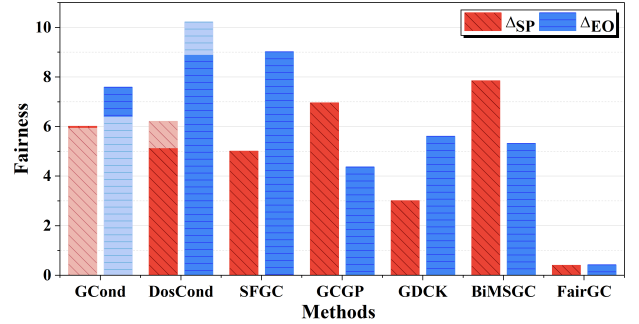


Fig. 1. Fairness evaluation of GC methods on the Pokec-n-G dataset. Both vanilla GC methods (darker bars) and their fairness-enhanced versions integrated with FairGNN (lighter bars) exhibit worse fairness compared to our proposed FairGC. Δ_{SP} and Δ_{EO} are employed as fairness metrics (refer to Section III-B), where lower values represent better fairness.

dictive utility while remaining largely indifferent to the underlying fairness constraints [6]. Since these methods operate in a bias-blind manner, they often inadvertently capture and even amplify demographic disparities present in the original data, leading to synthetic surrogates that are unsuitable for socially sensitive applications [7]–[9]. In practical tasks like credit risk assessment or social friendship recommendation, this lack of fairness awareness results in significant gaps in Statistical Parity (Δ_{SP}) [10] and Equal Opportunity (Δ_{EO}) [11], where protected groups may suffer from lower positive prediction rates or disproportionately higher misclassification errors. Ultimately, the failure to decouple predictive signals from sensitive attributes during the condensation process yields a biased proxy that fails to provide equitable decisions, highlighting a critical research gap in developing fairness-aware data reduction techniques.

To quantitatively demonstrate the fairness risks inherent in existing GC techniques, we evaluated six GC methods (GCond [12], DosCond [13], SFGC [14], GCGP [15], GDCK [16], BiMSGC [17]) on the Pokec-n-G dataset. Our experimental setup followed established fair graph learning protocols, employing Δ_{SP} and Δ_{EO} to measure performance gaps across sensitive groups. The results reveal that even when these frameworks are integrated with FairGNN [18], a widely used fairness-enhancing backbone, they still exhibit

*Corresponding author

significantly higher disparities compared to our proposed FairGC. This empirical evidence underscores a critical flaw: current methods treat data reduction and fairness as decoupled objectives, resulting in synthetic graphs that are fundamentally unsuitable for socially sensitive applications.

Crucially, existing GC paradigms fail to fundamentally resolve fairness concerns, and even the integration of external fairness-aware mechanisms often results in a degradation of overall performance rather than an improvement [19]. This occurs because standard condensation algorithms are designed to maximize utility by matching gradients or distributions, which inadvertently discards the subtle structural dependencies and feature correlations essential for fairness constraints. Since the condensation process lacks the inherent capability to carry fairness properties through the data reduction phase, the resulting synthetic nodes become biased proxies that downstream fair learners cannot rectify. These observations emphasize the urgent need for a unified framework that can integrate fairness as a core objective of the condensation process itself, ensuring that the compressed graph is both efficient and socially equitable.

To bridge this critical gap, we propose FairGC, a novel framework that first utilizes a Distribution-Preserving Condensation module to synchronize label and sensitive attribute distributions, effectively preventing bias propagation during the data reduction phase. Subsequently, it incorporates a Spectral Encoding module that leverages Laplacian eigen-decomposition and sinusoidal mapping to capture multi-scale topological footprints within the synthetic representation. These distilled spatial and spectral insights are then integrated through a Fairness-Enhanced Neural Architecture that employs multi-domain fusion and a label-smoothing curriculum to ensure equitable downstream predictions. Extensive experimental evaluations on diverse real-world datasets demonstrate that FairGC consistently achieves a superior balance between predictive utility and demographic fairness, significantly outperforming existing competitive GC paradigms. Our main contributions can be summarized as follows:

- We conduct an empirical study revealing that current GC methods amplify demographic bias and “crystallize” it within synthetic nodes, making it difficult for downstream models to fix. Our analysis proves that existing fairness-aware learners fail to rectify these disparities once the data reduction process is complete.
- We introduce FairGC, a unified framework that embeds fairness directly into the distillation pipeline through distribution-preserving condensation and spectral encoding. This design ensures synthetic graphs are both topologically accurate and socially equitable, balancing model utility with demographic fairness.
- Extensive evaluations on real-world datasets like Credit and Pokec demonstrate that FairGC outperforms state-of-the-art baselines across various sensitive attributes. The results confirm our model achieves a superior trade-off between classification accuracy and fairness metrics for responsible graph learning.

II. RELATED WORK

A. Graph Condensation

GC emerges as a pivotal technique to distill large-scale graph structures into a significantly smaller set of synthetic nodes, ensuring that models trained on this compact proxy achieve comparable performance to those trained on the full dataset [20]–[22]. GDCK [16] accelerates graph distillation by leveraging Neural Tangent Kernels and kernel ridge regression, bypassing the need for expensive GNN training and nested optimization loops. GCGP [15] leverages Gaussian Processes to estimate predictive posterior distributions, eliminating iterative GNN training and enabling efficient gradient-based optimization of discrete graph structures via Concrete relaxation. Traditional frameworks often maximize accuracy at the cost of exacerbating source graph biases, making their synthetic surrogates unsuitable for sensitive tasks. FairGC addresses this by coupling fairness with graph condensation.

