
Aggregation Buffer: Revisiting DropEdge with a New Parameter Block

DooHo Lee¹ Myeong Kong¹ Sagad Hamid^{1,2} Cheonwoo Lee¹ Jaemin Yoo¹

Abstract

We revisit DropEdge, a data augmentation technique for GNNs which randomly removes edges to expose diverse graph structures during training. While being a promising approach to effectively reduce overfitting on specific connections in the graph, we observe that its potential performance gain in supervised learning tasks is significantly limited. To understand why, we provide a theoretical analysis showing that the limited performance of DropEdge comes from the fundamental limitation that exists in many GNN architectures. Based on this analysis, we propose *Aggregation Buffer*, a parameter block specifically designed to improve the robustness of GNNs by addressing the limitation of DropEdge. Our method is compatible with any GNN model, and shows consistent performance improvements on multiple datasets. Moreover, our method effectively addresses well-known problems such as degree bias or structural disparity as a unifying solution. Code and datasets are available at <https://github.com/dooho00/agg-buffer>.

1. Introduction

Graph-structured data are pervasive across various research fields and real-world applications, as graphs naturally capture essential relationships among entities in complex systems. Graph neural networks (GNNs) have emerged as a powerful framework to effectively incorporate these relationships for graph-related tasks. In contrast to traditional multi-layer perceptrons (MLPs), which solely consider node features, GNNs additionally take advantage of edge information to incorporate crucial interrelations between node features (Kipf & Welling, 2017; Gasteiger et al., 2019; Hamilton et al., 2017; Veličković et al., 2018). As a consequence,

¹School of Electrical Engineering, KAIST, Daejeon, Republic of Korea ²Computer Science Department, University of Münster, Münster, Germany. Correspondence to: Jaemin Yoo <jaemin@kaist.ac.kr>.

GNNs are able to account for interaction patterns and structural dependencies, a source of knowledge that enables improving the performance in semi-supervised learning tasks, even with limited observations (Ying et al., 2018; Brody et al., 2022; Song et al., 2022).

While leveraging edge structure has proven highly effective, it often makes a GNN overfit to certain structural properties of nodes mainly observed in the training data. As a result, the model’s performance suffers considerably in the presence of *structural inconsistencies*. For example, it is widely known that GNNs perform worse on low-degree nodes than on high-degree nodes even when their features are highly informative, since high-degree nodes are the main source of information for their training (Tang et al., 2020; Liu et al., 2023; Subramonian et al., 2024). Moreover, GNNs exhibit poor accuracy on nodes whose neighbors have conflicting structural properties, such as heterophilous neighbors in homophilous graphs, or vice versa (Wang et al., 2024; Mao et al., 2024). These problems clearly highlight the two faces of GNNs—their reliance on edge structure is the key to their success, while also making them more vulnerable.

Common approaches to enhance robustness against input data variations in supervised learning are *random dropping* techniques such as DropOut (Srivastava et al., 2014). For GNNs, DropEdge (Rong et al., 2020) has been introduced as a means to increase the robustness against edge perturbations. DropEdge removes a random subset of edges at each iteration, exposing a GNN to diverse structural information. However, the performance gain by DropEdge is limited in practice, and DropEdge is typically excluded from the standard hyperparameter search space of GNNs in benchmark studies (Dwivedi et al., 2023; Luo et al., 2024).

In this work, we provide a theoretical analysis on the reason why DropEdge fails. We study the *objective shift* caused by DropEdge and highlight the implicit bias-robustness trade-off in its objective function. Then, we prove that the failure of DropEdge is not because of its algorithm, but the inductive bias existing in most GNN architectures, based on the concept of *discrepancy bound* in comparison to MLPs.

Building on these insights, we propose *Aggregation Buffer* (AGG_B), a new parameter block which can be integrated to to any trained GNN as a post-processing procedure. AGG_B effectively addresses the architectural limitation of GNNs,

allowing DropEdge to significantly enhance the robustness of GNNs compared to its original working mechanism. We demonstrate the effectiveness of AGG_B in improving the robustness and overall accuracy of GNNs across 12 node classification benchmarks. In addition, we show that AGG_B works as a unifying solution to structural inconsistencies such as degree bias and structural disparity, both of which arise from structural variations in graph datasets.

2. Preliminaries

Notation. Let $\mathcal{G} = (V, E)$ be an undirected graph, where V is the set of nodes and E is the set of edges. We denote the adjacency matrix by $\mathbf{A} \in \{0, 1\}^{|V| \times |V|}$, where $a_{ij} = 1$ if there is an edge between nodes i and j , and $a_{ij} = 0$ otherwise. The node feature matrix is denoted by $\mathbf{X} \in \mathbb{R}^{|V| \times d_0}$, where d_0 is the dimensionality of features.

Graph Neural Network. A graph neural network (GNN) consists of multiple layers, each performing two key operations: *aggregate* (AGG) and *update* (UPDATE) (Gilmer et al., 2017; Hu et al., 2020b). For each node, AGG gathers information from its neighboring nodes in the graph structure, while UPDATE combines the aggregated information with the node’s previous representation. With $\mathbf{H}^{(0)} = \mathbf{X}$, we formally define the l -th layer $\mathbf{H}^{(l)} \in \mathbb{R}^{|V| \times d_l}$ as

$$\begin{aligned} \mathbf{H}_N^{(l)} &= \text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}), \\ \mathbf{H}^{(l)} &= \text{UPDATE}^{(l)}(\mathbf{H}_N^{(l)}, \mathbf{H}^{(l-1)}), \end{aligned}$$

where d_l is the dimensionality of embeddings from the l -th layer, and a learnable weight matrix $\mathbf{W}^{(l)} \in \mathbb{R}^{d_{l-1} \times d_l}$ is typically used to transform representations between layers. We also denote $\mathbf{H}^{(s:t)} \in \mathbb{R}^{|V| \times (d_s + \dots + d_t)}$ as the concatenation of node embeddings from layer s to t along the feature dimension, as $\mathbf{H}^{(s:t)} = \mathbf{H}^{(s)} \parallel \dots \parallel \mathbf{H}^{(t)}$, where $\parallel$ denotes the concatenation operator and $s < t$.

3. Revisiting DropEdge

We give an overview of DropEdge and formalize its *objective shift* in node-level tasks. We then analyze the implicit bias-robustness trade-off in its objective and exhibit an unexpected failure of DropEdge on improving robustness.

3.1. Overview of DropEdge

DropEdge (Rong et al., 2020) is a data augmentation algorithm for improving the generalizability of GNNs by introducing stochasticity during its training. More specifically, it modifies the graph’s adjacency matrix $\mathbf{A}$ using a binary mask $\mathbf{M} \in \{0, 1\}^{|V| \times |V|}$, by creating an adjacency matrix $\tilde{\mathbf{A}} = \mathbf{M} \odot \mathbf{A}$, where $\odot$ is the element-wise multiplication between matrices. The matrix $\mathbf{M}$ is generated randomly to drop a subset of edges by setting their values to zero.



Figure 1. DropEdge generates various reduced rooted subgraphs for center nodes (*) by randomly removing edges.

In node-level tasks, a GNN can be considered as taking the k -hop subgraph of each node as its input. For each node i , the edge removal operation in DropEdge can be interpreted as transforming the rooted subgraph $\mathcal{G}_i$, centered on node i , into a reduced rooted subgraph, denoted as $\tilde{\mathcal{G}}_i$. Figure 1 illustrates how DropEdge modifies a rooted subgraph.

Definition 3.1 (Rooted Subgraph). A *rooted subgraph* $\mathcal{G}_i = (V_i, E_i)$ is a k -hop subgraph centered on node i , where V_i is the set of nodes within the k -hop neighborhood of node i and E_i denotes the set of edges between nodes in V_i .

Definition 3.2 (Reduced Rooted Subgraph). Given a rooted subgraph $\mathcal{G}_i$ as an input, a *reduced rooted subgraph* $\tilde{\mathcal{G}}_i = (\tilde{V}_i, \tilde{E}_i)$ is a subgraph of $\mathcal{G}_i$ created by DropEdge, where $\tilde{V}_i \subseteq V_i$ and $\tilde{E}_i \subseteq E_i$ is the edge set induced by $\tilde{V}_i$.

3.2. Bias-Robustness Trade-off

In a typical classification task, the objective is to minimize the Kullback-Leibler (KL) divergence between the true posterior $P(\mathbf{y}_i|\mathcal{G}_i)$ and the modeled one $Q(\mathbf{y}_i|\mathcal{G}_i)$ as

$$\mathcal{L}(\theta) = D_{\text{KL}}(P(\mathbf{y}_i|\mathcal{G}_i) \| Q(\mathbf{y}_i|\mathcal{G}_i)), \quad (1)$$

where θ is the set of parameters used for modeling Q . When DropEdge is used during the training of a GNN, it perturbs the given rooted subgraph and creates a reduced subgraph $\tilde{\mathcal{G}}_i$ which leads to the following shifted objective function:

$$\tilde{\mathcal{L}}(\theta) = D_{\text{KL}}(P(\mathbf{y}_i|\mathcal{G}_i) \| Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)). \quad (2)$$

The shifted objective $\tilde{\mathcal{L}}$ can be decomposed as follows:

$$\begin{aligned} \tilde{\mathcal{L}}(\theta) &= D_{\text{KL}}(P(\mathbf{y}_i|\mathcal{G}_i) \| Q(\mathbf{y}_i|\mathcal{G}_i)) + \\ &\quad \mathbb{E}_P[\log Q(\mathbf{y}_i|\mathcal{G}_i) - \log Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)]. \end{aligned} \quad (3)$$

The first term corresponds to the standard objective function $\mathcal{L}(\theta)$ in Equation 1, which ensures that optimizing $\tilde{\mathcal{L}}(\theta)$ remains aligned with its intended purpose. It particularly measures how well the model approximates the true posterior and can be referred to as *bias*, as it relies on observed labels collected to represent the true distribution.

The second term measures the expected difference between $\log Q(\mathbf{y}_i|\mathcal{G}_i)$ and $\log Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)$. This can be understood as measuring the *robustness* of a GNN against edge perturbations, as it is minimized s.t. the GNN produces consistent

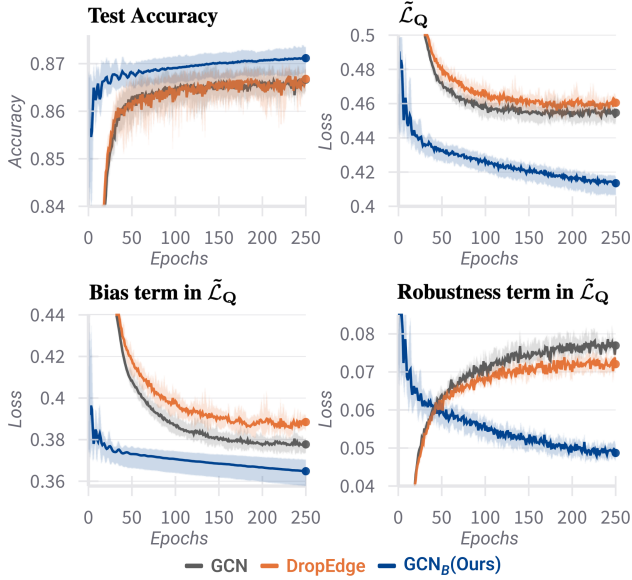


Figure 2. Accuracy and loss terms on test data during the training of a GCN at PubMed. We illustrate the average of 10 independent runs, with shaded regions representing the minimum and maximum values. While DropEdge decreases the robustness term compared to standard GNNs, it leads to increasing the bias term, eventually resulting in similar test accuracy to standard GNNs.

predictions for $\mathcal{G}_i$ and $\tilde{\mathcal{G}}_i$. However, as the expectation involves the true distribution P , it can only be computed if P is known. By assuming that Q is sufficiently close to P , we can rewrite $\tilde{\mathcal{L}}$ as follows:

$$\tilde{\mathcal{L}}_Q(\theta) = D_{\text{KL}}(P(\mathbf{y}_i|\mathcal{G}_i)\|Q(\mathbf{y}_i|\mathcal{G}_i)) + D_{\text{KL}}(Q(\mathbf{y}_i|\mathcal{G}_i)\|Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)). \quad (4)$$

We discuss the validity of this approximation in Appendix E. The interplay between the terms in $\tilde{\mathcal{L}}_Q$ naturally introduces a *bias-robustness trade-off*. The first term, which is equal to $\mathcal{L}$, enables learning the true posterior accurately. The second term works as a regularizer, promoting consistency across different reduced rooted subgraphs. Finding an optimal balance between bias and robustness is key to maximize the performance of GNNs on unseen test graphs.

3.3. Unexpected Failure of DropEdge

DropEdge is designed to enhance the robustness of GNNs against structural perturbations. To evaluate its effectiveness, we train a GNN with and without DropEdge and measure $\tilde{\mathcal{L}}_Q$ on the test set. As shown in Figure 2, DropEdge successfully regularizes the robustness term compared to standard GNNs. However, this comes at the cost of increasing the bias term, leading to a similar total $\tilde{\mathcal{L}}_Q$ and no overall performance improvement. It is notable that we have carefully tuned the drop ratio of DropEdge for this example, suggest-

ing that other drop ratios would lead to degradation.

This behavior is consistently observed across all datasets in our experiments, raising the question of whether DropEdge can truly improve the performance. While a trade-off between bias and robustness is expected, this outcome is unusual compared to data augmentation methods in other domains (Srivastava et al., 2014; DeVries, 2017; Hou et al., 2022). In most cases, small perturbations of data do not significantly interfere with the primary learning objective, allowing robustness optimization to improve generalization. However, in GNNs trained with DropEdge, optimizing robustness immediately increases the bias term on test data, preventing sufficient robustness to be achieved.

This phenomenon highlights a fundamental challenge: the minimization of $\tilde{\mathcal{L}}_Q$, in terms of both bias and robustness, is inherently difficult to achieve within the standard training framework of GNNs, limiting the effectiveness of DropEdge and similar techniques in improving edge-robustness.

3.4. Reason of the Failure: Core Limitations of GNNs

The robustness term in $\tilde{\mathcal{L}}_Q$ can be optimized only when a GNN is able to produce similar outputs for different adjacency matrices, namely $\mathbf{A}$ and $\tilde{\mathbf{A}}$. To study the poor efficacy of DropEdge, we analyze how well a GNN can bound the difference between its outputs given different inputs, which we refer to as the *discrepancy bound*. Our key observation is that the failure of DropEdge is not rooted in its algorithm but rather in the inductive bias of GNNs, suggesting that it cannot be addressed optimally with existing GNN layers.

Definition 3.3 (Discrepancy bound). Let $\mathbf{H}_1^{(l)}$ and $\mathbf{H}_2^{(l)}$ be the outputs of the l -th layer of a network f given different inputs $\mathbf{H}_1^{(l-1)}$ and $\mathbf{H}_2^{(l-1)}$. The discrepancy bound of f at the l -th layer is a constant C , such that

$$\|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 \leq C\|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2,$$

where C is independent of the specific inputs.

As a comparison, we first study the discrepancy bound of MLPs in Lemma 3.5 and move on to GNNs. Proofs for all theoretical results in this section are in Appendix C.

Lemma 3.4. *Commonly used activation functions—ReLU, Sigmoid, and GELU—and parameterized linear transformation satisfy Lipschitz continuity.*

Lemma 3.5. *Given an MLP with activation function σ , the discrepancy bound at the l -th layer is $C = L_\sigma\|\mathbf{W}^{(l)}\|_2$ where L_σ is the Lipschitz constant of σ .*

Theorem 3.6. *In an L -layer MLP with activation function σ , the discrepancy bound at the L -th layer can be derived for every intermediate layer $l < L$ as*

$$\|\mathbf{H}_1^{(L)} - \mathbf{H}_2^{(L)}\|_2 \leq C\|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2,$$

where $C = L_\sigma^{(L-l)} \prod_{i=l+1}^L \|\mathbf{W}^{(i)}\|_2$.

Theorem 3.6 implies that reducing discrepancies in intermediate representations can minimize discrepancies in the final output, allowing parameters in each layer to effectively contribute to the model’s robustness. On the other hand, the linear discrepancy bound does not hold for GNNs. We formalize this observation in **Theorem 3.8**.

Lemma 3.7. *Commonly used aggregation functions in GNNs—regular, random walk normalized, and symmetric normalized—satisfy Lipschitz continuity.*

Theorem 3.8. *Given a graph convolutional network (GCN) with any non-linear activation function σ and different adjacency matrices $\mathbf{A}_1$ and $\mathbf{A}_2$, the discrepancy bound cannot be established as a constant C independent of the input.*

Theorem 3.9. *Under the same conditions as **Theorem 3.8**, the discrepancy of a GCN at layer l is bounded as*

$$\|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 \leq C_1 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + C_2,$$

where $C_1 = L_\sigma \|\mathbf{W}^{(l)}\|_2$, $C_2 = C_1 \|V\| \|\hat{\mathbf{A}}_1 - \hat{\mathbf{A}}_2\|_2$, and $\hat{\mathbf{A}}$ is the normalized adjacency matrices of $\mathbf{A}$.

The key difference between GNNs and MLPs arises from the AGG operation in GNN layers. While the inclusion of the AGG operation enables a GNN to utilize the graph structure, it becomes problematic when aiming for robustness under different adjacency matrices. As demonstrated in **Theorem 3.9**, discrepancies can arise purely due to differences in the adjacency matrices, as a form of C_2 , even if the pre-aggregation representations are identical. This issue ultimately hinders the optimization of the robustness term in $\tilde{\mathcal{L}}_Q$, as observed in **Section 3.3**.

4. Achieving Edge-Robustness

Our analysis in **Section 3** shows the difficulty of optimizing $\tilde{\mathcal{L}}_Q$ due to the nature of GNNs. As a solution, we propose *Aggregation Buffer* (AGG_B), a new parameter block which can be integrated into a GNN’s backbone as illustrated in **Figure 3**. AGG_B is specifically designed to refine the output of the AGG operation, mitigating discrepancies caused by variations in the graph structure introduced by DropEdge.

4.1. Aggregation Buffer: A New Parameter Block

Unlike the standard training strategy, where an augmentation function is used during training, we propose a two-step approach; given a GNN trained without DropEdge, we integrate AGG_B into each GNN layer and train AGG_B with DropEdge while freezing the pre-trained parameters. This two-step procedure provides several advantages:

1. **Practical Usability.** Our approach can be applied to any trained GNN. Separate training of AGG_B enables modular application even to already deployed models.

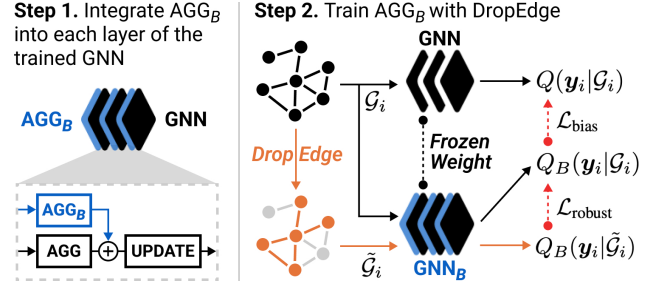


Figure 3. Illustration of AGG_B and its training scheme. After the integration into a pre-trained GNN, AGG_B is trained using $\mathcal{L}_{RC}$ with DropEdge, while the pre-trained parameters remain frozen.

2. **Effectiveness.** Pre-training without DropEdge avoids the suboptimal minimization of the bias term observed in **Section 3.3**. As a result, AGG_B can focus entirely on optimizing the robustness term.
3. **No Knowledge Loss.** Freezing the pre-trained parameters prevents any unintended loss of knowledge during the training of AGG_B . The integration of AGG_B can even be detached to get the original model back.

The main idea of our approach is to assign distinct roles to different sets of parameters: the pre-trained parameters focus on solving the primary classification task, while AGG_B is dedicated to mitigate representation changes caused by inconsistent graph structures. Given AGG_B , while its details will be discussed later, we modify a GNN layer as

$$\mathbf{H}_N^{(l)} = \text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}) + \text{AGG}_B^{(l)}(\mathbf{H}^{(0:l-1)}, \mathbf{A}),$$

where AGG_B can leverage all available resources until the current layer l , including the adjacency matrix $\mathbf{A}$ and the preceding representations $\mathbf{H}^{(0:l-1)}$. We henceforth refer to the GNN model augmented with AGG_B as GNN_B .

4.2. Essential Conditions for AGG_B

The important part of our approach is to decide the actual function of AGG_B . Existing methods for enhancing GNN layers, such as residual connections (He et al., 2016; Chen et al., 2020) and JKNet (Xu et al., 2018), are not considered as AGG_B since they fail to satisfy the essential conditions that AGG_B must meet to achieve its purpose. To derive our own approach that is better than existing methods, we first introduce the two essential conditions for AGG_B .

C1: Edge-Awareness. When the adjacency matrix $\mathbf{A}$ is perturbed to $\tilde{\mathbf{A}}$, AGG_B should produce distinct outputs to compensate for structural changes:

$$\text{AGG}_B^{(l)}(\mathbf{A}) \neq \text{AGG}_B^{(l)}(\tilde{\mathbf{A}}).$$

This condition ensures that AGG_B adapts to structural variations by modifying its output accordingly. Existing layers

that depend only on node representations, such as residual connections and JKNet, fail to meet this condition as they produce identical outputs regardless of structural perturbations when the input representations remain the same.

C2: Stability. For any perturbed adjacency matrix $\tilde{\mathbf{A}} \subset \mathbf{A}$ created by random edge dropping, AGG_B should produce outputs with a smaller deviation from the original output when given $\mathbf{A}$, compared to when given $\tilde{\mathbf{A}}$:

$$\|\text{AGG}_B^{(l)}(\mathbf{A})\|_F < \|\text{AGG}_B^{(l)}(\tilde{\mathbf{A}})\|_F.$$

This condition ensures the knowledge learned by the original GNN to be preserved, contained in the frozen pre-trained parameters, by minimizing unnecessary changes under the original graph structure. At the same time, it provides sufficient flexibility to adapt and correct for structural perturbations, thereby optimizing edge-robustness without compromising the integrity of the original representations.

Our Solution. We propose a simple structure-aware form of AGG_B which satisfies both conditions above:

$$g_B(\mathbf{H}^{(0:l-1)}, \mathbf{A}) = (\mathbf{D} + \mathbf{I})^{-1} \mathbf{H}^{(0:l-1)} \mathbf{W}^{(l)},$$

where $\mathbf{D}$ is the degree matrix of adjacency matrix $\mathbf{A}$. Since it is *degree-normalized linear transformation*, its computation is faster than the regular AGG operation. When computed in parallel, integrating AGG_B does not increase inference time, ensuring efficient execution.

Theorem 4.1. g_B satisfies the conditions C1 and C2.

Proof. The proof is in Appendix D. $\square$

4.3. Objective Function for Training AGG_B

We train the AGG_B to minimize an objective function, $\mathcal{L}_{\text{RC}}$, referred to as the *robustness-controlled loss*, which has a few adjustments from $\tilde{\mathcal{L}}_Q$. First, we introduce a hyperparameter λ to explicitly balance the strength between the bias term $\mathcal{L}_{\text{bias}}$ and the robustness term $\mathcal{L}_{\text{robust}}$:

$$\mathcal{L}_{\text{RC}}(\theta_B) = \mathcal{L}_{\text{bias}}(\theta_B) + \lambda \cdot \mathcal{L}_{\text{robust}}(\theta_B), \quad (5)$$

where θ_B refers to the set of parameters in AGG_B .