B. Fairness-aware GNNs

Mitigating algorithmic bias in GNNs has become a crucial research frontier to prevent discriminatory protected groups [23]. Pre-processing techniques [24], [25], aim to rectify data-level bias by generating de-biased graph augmentations. In contrast, in-processing methods integrate fairness constraints directly into the training objective. A seminal FairGNN, which leverages an adversarial debiasing framework to learn node representations that are invariant to sensitive attributes, even when such attributes are only partially observed. While other recent variants like FairSIN [26] and FMP [27] introduce sophisticated message-passing constraints, they are typically designed for static, full-scale graphs. These models often lack robustness for structural reduction, revealing a gap where fairness and data reduction are treated as isolated tasks. FairGC fills this void by intrinsically coupling graph compression with fairness preservation.

III. PRELIMINARIES

A. Notations

Let $\mathcal{G} = (\mathcal{V}, \mathcal{E}, \mathbf{X}, \mathbf{y}, \mathbf{s})$ denote the original graph, where $\mathcal{V}$ is the set of n nodes, $\mathcal{E}$ is the edge set, $\mathbf{X} \in \mathbb{R}^{n \times d}$ is the node feature matrix, $\mathbf{y} \in \{0, 1\}^n$ is the label vector, and $\mathbf{s} \in \{0, 1\}^n$ is the binary sensitive attribute vector. The condensed graph is denoted as $\mathcal{G}_{\text{syn}} = (\mathcal{V}_{\text{syn}}, \mathcal{E}_{\text{syn}}, \tilde{\mathbf{X}}, \tilde{\mathbf{y}}, \tilde{\mathbf{s}})$, where $\tilde{\mathbf{X}} \in \mathbb{R}^{n_{\text{syn}} \times d}$, $\tilde{\mathbf{y}} \in \{0, 1\}^{n_{\text{syn}}}$, $\tilde{\mathbf{s}} \in \{0, 1\}^{n_{\text{syn}}}$, and $\mathcal{E}_{\text{syn}}$ is represented by the adjacency matrix $\mathbf{A}_{\text{syn}} \in \{0, 1\}^{n_{\text{syn}} \times n_{\text{syn}}}$.

B. Fairness Evaluation Metrics

To quantitatively evaluate the fairness of our proposed FairGC, we focus on two widely recognized independence-based criteria: Statistical Parity and Equal Opportunity. In our notation, $y \in \{0, 1\}$ and $\hat{y} \in \{0, 1\}$ represent the ground-truth and predicted labels, respectively, while $s \in \{0, 1\}$ indicates the sensitive attribute (e.g., gender or age).

Statistical Parity (SP) aims for an ideal state where the model’s decision is decoupled from the sensitive group membership [10]. A lower Δ_{SP} indicates that the model is more

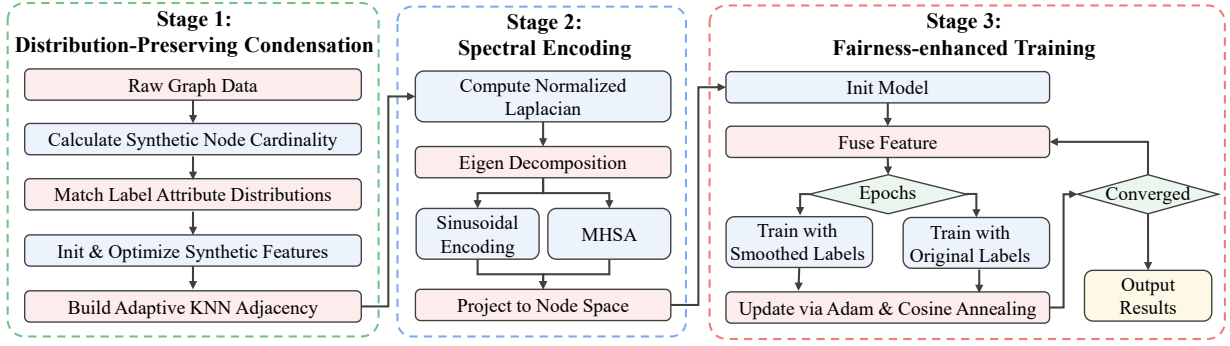


Fig. 2. The illustration of FairGC.

“group-blind” in its global outcome distribution. To measure the deviation from this equilibrium, we define it as follows:

$$\Delta_{SP} = |\mathbb{P}(\hat{y} = 1 | s = 0) - \mathbb{P}(\hat{y} = 1 | s = 1)|. \quad (1)$$

Equal Opportunity (EO) provides a more nuanced perspective by focusing on the fairness of qualified individuals [11]. Δ_{EO} ensures that the model does not disproportionately misclassify positive instances from any specific protected group. Both metrics are computed based on the empirical probabilities observed in the test set. Formally, the disparity in equal opportunity is captured by the difference in TPRs:

$$\Delta_{EO} = |\mathbb{P}(\hat{y} = 1 | y = 1, s = 0) - \mathbb{P}(\hat{y} = 1 | y = 1, s = 1)|. \quad (2)$$

C. Graph Condensation

GC aims to distill a massive source graph $\mathcal{G}$ into a compact, synthetic counterpart $\mathcal{G}_{syn}$ that retains equivalent training utility. Distinct from graph sparsification or coreset selection, which are constrained to a subset of existing nodes, GC generates entirely new feature matrices and connectivity patterns. This paradigm is typically formulated as an optimization problem intended to minimize the performance gap between models trained on original and synthetic data:

$$\min_{\mathcal{G}_{syn}} \mathcal{D}(\mathcal{M}(\mathcal{G}; \theta), \mathcal{M}(\mathcal{G}_{syn}; \tilde{\theta})), \quad (3)$$

where $\mathcal{M}(\cdot)$ denotes a GNN-based learner, while θ and $\tilde{\theta}$ represent the parameters optimized on $\mathcal{G}$ and $\mathcal{G}_{syn}$, respectively.