Then, we reformulate the bias term by replacing P with Q . Since our method involves two-stage training, we can safely assume that the modeled distribution Q of the pre-trained GNN is a good approximation of the true distribution P at least within the training data. As a result, the bias term can simulate knowledge distillation in training data:

$$\begin{aligned} \mathcal{L}_{\text{bias}}(\theta_B) &= \frac{1}{|V_{\text{trn}}|} \sum_{i \in V_{\text{trn}}} D_{\text{KL}}(Q(\mathbf{y}_i | \mathcal{G}_i) \| Q_B(\mathbf{y}_i | \mathcal{G}_i)), \\ \mathcal{L}_{\text{robust}}(\theta_B) &= \frac{1}{|V|} \sum_{i \in V} D_{\text{KL}}(Q_B(\mathbf{y}_i | \mathcal{G}_i) \| Q_B(\mathbf{y}_i | \tilde{\mathcal{G}}_i)), \end{aligned}$$

where V_{trn} refers to the set of (labeled) training nodes, and Q_B represents the modeled distribution of the GNN enhanced with AGG_B , which we refer to as GNN_B .

Unlike the bias term, the robustness term does not require access to the true distribution. This independence enables its application to all nodes, including unlabeled nodes, promoting comprehensive edge-robustness for the graph. On the other hand, it may not be effective to apply the bias term to all nodes as well, since it relies on an assumption that the pre-trained model distribution Q approximates P also in the unlabeled nodes, which is hardly true in practice.

5. Related Works

Random Dropping Methods for GNNs. Several random-dropping techniques were proposed for GNNs to improve their robustness, complementing the widely-used DropOut (Srivastava et al., 2014) method used in classical machine learning (You et al., 2020; Zhang et al., 2021; Li et al., 2023; Fang et al., 2023). DropEdge (Rong et al., 2020) removes a random subset of edges, while DropNode (Feng et al., 2020) removes nodes along with their connected edges. Existing graph sampling methods can also be seen as variants of these approaches. DropMessage (Fang et al., 2023) integrates DropNode, DropEdge, and DropOut by dropping propagated messages during the message-passing phase, offering higher information diversity. While these methods aim to reduce overfitting on edges in supervised learning, their performance improvements have been modest.

Sub-optimality of GNNs. Incorporating edge information for its prediction is the core idea of GNNs. However, it also makes GNNs vulnerable to structural inconsistencies in the graph, making it suffer from well-known problems like degree bias and structural disparity. *Degree bias* refers to the tendency of performing significantly better on high-degree (head) nodes than on low-degree (tail) nodes (Tang et al., 2020; Liu et al., 2023). Tail-GNN (Liu et al., 2021) transfers representation-translation from head to tail nodes, while Coldbrew (Zheng et al., 2021) uses existing nodes as virtual neighbors for tail nodes. While both approaches improve the performance on tail nodes, they degrade the performance on head nodes and rely on manual degree thresholds. TUNEUP (Hu et al., 2023) fine-tunes GNNs with pseudo-labels and DropEdge, differing from our method by not freezing pre-trained parameters, lacking AGG_B , and using a different loss function. GraphPatcher (Ju et al., 2024) attaches virtual nodes to enhance the representations of tail nodes. *Structural disparity* arises when neighboring nodes have conflicting properties, such as heterophilous nodes in homophilous graphs. Recent studies (Wang et al., 2024; Zhu et al., 2020; Mao et al., 2024) show that MLPs outperform GNNs in such scenarios, implying that avoiding edge-reliance is often more beneficial. Our work addresses both issues holistically, improving GNN generalization by enhancing edge-robustness through the idea of AGG_B .

Table 1. Accuracy (%) of all models for test nodes grouped by degree. Head nodes refer to the top 33% of nodes by degree, while tail nodes refer to the bottom 33%. **Bold** values indicate the best performance, and underlined values indicate the second-best performance. Standard deviations are shown as subscripts. Our GCN_B achieves at least the second-best in 31 out of 36 settings.

Method	Cora	Citeseer	PubMed	Wiki-CS	A.Photo	A.Computer	CS	Physics	Arxiv	Actor	Squirrel	Chameleon
OVERALL PERFORMANCE												
MLP	64.86 $\pm$ 1.21	65.55 $\pm$ 0.76	84.62 $\pm$ 0.28	75.98 $\pm$ 0.51	85.97 $\pm$ 0.81	80.81 $\pm$ 0.40	93.55$\pm$0.18	95.09 $\pm$ 0.12	56.41 $\pm$ 0.14	34.86$\pm$0.97	32.55 $\pm$ 1.51	32.10 $\pm$ 3.10
GCN	83.44 $\pm$ 1.44	72.45 $\pm$ 0.80	86.48 $\pm$ 0.17	80.26 $\pm$ 0.34	92.21 $\pm$ 1.36	88.24 $\pm$ 0.63	91.85 $\pm$ 0.29	95.18 $\pm$ 0.17	71.80 $\pm$ 0.10	30.16 $\pm$ 0.73	41.67 $\pm$ 2.42	40.19 $\pm$ 4.29
DropEdge	83.27 $\pm$ 1.55	72.29 $\pm$ 0.60	86.47 $\pm$ 0.21	80.22 $\pm$ 0.55	92.14 $\pm$ 1.42	88.08 $\pm$ 1.08	91.91 $\pm$ 0.16	95.13 $\pm$ 0.16	71.73 $\pm$ 0.21	29.86 $\pm$ 0.82	38.40 $\pm$ 2.57	40.51 $\pm$ 3.38
DropNode	83.65 $\pm$ 1.83	72.20 $\pm$ 0.67	86.55 $\pm$ 0.18	80.11 $\pm$ 0.61	91.89 $\pm$ 1.21	88.17 $\pm$ 0.40	91.93 $\pm$ 0.28	95.11 $\pm$ 0.16	71.72 $\pm$ 0.16	29.07 $\pm$ 0.93	38.01 $\pm$ 2.00	39.74 $\pm$ 2.79
DropMessage	83.45 $\pm$ 1.56	72.44 $\pm$ 0.76	<u>86.56$\pm$0.16</u>	80.30 $\pm$ 0.37	92.13 $\pm$ 1.56	<u>88.52$\pm$0.44</u>	92.08 $\pm$ 0.21	95.14 $\pm$ 0.18	71.93 $\pm$ 0.20	29.62 $\pm$ 1.05	38.75 $\pm$ 3.34	40.48 $\pm$ 3.07
TUNEUP	83.59 $\pm$ 1.26	<u>73.00$\pm$0.78</u>	86.43 $\pm$ 0.36	80.56 $\pm$ 0.47	92.11 $\pm$ 1.37	88.14 $\pm$ 0.95	90.89 $\pm$ 0.45	94.51 $\pm$ 0.25	71.81 $\pm$ 0.15	28.95 $\pm$ 1.48	41.49 $\pm$ 2.65	40.24 $\pm$ 4.24
GraphPatcher	83.57 $\pm$ 1.38	72.22 $\pm$ 0.73	86.21 $\pm$ 0.23	<u>80.64$\pm$0.51</u>	92.89$\pm$0.57	88.49 $\pm$ 0.71	91.74 $\pm$ 0.25	<u>95.25$\pm$0.24</u>	<u>72.06$\pm$0.06</u>	28.07 $\pm$ 0.67	<u>41.89$\pm$2.49</u>	40.35 $\pm$ 4.11
GCN _B (Ours)	84.84$\pm$1.39	73.32$\pm$0.85	87.56$\pm$0.27	80.75$\pm$0.42	<u>92.44$\pm$1.42</u>	88.76$\pm$0.65	<u>93.54$\pm$0.37</u>	95.79$\pm$0.17	72.43$\pm$0.16	<u>30.56$\pm$0.84</u>	42.39$\pm$2.19	40.96$\pm$4.83
ACCURACY ON HEAD NODES (HIGH-DEGREE)												
MLP	65.86 $\pm$ 1.56	70.99 $\pm$ 1.33	84.70 $\pm$ 0.32	80.06 $\pm$ 0.83	88.58 $\pm$ 1.12	86.09 $\pm$ 0.68	94.08$\pm$0.24	97.50 $\pm$ 0.14	63.93 $\pm$ 0.17	34.27$\pm$1.42	25.80 $\pm$ 3.72	29.74 $\pm$ 3.68
GCN	84.70 $\pm$ 1.60	79.10 $\pm$ 0.97	87.81 $\pm$ 0.36	85.13 $\pm$ 0.56	<u>94.85$\pm$2.01</u>	90.72 $\pm$ 0.75	93.15 $\pm$ 0.26	<u>97.64$\pm$0.12</u>	80.81 $\pm$ 0.10	27.63 $\pm$ 1.39	35.12 $\pm$ 3.80	36.51 $\pm$ 6.92
DropEdge	84.74 $\pm$ 2.01	78.92 $\pm$ 0.78	87.77 $\pm$ 0.38	84.99 $\pm$ 0.30	94.50 $\pm$ 1.75	90.10 $\pm$ 1.27	93.14 $\pm$ 0.13	97.61 $\pm$ 0.11	80.67 $\pm$ 0.26	27.51 $\pm$ 2.35	33.64 $\pm$ 4.98	37.58 $\pm$ 6.54
DropNode	84.82 $\pm$ 2.47	79.01 $\pm$ 1.34	87.80 $\pm$ 0.34	85.02 $\pm$ 0.55	91.89 $\pm$ 1.21	90.53 $\pm$ 0.58	93.19 $\pm$ 0.23	97.59 $\pm$ 0.11	80.73 $\pm$ 0.25	26.62 $\pm$ 1.15	32.33 $\pm$ 5.09	35.92 $\pm$ 6.81
DropMessage	84.86 $\pm$ 1.60	79.33 $\pm$ 1.10	<u>87.84$\pm$0.45</u>	84.96 $\pm$ 0.42	94.62 $\pm$ 2.24	<u>91.01$\pm$0.75</u>	93.28 $\pm$ 0.29	97.57 $\pm$ 0.11	80.77 $\pm$ 0.25	27.55 $\pm$ 1.70	30.42 $\pm$ 4.14	38.85$\pm$7.47
TUNEUP	84.58 $\pm$ 1.46	79.43$\pm$0.83	87.78 $\pm$ 0.54	85.35$\pm$0.51	94.73 $\pm$ 1.95	90.62 $\pm$ 1.12	92.12 $\pm$ 0.40	97.26 $\pm$ 0.15	80.74 $\pm$ 0.18	26.56 $\pm$ 1.43	34.85 $\pm$ 3.81	35.82 $\pm$ 5.38
GraphPatcher	85.21 $\pm$ 1.56	79.00 $\pm$ 0.66	87.66 $\pm$ 0.47	<u>85.22$\pm$0.65</u>	95.28$\pm$0.61	91.51$\pm$0.69	93.25 $\pm$ 0.42	97.46 $\pm$ 0.20	80.89$\pm$0.06	26.85 $\pm$ 1.38	35.72$\pm$4.41	36.40 $\pm$ 4.99
GCN _B (Ours)	85.82$\pm$1.31	<u>79.41$\pm$0.99</u>	88.14$\pm$0.60	85.04 $\pm$ 0.56	92.84 $\pm$ 2.05	90.70 $\pm$ 0.80	<u>93.87$\pm$0.26</u>	97.70$\pm$0.11	<u>80.85$\pm$0.13</u>	<u>27.65$\pm$1.48</u>	<u>35.38$\pm$4.28</u>	<u>37.68$\pm$7.39</u>
ACCURACY ON TAIL NODES (LOW-DEGREE)												
MLP	63.20 $\pm$ 1.36	60.27 $\pm$ 1.42	84.30 $\pm$ 0.43	73.02 $\pm$ 1.02	81.91 $\pm$ 0.90	75.51 $\pm$ 0.73	<u>92.96$\pm$0.28</u>	92.76 $\pm$ 0.21	49.71 $\pm$ 0.19	34.47$\pm$1.34	35.59 $\pm$ 3.33	28.94 $\pm$ 5.09
GCN	79.79 $\pm$ 1.75	65.77 $\pm$ 1.49	85.14 $\pm$ 0.25	77.83 $\pm$ 0.58	87.98 $\pm$ 0.88	83.35 $\pm$ 0.92	90.04 $\pm$ 0.53	92.74 $\pm$ 0.33	62.76 $\pm$ 0.21	<u>32.33$\pm$2.79</u>	45.85 $\pm$ 4.69	37.17 $\pm$ 6.51
DropEdge	79.61 $\pm$ 1.56	65.54 $\pm$ 1.32	85.21 $\pm$ 0.34	77.99 $\pm$ 0.55	88.13 $\pm$ 1.01	83.65 $\pm$ 1.13	90.09 $\pm$ 0.32	92.66 $\pm$ 0.36	62.65 $\pm$ 0.33	31.94 $\pm$ 1.91	43.20 $\pm$ 3.17	34.91 $\pm$ 5.93
DropNode	80.19 $\pm$ 1.63	65.50 $\pm$ 1.28	<u>85.33$\pm$0.24</u>	77.62 $\pm$ 0.67	87.69 $\pm$ 1.01	83.23 $\pm$ 0.54	90.12 $\pm$ 0.54	92.67 $\pm$ 0.34	62.69 $\pm$ 0.17	30.77 $\pm$ 1.51	42.76 $\pm$ 2.09	34.33 $\pm$ 5.88
DropMessage	79.71 $\pm$ 1.86	65.75 $\pm$ 1.42	85.31 $\pm$ 0.30	77.90 $\pm$ 0.56	88.07 $\pm$ 1.03	83.61 $\pm$ 0.52	90.35 $\pm$ 0.32	92.72 $\pm$ 0.38	63.20 $\pm$ 0.18	30.73 $\pm$ 2.05	44.44 $\pm$ 6.24	34.66 $\pm$ 6.55
TUNEUP	80.40 $\pm$ 1.77	66.35 $\pm$ 1.66	85.12 $\pm$ 0.28	78.13 $\pm$ 0.80	87.87 $\pm$ 0.97	83.45 $\pm$ 0.86	88.98 $\pm$ 0.59	91.64 $\pm$ 0.32	62.89 $\pm$ 0.19	31.09 $\pm$ 3.29	45.51 $\pm$ 4.66	37.50 $\pm$ 6.91
GraphPatcher	81.13 $\pm$ 1.91	65.39 $\pm$ 1.17	84.98 $\pm$ 0.24	<u>78.88$\pm$0.99</u>	89.28$\pm$0.66	83.24 $\pm$ 1.02	89.48 $\pm$ 0.49	93.03 $\pm$ 0.39	<u>63.56$\pm$0.13</u>	29.22 $\pm$ 1.71	46.24 $\pm$ 3.85	38.29$\pm$6.88
GCN _B (Ours)	82.05$\pm$1.75	67.17$\pm$1.37	86.85$\pm$0.22	79.25$\pm$0.58	<u>88.53$\pm$1.09</u>	84.61$\pm$0.98	92.97$\pm$0.69	94.07$\pm$0.27	64.40$\pm$0.20	32.25 $\pm$ 1.99	47.06$\pm$4.13	37.35 $\pm$ 6.99

6. Experiments

Datasets. We evaluate the accuracy of node classification for 12 widely-used benchmark graphs, including Cora, Citeseer, Pubmed, Computers, Photo, CS, Physics, Ogbn-arxiv, Actor, Squirrel and Chameleon (Shchur et al., 2018; Hu et al., 2020a; Pei et al., 2020). For Squirrel and Chameleon, we use the filtered versions following prior work (Platonov et al., 2023). These datasets are frequently used in prior works (Ju et al., 2024) and cover graphs from diverse domains with varying characteristics. Detailed descriptions of these datasets are provided in Appendix A.

Baselines. We compare GCN_B (GCN with AGG_B) with its pre-trained model GCN (Kipf & Welling, 2017) as well as closely related graph learning methods, grouped into two categories. The first category includes random-dropping data augmentation algorithms for GNNs such as DropEdge (Rong et al., 2020), DropNode (Feng et al., 2020), and DropMessage (Fang et al., 2023). The second category includes state-of-the-art methods for addressing the degree bias problem such as TUNEUP (Hu et al., 2023) and GraphPatcher (Ju et al., 2024), which are relevant since degree-robustness can be seen as a special case of edge-robustness. Lastly, we include standard MLPs, providing a baseline for full edge-robustness as no edge-information is utilized.

Setup. We adopt the public dataset splits for Ogbn-arxiv (Hu et al., 2020a), Actor, Squirrel and Chameleon (Pei et al., 2020; Platonov et al., 2023). For the remain-

ing eight datasets, we use an independently randomized 10%/10%/80% split for training, validation, and test, respectively. Our experiments are conducted using a two-layer GCN (Kipf & Welling, 2017), with hyperparameters selected via grid search based on validation accuracy across five runs, following prior works (Luo et al., 2024). For all baselines, we choose the hyperparameters as reported in the respective works or perform grid searches if no reference is available. Detailed hyperparameter settings and search spaces can be found in Appendix B.

Evaluation. All reported performances, including the ablation studies, are averaged over ten independent runs with different random seeds and splits, where we provide both the means and standard deviations. To assess edge-robustness, we evaluate the performance w.r.t. degree bias and structural disparity. For degree bias, we provide the performance on head nodes (the top 33% of nodes by degree) and tail nodes (bottom 33%), as defined in prior works (Ju et al., 2024). For structural disparity, nodes are grouped similarly based on their homophily ratio. Homophilous nodes are the top 33% with the highest homophily ratios, while heterophilous nodes comprise the bottom 33% with the lowest ratios.

6.1. Overall Performance

We compare GCN_B with several baselines and present the results in Table 1. In terms of overall performance, GCN_B achieves the highest accuracy in 9 and the second-best in 3

Table 2. Performance (%) of all models grouped by node homophily ratio. Homophilous nodes represent the top 33% of nodes with the highest homophily ratio, while heterophilous nodes represent the bottom 33%. **Bold** values indicate the best performance, and underlined values indicate the second-best performance. Standard deviations are shown as subscripts.

Method	Cora	Citeseer	PubMed	Wiki-CS	Photo	Computer	CS	Physics	Arxiv	Actor	Squirrel	Chameleon
ACCURACY ON HOMOPHILOUS NODES												
MLP	71.68 $\pm$ 1.77	76.37 $\pm$ 1.19	89.90 $\pm$ 0.58	86.30 $\pm$ 0.58	83.63 $\pm$ 1.41	86.01 $\pm$ 0.59	96.96 $\pm$ 0.25	98.02 $\pm$ 0.07	74.69 $\pm$ 0.14	36.88 $\pm$ 1.52	35.84 $\pm$ 2.73	33.50 $\pm$ 5.96
GCN	92.69 $\pm$ 1.53	87.96 $\pm$ 1.25	95.99 $\pm$ 0.22	94.05 $\pm$ 0.65	96.45 $\pm$ 3.76	94.47 $\pm$ 0.53	99.25 $\pm$ 0.15	99.32 $\pm$ 0.15	95.43 $\pm$ 0.08	39.47$\pm$1.62	48.71 $\pm$ 3.57	47.15 $\pm$ 5.79
DropEdge	92.38 $\pm$ 1.88	88.06 $\pm$ 0.90	96.12 $\pm$ 0.36	94.44 $\pm$ 0.42	96.35 $\pm$ 3.84	94.72 $\pm$ 0.43	99.28 $\pm$ 0.13	99.32 $\pm$ 0.12	95.68 $\pm$ 0.15	38.30 $\pm$ 1.08	41.25 $\pm$ 4.34	42.39 $\pm$ 4.22
DropNode	92.81 $\pm$ 1.63	87.84 $\pm$ 0.97	96.17 $\pm$ 0.32	93.99 $\pm$ 0.60	96.48 $\pm$ 3.71	94.29 $\pm$ 0.30	99.31 $\pm$ 0.17	99.29 $\pm$ 0.13	95.62 $\pm$ 0.13	38.00 $\pm$ 0.79	40.73 $\pm$ 5.13	42.67 $\pm$ 5.93
DropMessage	92.51 $\pm$ 1.67	88.06 $\pm$ 1.02	96.18$\pm$0.23	94.49 $\pm$ 0.34	96.35 $\pm$ 3.67	94.62 $\pm$ 0.47	99.37$\pm$0.15	99.32 $\pm$ 0.13	95.79$\pm$0.06	38.02 $\pm$ 1.64	45.22 $\pm$ 4.60	41.25 $\pm$ 4.98
TUNEUP	93.17 $\pm$ 1.60	<u>88.35$\pm$1.12</u>	96.04 $\pm$ 0.30	94.05 $\pm$ 0.65	96.30 $\pm$ 3.76	94.57 $\pm$ 0.55	99.14 $\pm$ 0.14	99.26 $\pm$ 0.13	95.67 $\pm$ 0.11	37.94 $\pm$ 2.57	48.58 $\pm$ 3.79	47.89$\pm$5.89
GraphPatcher	93.23 $\pm$ 1.24	87.38 $\pm$ 1.09	96.04 $\pm$ 0.28	94.08 $\pm$ 0.53	98.18$\pm$0.19	94.65 $\pm$ 0.64	98.33 $\pm$ 0.30	99.44$\pm$0.07	<u>95.72$\pm$0.09</u>	34.76 $\pm$ 1.18	48.71 $\pm$ 2.89	44.83 $\pm$ 5.22
GCN _B (Ours)	94.13$\pm$1.39	88.78$\pm$1.52	95.83 $\pm$ 0.28	94.65$\pm$0.57	<u>96.54$\pm$3.76</u>	95.30$\pm$0.49	98.64 $\pm$ 0.56	<u>99.33$\pm$0.17</u>	95.66 $\pm$ 0.13	38.26 $\pm$ 1.90	49.52$\pm$3.40	47.01 $\pm$ 5.60
ACCURACY ON HETEROPHILOUS NODES												
MLP	50.93 $\pm$ 0.98	44.88$\pm$1.82	73.22$\pm$0.71	57.90$\pm$0.93	81.71 $\pm$ 1.34	66.84 $\pm$ 0.69	85.96$\pm$0.29	89.08$\pm$0.37	34.53$\pm$0.18	31.66$\pm$2.79	32.56 $\pm$ 4.13	29.53 $\pm$ 4.83
GCN	64.18 $\pm$ 2.49	41.96 $\pm$ 1.24	67.34 $\pm$ 0.47	51.89 $\pm$ 1.08	81.74 $\pm$ 0.75	71.42 $\pm$ 1.25	76.81 $\pm$ 0.68	86.60 $\pm$ 0.37	32.51 $\pm$ 0.28	19.13 $\pm$ 1.55	42.19 $\pm$ 5.54	33.74 $\pm$ 7.61
DropEdge	64.09 $\pm$ 2.68	41.78 $\pm$ 1.27	67.12 $\pm$ 0.52	50.97 $\pm$ 1.49	81.50 $\pm$ 0.69	71.06 $\pm$ 1.95	76.92 $\pm$ 0.35	86.47 $\pm$ 0.40	31.70 $\pm$ 0.52	19.29 $\pm$ 1.72	41.59 $\pm$ 6.04	<u>37.01$\pm$5.02</u>
DropNode	<u>64.60$\pm$3.58</u>	41.59 $\pm$ 1.08	67.24 $\pm$ 0.51	51.66 $\pm$ 1.21	80.67 $\pm$ 0.97	71.38 $\pm$ 1.21	76.93 $\pm$ 0.65	86.46 $\pm$ 0.39	31.91 $\pm$ 0.57	18.93 $\pm$ 1.02	41.54 $\pm$ 4.52	36.78 $\pm$ 5.27
DropMessage	64.39 $\pm$ 2.77	41.84 $\pm$ 0.84	67.23 $\pm$ 0.39	51.48 $\pm$ 0.98	81.65 $\pm$ 0.82	<u>71.87$\pm$0.98</u>	77.27 $\pm$ 0.51	86.47 $\pm$ 0.44	32.29 $\pm$ 0.46	19.49 $\pm$ 1.18	40.79 $\pm$ 4.68	37.90$\pm$7.63
TUNEUP	63.59 $\pm$ 2.36	42.74 $\pm$ 1.07	67.09 $\pm$ 0.89	52.50 $\pm$ 0.72	81.58 $\pm$ 0.87	71.03 $\pm$ 1.17	74.16 $\pm$ 1.11	84.67 $\pm$ 0.65	31.68 $\pm$ 0.23	18.24 $\pm$ 0.92	42.13 $\pm$ 5.24	33.00 $\pm$ 6.69
GraphPatcher	64.17 $\pm$ 2.22	<u>44.47$\pm$0.89</u>	66.41 $\pm$ 0.34	<u>53.03$\pm$0.86</u>	81.46 $\pm$ 1.68	71.56 $\pm$ 1.96	78.67 $\pm$ 0.53	86.87 $\pm$ 0.64	33.38 $\pm$ 0.14	18.49 $\pm$ 1.33	<u>42.41$\pm$5.16</u>	34.45 $\pm$ 7.90
GCN _B (Ours)	65.54$\pm$2.34	43.24 $\pm$ 1.05	<u>70.77$\pm$0.71</u>	52.44 $\pm$ 1.27	82.29$\pm$1.05	72.02$\pm$1.25	<u>82.75$\pm$0.63</u>	88.43 $\pm$ 0.35	<u>34.02$\pm$0.34</u>	<u>19.96$\pm$1.49</u>	42.42$\pm$4.97	35.03 $\pm$ 7.53

out of 12 cases. Random-dropping methods such as DropEdge, DropNode, and DropMessage fail to consistently outperform GCN. This suggests that it is hard to enhance the edge-robustness of GNNs solely by performing data augmentation function, due to the inductive bias of GNNs as we have claimed in Section 3. Although GCN_B is also based on DropEdge, it consistently improves the performance of GCN since it effectively addresses its edge-vulnerability.