IV. METHODS

A. Distribution-Preserving Condensation Module

To alleviate computational overhead while upholding fairness constraints, we condense graph $\mathcal{G} = (\mathbf{X}, \mathbf{y}, \mathbf{s}, \mathcal{E})$ into a synthetic surrogate $\mathcal{G}_{syn} = (\tilde{\mathbf{X}}, \tilde{\mathbf{y}}, \tilde{\mathbf{s}}, \mathbf{A}_{syn})$. Given a budget ratio $\rho \in (0, 1)$, the synthetic cardinality is defined as:

$$n_{syn} = \max(10, \lfloor \rho \cdot n \rfloor). \quad (4)$$

1) *Multi-faceted Distribution Synchronization*: To prevent bias propagation during condensation, we synchronize the marginal distributions of labels and sensitive attributes. Let p_c and q_a denote the empirical proportions of class c and sensitive group a in the training set. Node allocation is governed by:

$$n_{c,syn} = n_{syn} \times p_c, \quad n_{a,syn} = n_{syn} \times q_a. \quad (5)$$

This ensures that $\tilde{\mathbf{y}}$ and $\tilde{\mathbf{s}}$ replicate the source statistical properties:

$$\mathbb{P}(\tilde{y} = c) = \mathbb{P}(y = c), \quad \mathbb{P}(\tilde{s} = a) = \mathbb{P}(s = a). \quad (6)$$

For irregular sensitive attributes, randomized assignment is utilized to maintain distributional equilibrium.

2) *Feature Distillation via Proxy Optimization*: Synthetic features $\tilde{\mathbf{X}}$ are initialized via stratified sampling and stabilized through Z-score normalization:

$$\tilde{\mathbf{X}}^{(0)} = \frac{\mathbf{X}_{sampled} - \mu(\mathbf{X}_{sampled})}{\sigma(\mathbf{X}_{sampled}) + \epsilon}, \quad (7)$$

where $\epsilon = 10^{-8}$. We optimize $\tilde{\mathbf{X}}$ as learnable parameters by minimizing the empirical risk of a proxy MLP $g\phi$:

$$\mathcal{L}_{cond} = -\frac{1}{n_{syn}} \sum_{i=1}^{n_{syn}} \log[g_\phi(\tilde{\mathbf{x}}_i)] \tilde{y}_i. \quad (8)$$

The optimization employs the Adam optimizer with gradient clipping to refine task-relevant utility.

3) *Hybrid Structural Reconstruction*: The adjacency matrix $\mathbf{A}_{syn}$ is reconstructed adaptively. For large-scale settings ($n_{syn} > 20,000$), we implement a sparse k -nearest neighbor (kNN) graph based on cosine similarity to ensure memory efficiency. Conversely, smaller graphs adopt a denser connectivity strategy. This dual-track approach balances topological fidelity with scalability.

B. Spectral Encoding Module

To preserve multi-scale structural characteristics, we propose a spectral encoding mechanism that captures global topological patterns. Given $\mathbf{A}_{syn}$, we derive the normalized Laplacian matrix:

$$\mathbf{L}_{syn} = \mathbf{I} - \mathbf{D}^{-\frac{1}{2}} \mathbf{A}_{syn} \mathbf{D}^{-\frac{1}{2}}, \quad (9)$$

where $\mathbf{D}$ and $\mathbf{I}$ represent the degree and identity matrices, respectively. Eigendecomposition is performed to extract the leading K eigenvalues and eigenvectors:

$$\mathbf{L}_{syn} = \mathbf{U} \mathbf{\Lambda} \mathbf{U}^\top, \quad \mathbf{\Lambda} = \text{diag}(\lambda_1, \dots, \lambda_K). \quad (10)$$

The hyperparameter K is calibrated per dataset via preliminary investigations to balance information retention and computational tractability.

Algorithm 1 Fair Graph Condensation (FairGC)

Require: Original graph $\mathcal{G} = (\mathbf{X}, \mathbf{y}, \mathbf{s}, \mathcal{E})$, compression ratio ρ , spectral components K , fusion layers L

Ensure: Compact surrogate $\mathcal{G}_{\text{syn}}$ and fairness-enhanced model f_θ