One notable observation is that MLPs surpass all models in the CS and Actor datasets. This indicates that edges can contribute negatively to node classification depending on the structural property. On these datasets, our GCN_B achieves the highest performance among all GNNs, demonstrating its effectiveness for enhancing edge-robustness.

TUNEUP and GraphPatcher generally improve the performance of GCN, demonstrating that addressing degree bias enhances the overall accuracy. However, their effectiveness compared to the base GCN is more limited than previously reported. Unlike previous works (Hu et al., 2023; Ju et al., 2024), which used basic hyperparameter settings, our experiments involve an extensive grid search to find optimal GCN configurations, making improvements more challenging. Despite this, GCN_B significantly outperforms the base GCN, highlighting that edge-robustness is a critical factor for performance improvements in general.

6.2. Addressing Degree Bias

We assess the performance on head and tail nodes to evaluate the impact of our method on degree bias, as shown in Table 1. GCN_B successfully mitigates degree bias, achieving at least the second-best accuracy on tail nodes in 10 and on head nodes in 9 datasets, demonstrating that enhancing edge-robustness effectively reduces degree bias. Random-

dropping approaches fail to consistently improve the tail performance. The models specifically designed for degree bias, TUNEUP and GraphPatcher, improve the tail performance but the improvement is relatively marginal.

Especially in heterophilous graphs (Actor and Squirrel), tail nodes generally outperform head nodes, contradicting typical degree bias trends. Degree-bias methods, which rely on the principle of transferring information from head to tail nodes, show limited effectiveness in these cases. However, GCN_B still improves the performance by enhancing edge-robustness without being restricted to specific information flow, showing that edge-robustness is a broader focus.

6.3. Addressing Structural Disparity

We report the accuracy of the methods on homophilous and heterophilous nodes in Table 2 to evaluate structural disparity. Consistently with recent findings (Mao et al., 2024), our experiments show that MLPs generally outperform GNNs on heterophilous nodes, while GNNs perform better on homophilous nodes. Methods for addressing degree bias, particularly GraphPatcher, show some improvements on heterophilous nodes but inconsistently, indicating a correlation between degree bias and structural disparity, although the two problems seem distinct. GCN_B achieves the highest GNN performance in heterophilous nodes on 9 datasets, demonstrating that enhancing edge-robustness can effectively mitigate structural disparity.

6.4. Generalization to Other GNN Architectures

An important advantage of AGG_B is its broad applicability to most GNN architectures due to its modular design. To evaluate its versatility, we conduct extensive experiments

Table 3. Accuracy of different GNN models before and after the integration with AGG_B. AGG_B achieves consistent and significant performance improvements across various architectures.

	Pubmed	CS	Arxiv	Chameleon
SAGE	87.07 $\pm$ 0.24	92.44 $\pm$ 0.60	70.92 $\pm$ 0.16	37.34 $\pm$ 3.56
SAGE _B	88.09$\pm$0.28	93.36$\pm$0.47	71.16$\pm$0.14	37.85$\pm$3.80
GAT	85.64 $\pm$ 0.24	90.50 $\pm$ 0.28	71.86 $\pm$ 0.14	38.54 $\pm$ 2.70
GAT _B	87.47$\pm$0.37	93.09$\pm$0.60	72.26$\pm$0.14	39.08$\pm$2.84
SGC	84.01 $\pm$ 0.76	90.89 $\pm$ 0.45	69.15 $\pm$ 0.05	38.24 $\pm$ 3.00
SGC _B	84.77$\pm$1.02	91.90$\pm$0.43	69.55$\pm$0.04	38.91$\pm$3.08
GIN	85.42 $\pm$ 0.20	87.88 $\pm$ 0.51	63.94 $\pm$ 0.53	39.84 $\pm$ 2.69
GIN _B	87.18$\pm$0.17	88.58$\pm$1.00	65.66$\pm$0.75	41.72$\pm$2.41

Table 4. Accuracy with different layer architectures used as AGG_B, with none of the alternatives outperforming our proposed design.

	Pubmed	CS	Arxiv	Chameleon
JKNet	87.45 $\pm$ 0.25	93.36 $\pm$ 0.56	72.19 $\pm$ 0.18	40.29 $\pm$ 4.68
Residual	87.46 $\pm$ 0.24	92.05 $\pm$ 0.28	72.29 $\pm$ 0.12	39.77 $\pm$ 4.57
AGG	86.82 $\pm$ 0.55	91.63 $\pm$ 0.28	72.27 $\pm$ 0.12	40.69 $\pm$ 2.69
AGG _B (ours)	87.56$\pm$0.27	93.54$\pm$0.37	72.43$\pm$0.16	40.96$\pm$4.83

on four well-known architectures: SAGE (Hamilton et al., 2017), GAT (Veličković et al., 2018), SGC (Wu et al., 2019), and GIN (Xu et al., 2019). In Table 3, we compare the accuracy of each model before and after the integration of AGG_B. These results show that AGG_B consistently delivers significant performance improvements across all architectures, demonstrating its wide applicability and effectiveness.

7. Ablation Studies

Different Layer Architectures. In line with Section 4.2, we evaluate alternative layer architectures for AGG_B, including residual connections, JKNet, and the AGG layer of GCN, while not changing other components of our method. Table 4 shows that our proposed design consistently outperforms these alternatives. Especially, the standard AGG shows no improvement on Pubmed and even degrades performance on CS from the base GCN. This suggests that using an additional AGG operation to resolve structural inconsistencies is ineffective as it introduces another inconsistency, as described in Theorem 3.9. These results highlight the importance of the conditions we propose in Section 4.2 for designing effective AGG_B. Comprehensive results across all datasets, accompanied by an in-depth discussion of this ablation study, are presented in Appendix G.

Different Loss Functions. Following Section 4.3, we explore alternative loss functions for training AGG_B. First, we evaluate a variant of $\mathcal{L}_{RC}$ by restricting the robustness term to the training set. While this approach is less effective than the proposed loss on all four datasets, it still consistently

Table 5. Accuracy with alternative loss functions to train AGG_B.

	Pubmed	CS	Arxiv	Chameleon
Pseudo-label	86.62 $\pm$ 0.37	92.48 $\pm$ 0.15	71.84 $\pm$ 0.18	40.14 $\pm$ 4.01
Self-distillation	86.15 $\pm$ 0.35	92.02 $\pm$ 0.25	72.18 $\pm$ 0.19	40.22 $\pm$ 4.22
Cross-entropy	86.67 $\pm$ 0.16	93.29 $\pm$ 0.13	71.94 $\pm$ 0.13	40.76 $\pm$ 4.19
$\mathcal{L}_{RC}$ (train only)	86.89 $\pm$ 0.33	92.67 $\pm$ 0.57	72.26 $\pm$ 0.15	40.31 $\pm$ 3.98
$\mathcal{L}_{RC}$ (ours)	87.56$\pm$0.27	93.54$\pm$0.37	72.43$\pm$0.16	40.96$\pm$4.83

Table 6. Accuracy with various architectural hyperparameters, AGG_B significantly enhances performance in all configurations.

	Pubmed		Arxiv	
	GCN	GCN _B	GCN	GCN _B
NUMBER OF LAYERS				
2	86.48 $\pm$ 0.17	87.56$\pm$0.27	71.80 $\pm$ 0.10	72.43$\pm$0.16
4	84.82 $\pm$ 0.34	87.36$\pm$0.37	71.53 $\pm$ 0.20	72.42$\pm$0.20
6	83.46 $\pm$ 0.24	86.64$\pm$0.58	70.77 $\pm$ 0.27	71.79$\pm$0.21
8	82.68 $\pm$ 0.19	86.18$\pm$0.62	70.17 $\pm$ 0.45	71.36$\pm$0.38
HIDDEN DIMENSION SIZE				
64	86.54 $\pm$ 0.26	87.18$\pm$0.32	70.12 $\pm$ 0.12	70.43$\pm$0.10
128	86.56 $\pm$ 0.25	87.36$\pm$0.25	70.92 $\pm$ 0.14	71.24$\pm$0.13
256	86.54 $\pm$ 0.17	87.54$\pm$0.23	71.39 $\pm$ 0.08	71.87$\pm$0.12
512	86.48 $\pm$ 0.17	87.56$\pm$0.27	71.80 $\pm$ 0.10	72.43$\pm$0.16
ACTIVATION FUNCTION				
ReLU	86.48 $\pm$ 0.17	87.56$\pm$0.27	71.80 $\pm$ 0.10	72.43$\pm$0.16
ELU	86.51 $\pm$ 0.19	86.95$\pm$0.26	71.50 $\pm$ 0.22	71.97$\pm$0.20
Sigmoid	85.66 $\pm$ 0.16	86.62$\pm$0.42	71.54 $\pm$ 0.14	72.05$\pm$0.20
Tanh	85.28 $\pm$ 0.19	86.12$\pm$0.17	71.72 $\pm$ 0.15	72.25$\pm$0.14

improves performance over the base GCN. Additionally, we test other loss functions, including the cross entropy using labels, knowledge distillation from the pre-trained model (Hinton, 2015; Zhang et al., 2022), and pseudo-labeling (Lee et al., 2013; Hu et al., 2023). Although these alternatives lead to performance improvements when combined with AGG_B, their effectiveness and consistency are limited compared to our objective function $\mathcal{L}_{RC}$.

Architectural Hyperparameters. For broader applicability, we expect AGG_B to enhance robustness across diverse architectural hyperparameters not only in 2-layer GNNs. In Table 6, we present experimental results with varying numbers of layers, hidden dimension sizes, and activation functions in GCN. AGG_B consistently improves performance in all configurations, even when the base model performs poorly due to overly deep layers or small hidden sizes. Notably in deep networks, where the performance typically degrades due to oversmoothing, AGG_B significantly mitigates this decline, suggesting that the lack of edge-robustness is a potential reason of oversmoothing. These results demonstrate that AGG_B is broadly applicable to GNNs regardless of their architectural hyperparameters. Further experimental results involving deeper architectures and additional datasets are presented in Appendix H.

Table 7. Effect of each key component in the proposed method.

	Pubmed	CS	Arxiv	Chameleon
GCN _B	87.56$\pm$0.27	93.54$\pm$0.37	72.43 $\pm$ 0.16	40.96$\pm$4.83
(-) Freezing	87.50 $\pm$ 0.31	93.50 $\pm$ 0.41	72.91$\pm$0.14	40.77 $\pm$ 3.86
(-) AGG _B	86.82 $\pm$ 0.60	91.67 $\pm$ 0.41	72.21 $\pm$ 0.11	40.35 $\pm$ 3.89
(-) Pre-train	87.22 $\pm$ 0.14	92.78 $\pm$ 0.17	71.15 $\pm$ 0.12	39.39 $\pm$ 3.06

Effect of Each Component. We study the effect of each key component of our approach in Table 7. First, the accuracy usually degrades without freezing the parameters of the pre-trained GNN, indicating that freezing these parameters prevents unintended loss of the original knowledge. Next, we test the performance without AGG_B, which is equivalent to fine-tuning the pre-trained GNN using $\mathcal{L}_{RC}$. This leads to a significant performance degradation, even performing worse than the pre-trained GNN on the CS dataset. This supports our theoretical findings that the original GNN architecture is inherently limited in optimizing the robustness term in $\hat{\mathcal{L}}_Q$. Finally, training GCN_B in an end-to-end manner without pre-training also leads to performance degradation, demonstrating that the two-step training approach effectively optimizes both the bias and robustness terms.

8. Conclusion

In this work, we revisited DropEdge, identifying a critical limitation—DropEdge fails to fully optimize its robustness objective during training. Our theoretical analysis revealed that this limitation arises from the inherent properties of the AGG operation in GNNs, which struggles to maintain consistent representations under structural perturbations. To address this issue, we proposed *Aggregation Buffer* (AGG_B), a new parameter block designed to improve the AGG operations of GNNs. By refining the aggregation process, AGG_B effectively optimizes the robustness objective, making the model significantly stronger to structural variations. Experiments on 12 node classification benchmarks and various GNN architectures demonstrated significant performance gains driven by AGG_B, especially for the problems related to structural inconsistencies, such as degree bias and structural disparity. Despite its effectiveness, our approach has limitations as a two-step approach; its performance relies on pre-trained knowledge, as AGG_B focuses primarily on improving robustness. A potential direction for future work is to design a framework that enables AGG_B to be trained end-to-end, allowing simultaneous optimization of both bias and robustness without dependency on pre-training.

Acknowledgements

This work was supported by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) (RS-2024-00341425 and RS-2024-00406985).

Impact Statement

In this paper, we aim to advance the field of graph neural networks. Our work is believed to improve the reliability and effectiveness of GNNs on various applications. We do not foresee any direct negative societal impacts.

References

- Brody, S., Alon, U., and Yahav, E. How attentive are graph attention networks? In *International Conference on Learning Representations*, 2022.
- Chen, M., Wei, Z., Huang, Z., Ding, B., and Li, Y. Simple and deep graph convolutional networks. In *International conference on machine learning*, pp. 1725–1735. PMLR, 2020.
- DeVries, T. Improved regularization of convolutional neural networks with cutout. *arXiv preprint arXiv:1708.04552*, 2017.
- Dwivedi, V. P., Joshi, C. K., Luu, A. T., Laurent, T., Bengio, Y., and Bresson, X. Benchmarking graph neural networks. *Journal of Machine Learning Research*, 24 (43):1–48, 2023.
- Fang, T., Xiao, Z., Wang, C., Xu, J., Yang, X., and Yang, Y. Dropmessage: Unifying random dropping for graph neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pp. 4267–4275, 2023.
- Feng, W., Zhang, J., Dong, Y., Han, Y., Luan, H., Xu, Q., Yang, Q., Kharlamov, E., and Tang, J. Graph random neural networks for semi-supervised learning on graphs. *Advances in neural information processing systems*, 33: 22092–22103, 2020.
- Gasteiger, J., Bojchevski, A., and Günnemann, S. Combining neural networks with personalized pagerank for classification on graphs. In *International Conference on Learning Representations*, 2019.
- Gilmer, J., Schoenholz, S. S., Riley, P. F., Vinyals, O., and Dahl, G. E. Neural message passing for quantum chemistry. In *International conference on machine learning*, pp. 1263–1272. PMLR, 2017.
- Hamilton, W., Ying, Z., and Leskovec, J. Inductive representation learning on large graphs. *Advances in neural information processing systems*, 30, 2017.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.

- Hinton, G. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- Hou, L., Pang, R. Y., Zhou, T., Wu, Y., Song, X., Song, X., and Zhou, D. Token dropping for efficient BERT pretraining. In Muresan, S., Nakov, P., and Villavicencio, A. (eds.), *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 3774–3784, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.262.
- Hu, W., Fey, M., Zitnik, M., Dong, Y., Ren, H., Liu, B., Catasta, M., and Leskovec, J. Open graph benchmark: Datasets for machine learning on graphs. *Advances in neural information processing systems*, 33:22118–22133, 2020a.
- Hu, W., Liu, B., Gomes, J., Zitnik, M., Liang, P., Pande, V., and Leskovec, J. Strategies for pre-training graph neural networks. In *International Conference on Learning Representations*, 2020b.
- Hu, W., Cao, K., Huang, K., Huang, E. W., Subbian, K., and Leskovec, J. Tuneup: A training strategy for improving generalization of graph neural networks, 2023. URL <https://openreview.net/forum?id=8xuFDlyCoH>.
- Ju, M., Zhao, T., Yu, W., Shah, N., and Ye, Y. Graph-patcher: mitigating degree bias for graph neural networks via test-time augmentation. *Advances in Neural Information Processing Systems*, 36, 2024.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. In Bengio, Y. and LeCun, Y. (eds.), *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015.
- Kipf, T. N. and Welling, M. Semi-supervised classification with graph convolutional networks. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017.
- Langley, P. Crafting papers on machine learning. In Langley, P. (ed.), *Proceedings of the 17th International Conference on Machine Learning (ICML 2000)*, pp. 1207–1216, Stanford, CA, 2000. Morgan Kaufmann.
- Lee, D.-H. et al. Pseudo-label: The simple and efficient semi-supervised learning method for deep neural networks. In *Workshop on challenges in representation learning, ICML*, volume 3, pp. 896. Atlanta, 2013.
- Li, J., Wu, R., Sun, W., Chen, L., Tian, S., Zhu, L., Meng, C., Zheng, Z., and Wang, W. What’s behind the mask: Understanding masked graph modeling for graph autoencoders. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pp. 1268–1279, 2023.
- Liu, Z., Nguyen, T.-K., and Fang, Y. Tail-gnn: Tail-node graph neural networks. In *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery & Data Mining*, pp. 1109–1119, 2021.
- Liu, Z., Nguyen, T.-K., and Fang, Y. On generalized degree fairness in graph neural networks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pp. 4525–4533, 2023.
- Luo, Y., Shi, L., and Wu, X.-M. Classic GNNs are strong baselines: Reassessing GNNs for node classification. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024.
- Mao, H., Chen, Z., Jin, W., Han, H., Ma, Y., Zhao, T., Shah, N., and Tang, J. Demystifying structural disparity in graph neural networks: Can one size fit all? *Advances in neural information processing systems*, 36, 2024.
- Pei, H., Wei, B., Chang, K. C.-C., Lei, Y., and Yang, B. Geom-gcn: Geometric graph convolutional networks. *arXiv preprint arXiv:2002.05287*, 2020.
- Platonov, O., Kuznedelev, D., Diskin, M., Babenko, A., and Prokhorenkova, L. A critical look at the evaluation of GNNs under heterophily: Are we really making progress? In *The Eleventh International Conference on Learning Representations*, 2023.
- Rong, Y., Huang, W., Xu, T., and Huang, J. Dropedge: Towards deep graph convolutional networks on node classification. In *International Conference on Learning Representations*, 2020.
- Shchur, O., Mumme, M., Bojchevski, A., and Günnemann, S. Pitfalls of graph neural network evaluation. *arXiv preprint arXiv:1811.05868*, 2018.
- Song, Z., Yang, X., Xu, Z., and King, I. Graph-based semi-supervised learning: A comprehensive review. *IEEE Transactions on Neural Networks and Learning Systems*, 34(11):8174–8194, 2022.
- Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., and Salakhutdinov, R. Dropout: a simple way to prevent neural networks from overfitting. *The journal of machine learning research*, 15(1):1929–1958, 2014.
- Subramonian, A., Kang, J., and Sun, Y. Theoretical and empirical insights into the origins of degree bias in graph

- neural networks. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- Tang, X., Yao, H., Sun, Y., Wang, Y., Tang, J., Aggarwal, C., Mitra, P., and Wang, S. Investigating and mitigating degree-related biases in graph convolutional networks. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*, pp. 1435–1444, 2020.
- Veličković, P., Cucurull, G., Casanova, A., Romero, A., Liò, P., and Bengio, Y. Graph attention networks. In *International Conference on Learning Representations*, 2018.
- Wang, J., Guo, Y., Yang, L., and Wang, Y. Understanding heterophily for graph neural networks. In Salakhutdinov, R., Kolter, Z., Heller, K., Weller, A., Oliver, N., Scarlett, J., and Berkenkamp, F. (eds.), *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pp. 50489–50529. PMLR, 21–27 Jul 2024.
- Wang, M., Zheng, D., Ye, Z., Gan, Q., Li, M., Song, X., Zhou, J., Ma, C., Yu, L., Gai, Y., et al. Deep graph library: A graph-centric, highly-performant package for graph neural networks. *arXiv preprint arXiv:1909.01315*, 2019.
- Wu, F., Souza, A., Zhang, T., Fifty, C., Yu, T., and Weinberger, K. Simplifying graph convolutional networks. In *International conference on machine learning*, pp. 6861–6871. PMLR, 2019.
- Xu, K., Li, C., Tian, Y., Sonobe, T., Kawarabayashi, K.-i., and Jegelka, S. Representation learning on graphs with jumping knowledge networks. In *International conference on machine learning*, pp. 5453–5462. PMLR, 2018.
- Xu, K., Hu, W., Leskovec, J., and Jegelka, S. How powerful are graph neural networks? In *International Conference on Learning Representations*, 2019.
- Ying, R., He, R., Chen, K., Eksombatchai, P., Hamilton, W. L., and Leskovec, J. Graph convolutional neural networks for web-scale recommender systems. In *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining*, pp. 974–983, 2018.
- You, Y., Chen, T., Sui, Y., Chen, T., Wang, Z., and Shen, Y. Graph contrastive learning with augmentations. *Advances in neural information processing systems*, 33:5812–5823, 2020.
- Zeng, H., Zhou, H., Srivastava, A., Kannan, R., and Prasanna, V. GraphSAINT: Graph sampling based inductive learning method. In *International Conference on Learning Representations*, 2020.
- Zhang, H., Wu, Q., Yan, J., Wipf, D., and Yu, P. S. From canonical correlation analysis to self-supervised graph neural networks. *Advances in Neural Information Processing Systems*, 34:76–89, 2021.
- Zhang, S., Liu, Y., Sun, Y., and Shah, N. Graph-less neural networks: Teaching old mlps new tricks via distillation. In *International Conference on Learning Representations*, 2022.
- Zheng, W., Huang, E. W., Rao, N., Katariya, S., Wang, Z., and Subbian, K. Cold brew: Distilling graph node representations with incomplete or missing neighborhoods. *arXiv preprint arXiv:2111.04840*, 2021.
- Zhu, J., Yan, Y., Zhao, L., Heimann, M., Akoglu, L., and Koutra, D. Beyond homophily in graph neural networks: Current limitations and effective designs. *Advances in neural information processing systems*, 33:7793–7804, 2020.

A. Dataset Overview and Training Configuration for Base GCN

We selected the 12 datasets based on prior work (Ju et al., 2024). For Squirrel and Chameleon, we use the filtered versions provided by (Platonov et al., 2023) via their public repository: <https://github.com/yandex-research/heterophilous-graphs>. The remaining 10 datasets are sourced from the Deep Graph Library (DGL) (Wang et al., 2019). All graphs are treated as undirected. Detailed statistics are provided in Table 8.

Table 8. Statistics of datasets and hyperparameters used for training the base 2-layer GCN.