- 1: {Phase 1: Distribution-Preserving Condensation}
- 2: Calculate $n_{\text{syn}} = \max(10, \lfloor \rho \cdot n \rfloor)$;
- 3: Compute p_c, q_a from $\mathbf{y}, \mathbf{s}$;
- 4: Allocate $\tilde{\mathbf{y}}, \tilde{\mathbf{s}}$ via $n_{c,\text{syn}} = n_{\text{syn}} \times p_c, \quad n_{a,\text{syn}} = n_{\text{syn}} \times q_a$;
- 5: Initialize $\tilde{\mathbf{X}}^{(0)} = \frac{\mathbf{X}_{\text{sampled}} - \mu(\mathbf{X}_{\text{sampled}})}{\sigma(\mathbf{X}_{\text{sampled}}) + \epsilon}$, where $\epsilon = 10^{-8}$;
- 6: **while** not converged **do**
- 7: Compute $\mathcal{L}_{\text{cond}} = -\frac{1}{n_{\text{syn}}} \sum_{i=1}^{n_{\text{syn}}} \log[g_\phi(\tilde{\mathbf{x}}_i)] \tilde{y}_i$;
- 8: Update $\tilde{\mathbf{X}}$ via Adam with gradient clipping;
- 9: **end while**
- 10: Construct $\mathbf{A}_{\text{syn}}$ via adaptive kNN (sparse if $n_{\text{syn}} > 20,000$, else dense);
- 11: {Phase 2: Spectral Encoding}
- 12: Compute $\mathbf{L}_{\text{syn}} = \mathbf{I} - \mathbf{D}^{-\frac{1}{2}} \mathbf{A}_{\text{syn}} \mathbf{D}^{-\frac{1}{2}}$;
- 13: Eigendecompose $\mathbf{L}_{\text{syn}}$ to get $\{\lambda_i, \mathbf{u}_i\}_{i=1}^K$;
- 14: Generate $\mathbf{E}^{(0)}$ via sinusoidal encoding of λ_i ;
- 15: Refine $\mathbf{E} = \text{LN}(\mathbf{E}' + \text{FFN}(\text{LN}(\mathbf{E}^{(0)} + \text{MHSA}(\mathbf{E}^{(0)}))))$;
- 16: Project: $\mathbf{z}_v^{\text{spec}} = \sum_{i=1}^K \mathbf{U}_{vi} \mathbf{E}_i$;
- 17: {Phase 3: Fairness-Enhanced Training}
- 18: Initialize f_θ with FULayers, $\epsilon = 0.1$;
- 19: **for** $t = 1$ to T **do**
- 20: Fuse $\mathbf{H}$ and $\mathbf{Z}^{\text{spec}}$;
- 21: $\mathcal{L} = \begin{cases} \text{Smoothed NLL} & \text{if } t \leq 40 \\ \text{Standard NLL} & \text{otherwise} \end{cases}$
- 22: where $\tilde{\mathbf{y}}_i = (1 - \epsilon)\mathbf{y}_i + \frac{\epsilon}{C}$;
- 23: Update θ via AdamW with cosine annealing;
- 24: Monitor $\Delta_{\text{SP}}, \Delta_{\text{EO}}$;
- 25: **end for**
- 26: **return** $\mathcal{G}_{\text{syn}} = (\tilde{\mathbf{X}}, \tilde{\mathbf{y}}, \tilde{\mathbf{s}}, \mathbf{A}_{\text{syn}}), f_\theta$.

1) *Sinusoidal Eigenvalue Encoding:* To transform discrete spectral values into a continuous space, we adopt a sinusoidal encoding scheme:

$$\begin{aligned} \text{PE}(\lambda_i, 2j) &= \sin\left(\frac{\lambda_i}{10000^{2j/d}}\right), \\ \text{PE}(\lambda_i, 2j+1) &= \cos\left(\frac{\lambda_i}{10000^{2j/d}}\right), \end{aligned} \quad (11)$$

where $j \in [0, \dots, d/2 - 1]$. This mapping encapsulates the frequency characteristics of the graph spectrum.

2) *Attention-driven Spectral Refinement:* To model frequency dependencies, eigenvalue encodings $\mathbf{E}^{(0)}$ are processed via a multi-head self-attention (MHSA) layer:

$$\mathbf{E} = \text{LN}\left(\mathbf{E}' + \text{FFN}(\text{LN}(\mathbf{E}^{(0)} + \text{MHSA}(\mathbf{E}^{(0)})))\right), \quad (12)$$

where LN and FFN denote Layer Normalization and Feed-Forward Network, respectively.

3) *Spectral Feature Projection:* Refined embeddings are projected onto the node space via the eigenvector basis:

$$\mathbf{z}_v^{\text{spec}} = \sum_{i=1}^K \mathbf{U}_{vi} \mathbf{E}_i. \quad (13)$$

This projection bridges macroscopic topology and microscopic node features, ensuring the latent space is structurally informative.

C. Fairness-Enhanced Neural Architecture

We design a specialized architecture centered on multi-domain fusion and fairness-aware regularization to leverage distilled spatial and spectral knowledge.

1) *Feature Integration and Fusion:* Initial node features are projected via a standardized encoder: $\mathbf{H}^{(0)} = \sigma(\text{BN}(\mathbf{W}_0 \mathbf{X}))$. To facilitate cross-domain interaction, we introduce Fairness-aware Unified Layers (FULayer):

$$\mathbf{H}^{(l)} = \text{Norm}(\text{Dropout}(\sigma(\mathbf{W}_1^{(l)} \mathbf{H}^{(l-1)} + \mathbf{W}_2^{(l)} \mathbf{Z}^{\text{spec}}))), \quad (14)$$

for $l = 1, \dots, L$. This iterative additive fusion ensures the final representation $\mathbf{H}^{(L)}$ incorporates global structural priors while maintaining node-level signals.