	Cora	Citeseer	PubMed	Wiki-CS	Photo	Computer	CS	Physics	Arxiv	Actor	Squirrel	Chameleon
# nodes	2,708	3,327	19,717	11,701	7,650	13,752	18,333	34,493	169,343	7,600	2,334	890
# edges	10,556	9,228	88,651	431,726	238,162	491,722	163,788	495,924	1,166,243	33,391	93,996	18,598
# features	1,433	3,703	500	300	745	767	6,805	8,415	128	932	2,089	2,325
# classes	7	6	3	10	8	10	15	5	40	5	5	5
Homophily Ratio	0.8100	0.7355	0.8024	0.6543	0.8272	0.7772	0.8081	0.9314	0.6542	0.2167	0.2072	0.2361
Hidden Dim	512	512	512	512	512	512	256	64	512	64	256	256
Learning Rate	1e-2	1e-2	1e-2	1e-2	1e-2	1e-2	1e-3	1e-3	1e-2	1e-3	1e-2	1e-2
Weight Decay	5e-4	5e-4	5e-4	5e-4	5e-4	5e-4	5e-4	5e-5	5e-5	5e-4	5e-4	5e-4
Dropout	0.5	0.7	0.7	0.3	0.5	0.3	0.2	0.5	0.5	0.7	0.2	0.2
AGG Scheme	Sym	Sym	Sym	RW	Sym	Sym	Sym	Sym	RW	Sym	Sym	Sym

To train the base GCN, we conduct a grid search across five independent runs for each dataset, selecting the best hyperparameter configuration based on the highest validation accuracy, following the search space outlined in (Luo et al., 2024). The search space included hidden dimensions [64, 256, 512], dropout ratios [0.2, 0.3, 0.5, 0.7], weight decay values [0, 5e-4, 5e-5], and learning rates [1e-2, 1e-3, 5e-3]. We use the Adam optimizer (Kingma & Ba, 2015) for training with early stopping based on validation accuracy, using a patience of 100 epochs across all datasets.

We also consider two GCN aggregation schemes following prior work (Ju et al., 2024): (i) symmetric normalization, typically used in transductive settings, formulated as

$$\text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}) = \mathbf{D}^{-\frac{1}{2}} \mathbf{A} \mathbf{D}^{-\frac{1}{2}} \mathbf{H}^{(l-1)} \mathbf{W}^{(l)},$$

and (ii) random-walk normalization, commonly used in inductive settings, given by

$$\text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}) = \mathbf{D}^{-1} \mathbf{A} \mathbf{H}^{(l-1)} \mathbf{W}^{(l)}.$$

We select the aggregation scheme that achieves higher validation accuracy.

For the Squirrel and Chameleon datasets, we observe significant performance degradation when using standard GCN architectures. Therefore, guided by recommendations from (Platonov et al., 2023), which highlights data leakage issues and proposes filtered versions of these datasets, we incorporate residual connections and layer normalization into GCN. For reproducibility, detailed hyperparameter settings used for training the base GCN on each dataset are provided in Table 8.

B. Experiment Configurations for Baselines and Proposed Method

All experiments are conducted on an NVIDIA RTX A6000 GPU with 48 GB of memory. We sincerely thank all the authors of baseline methods for providing open-source implementations, which greatly facilitated reproducibility and comparison.

MLP. For MLPs, we perform a grid search using the exact same hyperparameter search space as the base GCN. This extensive search, which is often overlooked for MLPs, leads to a unique observation: well-tuned MLPs can outperform GNNs on certain datasets, such as Actor and CS.

Random Dropping Methods . For DropEdge (Rong et al., 2020), DropNode (Feng et al., 2020), and DropMessage (Fang et al., 2023), we use the official repository of DropMessage: <https://github.com/zjunet/DropMessage>, which offers a unified framework for empirical comparison of the random dropping techniques. We conduct a grid search over drop ratios from 0.1 to 1.0 in increments of 0.1 for each method.

GraphPatcher. For GraphPatcher (Ju et al., 2024), we use the official repository: <https://github.com/jumxglhf/GraphPatcher>. We adopt the provided hyperparameter settings for overlapping datasets. For the remaining datasets, we perform a hyperparameter search over five independent runs, following the search space suggested in the original paper.

Table 9. Hyperparameters used to train AGG_B integrated with a pre-trained 2-layer GCN

	Cora	Citeseer	PubMed	Wiki-CS	Photo	Computer	CS	Physics	Arxiv	Actor	Squirrel	Chameleon
λ	0.5	1.0	0.5	0.5	1.0	0.5	0.1	1.0	0.5	0.5	0.1	0.1
DropOut	0.7	0.7	0.0	0.5	0.0	0.2	0.7	0.7	0.0	0.2	0.2	0.0
DropEdge	0.5	0.2	0.2	0.2	0.5	0.5	0.2	0.2	0.5	0.5	0.7	0.7

TUNEUP. For TUNEUP (Hu et al., 2023), as the official implementation is not publicly available, we implemented the method ourselves. For the second training stage of TUNEUP, we conduct a grid search over DropEdge ratios from 0.1 to 1.0, and use the same search space for learning rate, dropout ratio, and weight decay as in the base GCN. Although TUNEUP was also manually implemented by the GraphPatcher authors, our extensively tuned implementation consistently yields higher performance across the most of datasets.

Aggregation Buffer. AGG_B is trained after being integrated into a pre-trained GNN. For training AGG_B , we use the Adam optimizer with a fixed learning rate of 1e-2 and weight decay of 0.0 across all datasets. It is noteworthy that further performance gains may be achievable by tuning these hyperparameters for each dataset individually. Since the hidden dimension and number of layers are determined by the pre-trained model, they are not tunable hyperparameters for AGG_B . Training of AGG_B is early stopped based on validation accuracy, with a patience of 100 epochs across all datasets.

AGG_B requires tuning on three key hyperparameters: the dropout ratio, DropEdge ratio, and the coefficient λ , which balances the bias and robustness terms in $\mathcal{L}_{RC}$, as described in Equation (5). The search space used for these hyperparameters in our experiments is as follows:

- λ values: [1, 0.5, 0.1],
- DropEdge ratio: [0.2, 0.5, 0.7, 1.0],
- Dropout ratio: [0, 0.2, 0.5, 0.7].

For hyperparameter tuning, we follow the same process used for training the base GCN, conducting a search across five independent runs and selecting the configuration with the highest validation accuracy. To ensure reproducibility, we provide the detailed hyperparameters for training AGG_B across datasets in Table 9, and release our implementation as open-source at <https://github.com/dooho00/agg-buffer>.

C. Proofs in Section 3

C.1. Proof of Lemma 3.4

Activation functions play a pivotal role in Graph Neural Networks (GNNs) by introducing non-linearity, which enables the network to model complex relationships within graph-structured data. Ensuring that these activation functions are Lipschitz continuous is essential for guaranteeing that similarly aggregated representation can result a similar output after applying the activation function. In this section, we formally derive the Lipschitz continuity of three widely used activation functions: Rectified Linear Unit (ReLU), Sigmoid, and Gaussian Error Linear Unit (GELU).

C.1.1. DEFINITION OF LIPSCHITZ CONTINUITY

A function $f : \mathbb{R} \rightarrow \mathbb{R}$ is said to be **Lipschitz continuous** if there exists a constant $L \geq 0$ such that for all $x, y \in \mathbb{R}$,

$$|f(x) - f(y)| \leq L|x - y|.$$

C.1.2. RECTIFIED LINEAR UNIT (ReLU)

The Rectified Linear Unit (ReLU) activation function is defined as:

$$\text{ReLU}(x) = \max(0, x).$$

To prove that ReLU is 1-Lipschitz continuous, we need to show that:

$$|\text{ReLU}(x) - \text{ReLU}(y)| \leq |x - y| \quad \forall x, y \in \mathbb{R}.$$

Case 1: $x \geq 0$ and $y \geq 0$

In this case,

$$\text{ReLU}(x) = x \quad \text{and} \quad \text{ReLU}(y) = y.$$

Thus,

$$|\text{ReLU}(x) - \text{ReLU}(y)| = |x - y| \leq |x - y|.$$

Case 2: $x < 0$ and $y < 0$

Here,

$$\text{ReLU}(x) = 0 \quad \text{and} \quad \text{ReLU}(y) = 0.$$

Therefore,

$$|\text{ReLU}(x) - \text{ReLU}(y)| = |0 - 0| = 0 \leq |x - y|.$$

Case 3: $x \geq 0$ and $y < 0$ (without loss of generality)

In this scenario,

$$\text{ReLU}(x) = x \quad \text{and} \quad \text{ReLU}(y) = 0.$$

Thus,

$$|\text{ReLU}(x) - \text{ReLU}(y)| = |x - 0| = |x| \leq |x - y|.$$

This inequality holds because $x \geq 0$ and $y < 0$, implying $|x| \leq |x - y|$.

In all cases, $|\text{ReLU}(x) - \text{ReLU}(y)| \leq |x - y|$. Therefore, ReLU is **1-Lipschitz continuous**.

C.1.3. SIGMOID FUNCTION

The Sigmoid activation function is defined as:

$$\sigma(x) = \frac{1}{1 + e^{-x}}.$$

The derivative of the Sigmoid function is:

$$\sigma'(x) = \sigma(x)(1 - \sigma(x)).$$

Using the fact that $\forall x \in \mathbb{R}, \sigma(x) > 0, 1 - \sigma(x) > 0$, we apply AM-GM inequality:

$$\frac{\sigma(x) + (1 - \sigma(x))}{2} = \frac{1}{2} \geq \sqrt{\sigma(x)(1 - \sigma(x))}.$$

Squaring both sides,

$$\left(\frac{1}{2}\right)^2 = \frac{1}{4} \geq \sigma(x)(1 - \sigma(x)).$$

Thus,

$$0 \leq \sigma'(x) = \sigma(x)(1 - \sigma(x)) \leq \frac{1}{4}.$$

By the Mean Value Theorem, for any $x, y \in \mathbb{R}$, there exists some c between x and y such that:

$$|\sigma(x) - \sigma(y)| = |\sigma'(c)||x - y|.$$

Using that $|\sigma'(c)| \leq \frac{1}{4}$, we have for all $x, y \in \mathbb{R}$

$$|\sigma(x) - \sigma(y)| \leq \frac{1}{4}|x - y|.$$

Therefore, Sigmoid is **$\frac{1}{4}$ -Lipschitz continuous**.

C.1.4. GAUSSIAN ERROR LINEAR UNIT(GELU)

The GELU activation function is expressed as:

$$\text{GELU}(x) = x\Phi(x),$$

where $\Phi(x)$ is the cumulative distribution function (CDF) of the standard normal distribution:

$$\Phi(x) = \frac{1}{2} \left(1 + \operatorname{erf} \left(\frac{x}{\sqrt{2}} \right) \right).$$

First, we compute the derivative of $\text{GELU}(x)$:

$$\frac{d}{dx} \text{GELU}(x) = \Phi(x) + x\phi(x),$$

where $\phi(x)$ is the probability density function (PDF) of the standard normal distribution:

$$\phi(x) = \frac{1}{\sqrt{2\pi}} e^{-x^2/2}.$$

In order to show the boundedness of the derivative, we examine the second derivative:

$$\frac{d^2}{dx^2} \text{GELU}(x) = 2\phi(x) - x^2\phi(x).$$

Setting the second derivative equal to zero to find critical points:

$$\frac{d^2}{dx^2} \text{GELU}(x) = 0 \implies \phi(x)(2 - x^2) = 0.$$

Since $\phi(x) > 0$ for all $x \in \mathbb{R}$, the extrema of $\frac{d}{dx} \text{GELU}(x)$ occurs at $x = \pm\sqrt{2}$.