2) *Optimization and Fairness Alignment:* To balance utility and demographic alignment, we implement a Label Smoothing curriculum ($\epsilon = 0.1$) for the initial 40 epochs:

$$\tilde{\mathbf{y}}_i = (1 - \epsilon)\mathbf{y}_i + \frac{\epsilon}{C}, \quad (15)$$

This soft-target approach encourages the learning of generalized, non-discriminatory decision boundaries. Subsequent training utilizes standard negative log-likelihood (NLL). Final predictions $\hat{\mathbf{y}}$ are generated via $\hat{\mathbf{y}} = \text{Softmax}(\text{MLP}(\mathbf{H}^{(L)}))$. We employ the AdamW optimizer with cosine annealing and monitor metrics to ensure demographic unbiasedness.

V. EXPERIMENTS

A. Datasets

We evaluate the performance of FairGC using four diverse real-world graphs: Credit [28], Pokec-z, Pokec-n [29], and AMiner-L [30]. Credit is a financial transaction network where edges indicate spending similarities between users; binarized Age is utilized as the sensitive attribute for default risk prediction. Pokec-z and Pokec-n are social friendship networks from Slovakian regions, Zilinsky and Nitriansky. Following common fairness benchmarks, we predict users' employment categories while considering Gender (Pokec-z-G and Pokec-n-G) and Region (Pokec-z-R and Pokec-n-R) as sensitive features. AMiner-L is a large-scale co-authorship network where the task is to predict research fields, using Affiliation Continent as the protected attribute. All experiments were conducted on a workstation equipped with a server equipped with 3x NVIDIA L40 GPUs with the driver version 580.95.05.

B. Baselines

In our experiments, we compare FairGC against a comprehensive set of SOTA baselines, which are categorized into two groups: 1) GC methods that focus on distilling large-scale graph structures into a significantly smaller set of synthetic nodes, and 2) fairness-aware GNNs that aim to mitigate bias in graph learning. For the fairness-aware GNNs, since they are not originally designed for GC, we adapt each model using the different GC framework, to enable a fair comparison under condensation scenarios, ensuring consistency in handling distilling requests.

TABLE I

EXPERIMENTAL RESULTS ON DIFFERENT DATASETS. $\uparrow$ DENOTES THE LARGER, THE BETTER; $\downarrow$ DENOTES THE OPPOSITE. THE BEST RESULTS ARE **RED AND BOLD-FACED**. THE RUNNER-UPS ARE BLUE AND UNDERLINED. OOM MEANS OUT-OF-MEMORY.

Methods	Pokey-n-G (r: 0.05)			Pokey-n-R (r: 0.05)			Pokey-z-G (r: 0.05)		
	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$
GCond (ICLR '22)	69.36 $\pm$ 0.94	6.01 $\pm$ 2.39	7.59 $\pm$ 3.33	69.32 $\pm$ 0.81	3.81 $\pm$ 0.52	6.57 $\pm$ 1.01	63.21 $\pm$ 1.23	5.40 $\pm$ 1.05	5.90 $\pm$ 1.64
DosCond (KDD '22)	68.83 $\pm$ 1.19	5.14 $\pm$ 0.97	8.91 $\pm$ 2.27	68.54 $\pm$ 1.62	6.14 $\pm$ 0.64	6.94 $\pm$ 2.60	64.30 $\pm$ 0.69	5.87 $\pm$ 1.29	5.60 $\pm$ 1.14
SFGC (NeurIPS '23)	<u>70.45$\pm$1.24</u>	5.01 $\pm$ 0.54	9.02 $\pm$ 2.01	69.82 $\pm$ 1.58	3.61 $\pm$ 0.86	1.83 $\pm$ 0.67	63.46 $\pm$ 1.57	<u>1.47$\pm$0.24</u>	<u>2.84$\pm$1.46</u>
GCGP (ArXiv '25)	<u>70.45$\pm$0.95</u>	6.96 $\pm$ 0.98	<u>4.37$\pm$0.43</u>	<u>70.39$\pm$2.03</u>	4.68 $\pm$ 0.08	3.28 $\pm$ 0.36	64.40 $\pm$ 0.49	10.86 $\pm$ 0.31	8.35 $\pm$ 0.25
GDCK (PAKDD '25)	70.13 $\pm$ 0.54	<u>3.01$\pm$0.59</u>	5.61 $\pm$ 0.83	70.31 $\pm$ 0.68	2.52 $\pm$ 0.57	<u>2.97$\pm$0.86</u>	<u>64.48$\pm$0.34</u>	2.47 $\pm$ 0.13	6.30 $\pm$ 0.17
BiMSGC (AAAI '25)	69.53 $\pm$ 1.00	7.85 $\pm$ 1.72	5.32 $\pm$ 2.01	69.56 $\pm$ 1.25	<u>1.49$\pm$0.40</u>	6.16 $\pm$ 2.10	63.48 $\pm$ 1.24	1.53 $\pm$ 0.14	4.86 $\pm$ 1.29
GCond _{+FairGNN}	69.85 $\pm$ 0.56	5.96 $\pm$ 1.89	6.41 $\pm$ 3.27	69.75 $\pm$ 0.47	4.73 $\pm$ 0.81	6.19 $\pm$ 2.07	63.13 $\pm$ 0.64	6.72 $\pm$ 0.77	7.58 $\pm$ 1.24
DosCond _{+FairGNN}	69.41 $\pm$ 0.18	6.21 $\pm$ 0.37	10.22 $\pm$ 2.15	69.15 $\pm$ 0.58	4.44 $\pm$ 1.51	3.26 $\pm$ 0.85	64.15 $\pm$ 0.67	7.73 $\pm$ 4.21	6.08 $\pm$ 5.11
FairGC	70.87$\pm$0.28	0.40$\pm$0.23	0.42$\pm$0.11	70.47$\pm$0.14	0.58$\pm$0.08	0.84$\pm$0.13	64.78$\pm$0.28	0.82$\pm$0.62	0.72$\pm$0.51

Methods	Pokey-z-R (r: 0.05)			Credit (r: 0.01)			AMiner-L (r: 0.05)		
	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$	ACC(%) $\uparrow$	$\Delta_{SP}(\%) \downarrow$	$\Delta_{EO}(\%) \downarrow$
GCond (ICLR '22)	62.51 $\pm$ 1.41	3.35 $\pm$ 1.35	6.65 $\pm$ 0.88	77.77 $\pm$ 0.05	0.59 $\pm$ 0.27	0.73 $\pm$ 0.28	89.29 $\pm$ 0.28	8.06 $\pm$ 3.63	9.28 $\pm$ 1.09
DosCond (KDD '22)	64.34 $\pm$ 0.66	4.12 $\pm$ 1.61	4.94 $\pm$ 2.19	77.74 $\pm$ 0.05	0.96 $\pm$ 1.14	1.11 $\pm$ 1.25	90.87$\pm$0.24	10.80 $\pm$ 1.12	8.46 $\pm$ 2.06
SFGC (NeurIPS '23)	<u>64.35$\pm$0.86</u>	2.99 $\pm$ 1.53	10.58 $\pm$ 1.80	77.14 $\pm$ 1.26	6.14 $\pm$ 1.35	4.87 $\pm$ 1.32	90.30 $\pm$ 1.44	6.68 $\pm$ 1.32	8.35 $\pm$ 1.44
GCGP (ArXiv '25)	<u>64.35$\pm$1.12</u>	10.86 $\pm$ 1.79	8.35 $\pm$ 1.24	77.82 $\pm$ 0.76	6.82 $\pm$ 1.65	3.92 $\pm$ 0.04	OOM	OOM	OOM
GDCK (PAKDD '25)	64.07 $\pm$ 0.34	<u>0.86$\pm$0.53</u>	<u>2.57$\pm$0.77</u>	76.94 $\pm$ 2.16	6.03 $\pm$ 0.57	9.32 $\pm$ 1.60	OOM	OOM	OOM
BiMSGC (AAAI '25)	63.74 $\pm$ 0.76	1.91 $\pm$ 0.02	3.48 $\pm$ 1.02	78.13$\pm$0.06	<u>0.56$\pm$0.02</u>	1.61 $\pm$ 0.14	89.97 $\pm$ 1.43	8.58 $\pm$ 2.19	6.94 $\pm$ 2.15
GCond _{+FairGNN}	63.21 $\pm$ 0.80	3.46 $\pm$ 2.68	8.39 $\pm$ 3.12	77.06 $\pm$ 0.63	0.87 $\pm$ 0.92	<u>1.12$\pm$1.36</u>	86.21 $\pm$ 3.54	6.18 $\pm$ 6.98	9.40 $\pm$ 5.02
DosCond _{+FairGNN}	64.09 $\pm$ 0.35	3.88 $\pm$ 2.64	4.90 $\pm$ 2.76	77.04 $\pm$ 0.38	3.24 $\pm$ 2.93	3.50 $\pm$ 3.69	90.66 $\pm$ 0.09	9.48 $\pm$ 1.27	5.88 $\pm$ 2.31
FairGC	64.62$\pm$0.30	0.10$\pm$0.09	0.29$\pm$0.28	<u>77.94$\pm$0.10</u>	0.20$\pm$0.11	0.22$\pm$0.06	<u>90.80$\pm$0.33</u>	1.06$\pm$0.71	3.28$\pm$1.26

We list six GC baselines: GCond [12], DosCond [13], SFGC [14], GCGP [15], GDCK [16], BiMSGC [17]. We use one fairness-aware GNN baseline with GC adaptation (some of them could not suit for fairness-aware controlling): FairGNN [18]. These baselines represent the current SOTA in GC and fairness-aware graph learning. By adapting the FairGNN with GC capability, we ensure a comprehensive and fair evaluation of FairGC against methods that address both fairness and condensation challenges.

C. Comparison Results

In this section, we present a comprehensive evaluation of FairGC against state-of-the-art GC baselines across several real-world datasets: Pokey-n, Pokey-z, Credit, and AMiner-L. The comparison focuses on the balance between predictive accuracy (ACC) and fairness metrics (Δ_{SP} and Δ_{EO}). The results, summarized in Table I, demonstrate that FairGC consistently achieves a superior trade-off between utility and fairness compared to all existing methods.

Predictive Accuracy. FairGC maintains highly competitive node classification accuracy across all datasets, frequently achieving the best or second-best performance. For instance, on Pokey-n-G and Pokey-z-G, FairGC achieves the highest accuracy (70.87% and 64.78%, respectively), outperforming advanced baselines like BiMSGC and GCond. On the Credit and AMiner-L datasets, while FairGC's accuracy is marginally lower than the top-performing utility-centric methods (e.g., 90.80% vs. 90.87% for DosCond on AMiner-L), this minor difference is negligible. This indicates that our fairness-aware condensation process preserves the essential informative features of the original graph without significant degradation in predictive power.

Fairness Performance. FairGC significantly enhances fairness by dramatically reducing both statistical parity and equal opportunity disparities. Across all evaluated datasets, FairGC achieves the lowest Δ_{SP} and Δ_{EO} values, often by a substantial margin. For example, on Pokey-n-G, FairGC reduces Δ_{EO} to 0.42%, a nearly 18x improvement over the standard GCond (7.59%). Notably, even when baselines are augmented with fairness-aware backbones (e.g., GCond_{+FairGNN}), they still struggle to match FairGC's performance. On the Credit dataset, FairGC achieves a Δ_{EO} of 0.22%, while most baselines fluctuate between 0.5% and 1.1%. These results underscore that FairGC effectively eliminates discriminatory patterns that traditional condensation methods often inadvertently preserve or amplify.

Collectively, FairGC emerges as a robust solution that excels in both accuracy and fairness. Unlike utility-focused methods (e.g., DosCond, SFGC) that achieve high accuracy at the cost of high bias, or simple fairness-augmented baselines that fail to achieve optimal equity, FairGC provides a holistic framework. It effectively addresses the challenges of preserving distribution-level fairness during graph reduction. The consistent gains across diverse datasets and condensation ratios validate FairGC's scalability and its practical value for socially sensitive applications where ethical considerations are as critical as model performance.

D. Ablation Study

To evaluate FairGC's components, we conduct an ablation study focusing on distribution-preserving condensation (FairGC w/o C) and fairness-enhanced architecture (FairGC w/o F), with results in Table II. The FairGC w/o F variant consistently exhibits the poorest fairness, as Δ_{SP} and

TABLE II

ABLATION STUDY RESULTS FOR FAIRGC. “P” AND “AM” DENOTE THE POKEC AND AMINER DATASETS, RESPECTIVELY.

Dataset	Metric	FairGC	FairGC w/o C	FairGC w/o F
P-n-G	ACC (%)	70.87$\pm$0.28	70.17 $\pm$ 0.05	69.36 $\pm$ 0.94
	Δ_{SP} (%)	0.40$\pm$0.23	0.48 $\pm$ 0.26	6.01 $\pm$ 2.39
	Δ_{EO} (%)	0.42$\pm$0.11	1.05 $\pm$ 0.85	7.59 $\pm$ 3.33
P-n-R	ACC (%)	70.47$\pm$0.14	70.23 $\pm$ 1.06	69.32 $\pm$ 0.81
	Δ_{SP} (%)	0.58$\pm$0.08	0.67 $\pm$ 0.25	3.81 $\pm$ 0.52
	Δ_{EO} (%)	0.84$\pm$0.13	1.06 $\pm$ 0.66	6.57 $\pm$ 1.01
P-z-G	ACC (%)	64.78$\pm$0.28	63.48 $\pm$ 0.97	63.21 $\pm$ 1.23
	Δ_{SP} (%)	0.82$\pm$0.62	2.23 $\pm$ 1.18	5.90 $\pm$ 1.64
	Δ_{EO} (%)	0.72$\pm$0.51	2.23 $\pm$ 1.18	5.90 $\pm$ 1.64
P-z-R	ACC (%)	64.62$\pm$0.30	63.44 $\pm$ 1.02	62.51 $\pm$ 1.41
	Δ_{SP} (%)	0.10$\pm$0.09	1.86 $\pm$ 1.03	3.35 $\pm$ 1.35
	Δ_{EO} (%)	0.29$\pm$0.28	4.23 $\pm$ 2.17	6.65 $\pm$ 0.88
Credit	ACC (%)	77.94$\pm$0.10	77.54 $\pm$ 0.14	77.77 $\pm$ 0.05
	Δ_{SP} (%)	0.20$\pm$0.11	0.42 $\pm$ 0.33	0.59 $\pm$ 0.27
	Δ_{EO} (%)	0.22$\pm$0.06	0.45 $\pm$ 0.30	0.73 $\pm$ 0.28
AM-L	ACC (%)	90.80$\pm$0.33	88.71 $\pm$ 0.58	89.29 $\pm$ 0.11
	Δ_{SP} (%)	1.06$\pm$0.71	1.56 $\pm$ 1.08	8.06 $\pm$ 3.63
	Δ_{EO} (%)	3.28$\pm$1.26	4.08 $\pm$ 2.15	9.28 $\pm$ 1.09

Δ_{EO} increase significantly across all datasets without explicit fairness-specific design. Similarly, the FairGC w/o C variant underscores the necessity of our condensation strategy; while traditional methods reduce graph size, they fail to preserve the nuanced statistical distributions required for fair downstream classification.

In summary, the ablation study underscores that both the distribution-preserving condensation module and the fairness-enhanced neural architecture are indispensable. While replacing the fairness architecture may occasionally boost accuracy, it leads to unacceptable fairness degradation. Likewise, traditional condensation fails to preserve the nuanced distributions necessary for fair representation.

VI. CONCLUSION

In this paper, we address the critical research gap where existing GC methods prioritize predictive accuracy but inadvertently amplify demographic biases during data reduction. We propose FairGC, a novel framework that integrates fairness as a core objective by combining distribution-preserving condensation, spectral encoding for structural fidelity, and a fairness-enhanced neural architecture. Through extensive empirical analysis on real-world datasets like Credit and Pokec-n, we demonstrate that FairGC effectively decouples predictive signals from sensitive attributes, preventing the crystallization of bias in synthetic nodes. Our results confirm that FairGC achieves a superior balance between training utility and demographic equity, providing a robust and socially responsible solution for deploying graph intelligence in resource-constrained environments.

ACKNOWLEDGMENTS

This work was supported by the National Natural Science Foundation of China (NSFC) under Grant No.62406054.

REFERENCES

- [1] F. Xia, C. Peng, J. Ren *et al.*, “Graph learning,” *Foundations and Trends in Signal Processing*, pp. 362–519, 2025.
- [2] X. Du, J. Li, D. Cheng *et al.*, “Telling peer direct effects from indirect effects in observational network data,” in *ICML*, 2025.
- [3] R. Luo, H. Huang, L. Ivan *et al.*, “FairGP: A scalable and fair graph transformer using graph partitioning,” in *AAAI*, 2025.
- [4] X. Gao, H. Yin, T. Chen *et al.*, “RobGC: Towards robust graph condensation,” *TKDE*, pp. 4791–4804, 2023.
- [5] Q. Zhang, Y. Sun, C. Jin *et al.*, “Effective policy learning for multi-agent online coordination beyond submodular objectives,” in *NeurIPS*, 2025.
- [6] R. Luo, T. Tang, F. Xia *et al.*, “Algorithmic fairness: A tolerance perspective,” *arXiv:2405.09543*, 2024.
- [7] Y. Dong, J. Ma, S. Wang *et al.*, “Fairness in graph mining: A survey,” *TKDE*, pp. 10 583–10 602, 2023.
- [8] Z. Xu, Z. Xu, J. Liu *et al.*, “Assessing classifier fairness with collider bias,” in *PAKDD*, 2022.
- [9] B. Oldfield, Z. Xu, and S. Kandanaarachchi, “Revisiting pre-processing group fairness: A modular benchmarking framework,” in *CIKM*, 2025.
- [10] C. Dwork, M. Hardt, T. Pitassi *et al.*, “Fairness through awareness,” in *ITCS*, 2012.
- [11] M. Hardt, E. Price, and N. Srebro, “Equality of opportunity in supervised learning,” in *NeurIPS*, 2016.
- [12] W. Jin, L. Zhao, S. Zhang *et al.*, “Graph condensation for graph neural networks,” in *ICLR*, 2022.
- [13] W. Jin, X. Tang, H. Jiang *et al.*, “Condensing graphs via one-step gradient matching,” in *KDD*, 2022.
- [14] X. Zheng, M. Zhang, C. Chen *et al.*, “Structure-free graph condensation: From large-scale graphs to condensed graph-free data,” in *NeurIPS*, 2023.
- [15] L. Wang and Q. Li, “Efficient graph condensation via Gaussian process,” *arXiv:2501.02565*, 2025.
- [16] Y. Zhang, Z. Chen, Q. Liu *et al.*, “GDCK: Efficient large-scale graph distillation utilizing a model-free kernelized approach,” in *PAKDD*, 2025.
- [17] X. Fu, Y. Gao, B. Yang *et al.*, “Bi-directional multi-scale graph dataset condensation via information bottleneck,” in *AAAI*, 2025.
- [18] E. Dai and S. Wang, “Learning fair graph neural networks with limited and private sensitive attribute information,” *TKDE*, pp. 7103–7117, 2023.
- [19] R. Luo, Z. Xu, X. Zhang *et al.*, “Fairness in graph learning augmented with machine learning: A survey,” *arXiv:2504.21296*, 2025.
- [20] X. Zhang, D. Song, and D. Tao, “Sparsified subgraph memory for continual graph representation learning,” in *ICDM*, 2022.
- [21] Q. Zhang, Z. Wan, Y. Yang *et al.*, “Near-optimal online learning for multi-agent submodular coordination: Tight approximation and communication efficiency,” in *ICLR*, 2025.
- [22] X. Zhang, D. Song, and D. Tao, “Ricci curvature-based graph sparsification for continual graph representation learning,” *TNNLS*, 2023.
- [23] Z. Xu, S. Kandanaarachchi, C. S. Ong *et al.*, “Fairness evaluation with item response theory,” in *WWW*, 2025.
- [24] Z. Zhang, M. Zhang, Y. Yu *et al.*, “Endowing pre-trained graph models with provable fairness,” in *WWW*, 2024.
- [25] Y. Zhu, J. Li, L. Chen *et al.*, “The devil is in the data: Learning fair graph neural networks via partial knowledge distillation,” in *WSDM*, 2024.
- [26] C. Yang, J. Liu, Y. Yan *et al.*, “FairSIN: Achieving fairness in graph neural networks through sensitive information neutralization,” in *AAAI*, 2024.
- [27] Z. Jiang, X. Han, C. Fan *et al.*, “Chasing fairness in graphs: A GNN architecture perspective,” in *AAAI*, 2024.
- [28] C. Agarwal, H. Lakkaraju, and M. Zitnik, “Towards a unified framework for fair and stable graph representation learning,” in *UAI*, 2021.
- [29] L. Takac and Z. Michal, “Data analysis in public social networks,” in *Proceedings of International Scientific Conference and International Workshop Present Day Trends of Innovations*, 2012.
- [30] J. Tang, J. Zhang, L. Yao *et al.*, “Arnetminer: Extraction and mining of academic social networks,” in *KDD*, 2008.