Hence, it is enough to examine the value of $\frac{d}{dx} \text{GELU}(x) = \Phi(x) + x\phi(x)$ at $\pm\infty, \pm\sqrt{2}$:

$$\lim_{x \rightarrow \infty} \frac{d}{dx} \text{GELU}(x) = \lim_{x \rightarrow \infty} \Phi(x) + \lim_{x \rightarrow \infty} x\phi(x) = 1 + 0 = 1$$

$$\lim_{x \rightarrow -\infty} \frac{d}{dx} \text{GELU}(x) = \lim_{x \rightarrow -\infty} \Phi(x) + \lim_{x \rightarrow -\infty} x\phi(x) = 0 + 0 = 0$$

$$\frac{d}{dx} \text{GELU}(\sqrt{2}) = \Phi(\sqrt{2}) + \sqrt{2} \cdot \phi(\sqrt{2}) = \frac{1}{2}(1 + \operatorname{erf}(1)) + \sqrt{2} \frac{1}{\sqrt{2\pi}} e^{-1} \approx 1.129$$

$$\frac{d}{dx} \text{GELU}(-\sqrt{2}) = \Phi(-\sqrt{2}) - \sqrt{2} \cdot \phi(-\sqrt{2}) = \frac{1}{2}(1 + \operatorname{erf}(-1)) - \sqrt{2} \frac{1}{\sqrt{2\pi}} e^{-1} \approx -0.129$$

using that

$$\lim_{x \rightarrow \infty} \Phi(x) = 1, \lim_{x \rightarrow -\infty} \Phi(x) = 0, \lim_{x \rightarrow \pm\infty} x\phi(x) = \lim_{x \rightarrow \pm\infty} \frac{1}{\sqrt{2\pi}} x e^{-x^2/2} = 0$$

Thus,

$$\left| \frac{d}{dx} \text{GELU}(x) \right| \leq 1.13, \quad \forall x \in \mathbb{R}.$$

Therefore, GELU is **1.13-Lipschitz continuous**.

C.2. Proof of Lemma 3.5

The spectral norm of a matrix satisfies the following sub-multiplicative property:

$$\|\mathbf{AB}\|_2 \leq \|\mathbf{A}\|_2 \|\mathbf{B}\|_2$$

Using this property, we establish the discrepancy bound for 1-layer propagation in standard neural networks. For two intermediate representations $\mathbf{H}_1^{(l)}$ and $\mathbf{H}_2^{(l)}$, we have:

$$\begin{aligned} \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 &= \|\sigma(\mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)} + \mathbf{b}^{(l)}) - \sigma(\mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)} + \mathbf{b}^{(l)})\|_2 \\ &\leq L_\sigma \|(\mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)} + \mathbf{b}^{(l)}) - (\mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)} + \mathbf{b}^{(l)})\|_2 \\ &\leq L_\sigma \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 \|\mathbf{W}^{(l)}\|_2, \end{aligned}$$

where L_σ is the Lipschitz constant for activation function σ .

C.3. Proof of Theorem 3.6

The discrepancy in the final output representation can be bounded recursively as follows:

$$\begin{aligned} \|\mathbf{H}_1^{(L)} - \mathbf{H}_2^{(L)}\|_2 &\leq L_\sigma \|\mathbf{W}^{(L)}\|_2 \|\mathbf{H}_1^{(L-1)} - \mathbf{H}_2^{(L-1)}\|_2 \\ &\leq L_\sigma^2 \|\mathbf{W}^{(L)}\|_2 \|\mathbf{W}^{(L-1)}\|_2 \|\mathbf{H}_1^{(L-2)} - \mathbf{H}_2^{(L-2)}\|_2 \\ &\leq \dots \\ &\leq L_\sigma^{(L-l)} \left(\prod_{i=l+1}^L \|\mathbf{W}^{(i)}\|_2 \right) \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 \\ &= C \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2, \end{aligned}$$

where $C = L_\sigma^{(L-l)} \prod_{i=l+1}^L \|\mathbf{W}^{(i)}\|_2$ represents the cascade constant.

C.4. Proof of Lemma 3.7

In this proof, we aim to show how minimizing the discrepancy at each aggregation step effectively bounds the final representation discrepancy. To ensure consistent analysis, we assume that the representation matrix $\mathbf{H}_*$ is normalized, satisfying $\|\mathbf{H}_*\|_2 \leq |V|$. By quantifying propagation of discrepancy across linear transformations and various aggregation operations, we demonstrate that controlling intermediate discrepancies reduce the discrepancy in the final output.

C.4.1. REGULAR AGGREGATION

For regular aggregation, the representation discrepancy satisfies:

$$\begin{aligned} &\|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2)\|_2 \\ &= \|\mathbf{A}_1 \mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)} - \mathbf{A}_2 \mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)}\|_2 \\ &= \|(\mathbf{A}_1 \mathbf{H}_1^{(l-1)} - \mathbf{A}_1 \mathbf{H}_2^{(l-1)} + \mathbf{A}_1 \mathbf{H}_2^{(l-1)} - \mathbf{A}_2 \mathbf{H}_2^{(l-1)}) \mathbf{W}^{(l)}\|_2 \\ &\leq (\|\mathbf{A}_1 \mathbf{H}_1^{(l-1)} - \mathbf{A}_1 \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{A}_1 \mathbf{H}_2^{(l-1)} - \mathbf{A}_2 \mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\ &\leq (\|\mathbf{A}_1\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{A}_1 - \mathbf{A}_2\|_2 \|\mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\ &\leq \|\mathbf{A}_1\|_2 \|\mathbf{W}^{(l)}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + |V| \|\mathbf{W}^{(l)}\|_2 \|\mathbf{A}_1 - \mathbf{A}_2\|_2 \end{aligned}$$

This shows that if $\mathbf{A}_1 = \mathbf{A}_2$, the representation discrepancy is linearly bounded by the input difference.

C.4.2. ROW-NORMALIZED AGGREGATION

For row-normalized aggregation, the representation discrepancy satisfies:

$$\begin{aligned}
 & \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2)\|_2 \\
 &= \|\mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)} - \mathbf{D}_2^{-1} \mathbf{A}_2 \mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)}\|_2 \\
 &= \|(\mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_1^{(l-1)} - \mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_2^{(l-1)} + \mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_2^{(l-1)} - \mathbf{D}_2^{-1} \mathbf{A}_2 \mathbf{H}_2^{(l-1)}) \mathbf{W}^{(l)}\|_2 \\
 &\leq (\|\mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_1^{(l-1)} - \mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{D}_1^{-1} \mathbf{A}_1 \mathbf{H}_2^{(l-1)} - \mathbf{D}_2^{-1} \mathbf{A}_2 \mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\
 &\leq (\|\mathbf{D}_1^{-1} \mathbf{A}_1\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{D}_1^{-1} \mathbf{A}_1 - \mathbf{D}_2^{-1} \mathbf{A}_2\|_2 \|\mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\
 &\leq \|\mathbf{W}^{(l)}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + |V| \|\mathbf{W}^{(l)}\|_2 \|\mathbf{D}_1^{-1} \mathbf{A}_1 - \mathbf{D}_2^{-1} \mathbf{A}_2\|_2
 \end{aligned}$$

noting that $\|\mathbf{D}_1^{-1} \mathbf{A}_1\|_2 = 1$ because row-normalized matrices have their largest eigenvalue equal to 1. This demonstrates that the discrepancy is linearly bounded if $\mathbf{A}_1 = \mathbf{A}_2$.

C.4.3. SYMMETRIC-NORMALIZED AGGREGATION

For symmetric-normalized aggregation, the representation discrepancy satisfies:

$$\begin{aligned}
 & \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2)\|_2 \\
 &= \|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)} - \mathbf{D}_2^{-\frac{1}{2}} \mathbf{A}_2 \mathbf{D}_2^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)}\|_2 \\
 &= \|(\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_1^{(l-1)} - \mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)} + \mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)} - \mathbf{D}_2^{-\frac{1}{2}} \mathbf{A}_2 \mathbf{D}_2^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)}) \mathbf{W}^{(l)}\|_2 \\
 &\leq (\|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_1^{(l-1)} - \mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)} - \mathbf{D}_2^{-\frac{1}{2}} \mathbf{A}_2 \mathbf{D}_2^{-\frac{1}{2}} \mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\
 &\leq (\|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + \|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} - \mathbf{D}_2^{-\frac{1}{2}} \mathbf{A}_2 \mathbf{D}_2^{-\frac{1}{2}}\|_2 \|\mathbf{H}_2^{(l-1)}\|_2) \|\mathbf{W}^{(l)}\|_2 \\
 &\leq \|\mathbf{W}^{(l)}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + |V| \|\mathbf{W}^{(l)}\|_2 \|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} - \mathbf{D}_2^{-\frac{1}{2}} \mathbf{A}_2 \mathbf{D}_2^{-\frac{1}{2}}\|_2
 \end{aligned}$$

where $\|\mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}}\|_2 = 1$ due to its normalization. This follows from the fact that $\mathbf{I} - \mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}}$ is positive semi-definite, and there exists $\mathbf{x} = \mathbf{D}_1^{\frac{1}{2}} \mathbf{y}$ such that $\mathbf{x}^\top \mathbf{x} = \mathbf{x}^\top \mathbf{D}_1^{-\frac{1}{2}} \mathbf{A}_1 \mathbf{D}_1^{-\frac{1}{2}} \mathbf{x}$ where $\mathbf{y}$ is eigenvector for $\mathbf{D}_1 - \mathbf{A}_1$ with corresponding eigenvalue 0. This implies that if $\mathbf{A}_1 = \mathbf{A}_2$, the representation discrepancy is linearly bounded by the input difference, similar to the case of regular aggregation.

C.5. Proof of Theorem 3.8

Before we proceed with the proof, let us first establish the following lemma.

Lemma C.1. *Let $\mathbf{A}, \mathbf{B} \in \mathbb{R}_+^{m \times m}$, be two distinct matrices ($\mathbf{A} \neq \mathbf{B}$), and let $\sigma : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}^{m \times n}$ be a non-constant, continuous, element-wise function. Then there exists some $\mathbf{Z} \in \mathbb{R}^{m \times n}$ such that*

$$\sigma(\mathbf{AZ}) \neq \sigma(\mathbf{BZ})$$

Equivalently, no such σ can satisfy $\sigma(\mathbf{AZ}) = \sigma(\mathbf{BZ})$ for all $\mathbf{Z}$ if $\mathbf{A} \neq \mathbf{B}$.

Proof. Since $\mathbf{A} \neq \mathbf{B}$, there exists at least one index (i, j) such that $a_{ij} \neq b_{ij}$. Denote $a_{ij} = a$ and $b_{ij} = b$ ($a > 0, b > 0$). We will construct a particular $\mathbf{Z}$ that reveals the difference under σ . Define $\mathbf{Z}$ so that its j -th row is a scalar variable $z \in \mathbb{R}$ and all other entries are zero. Then, each (i, l) -entry of $\mathbf{AZ}$ is az , while the corresponding entry of $\mathbf{BZ}$ is bz for $l = 1, \dots, n$. Because σ is applied element-wise, $\sigma(\mathbf{AZ}) = \sigma(\mathbf{BZ})$ implies $\sigma(az) = \sigma(bz)$. We analyze two cases:

- **Case 1:** $a = 0$ or $b = 0$: Without loss of generality, let $b = 0$. Then $\sigma(az) = \sigma(0)$ must hold for all z . This forces σ to be the constant, contradicting the given assumption.

- **Case 2:** $a \neq 0$ and $b \neq 0$: Without loss of generality, let $a > b$. Then, $\sigma(az) = \sigma(bz)$ for all z implies

$$\sigma(z) = \sigma\left(\frac{b}{a}z\right) = \sigma\left(\left(\frac{b}{a}\right)^2 z\right) = \cdots = \sigma\left(\left(\frac{b}{a}\right)^n z\right)$$

for any $n > 1$. Since $b/a < 1$, by the continuity of σ , $\sigma(z) = \lim_{n \rightarrow \infty} \sigma\left(\left(\frac{b}{a}\right)^n z\right) = \sigma(0)$. Hence σ would be constant on that range, again contradicting the non-constant assumption.

Thus, in either case, we find that the hypothesis $\sigma(\mathbf{AZ}) = \sigma(\mathbf{BZ})$ for all $\mathbf{Z}$ forces σ to be constant, which is a contradiction. Therefore, there must exist some $\mathbf{Z}$ for which $\sigma(\mathbf{AZ}) \neq \sigma(\mathbf{BZ})$. $\square$

Now, we use the above lemma to show that, if $\hat{\mathbf{A}}_1 \neq \hat{\mathbf{A}}_2$, then no constant C can bound the difference of GCN outputs in terms of the difference of inputs. Suppose, for contradiction, that there exists $C > 0$ such that for every pair $\mathbf{H}_1^{(l-1)}, \mathbf{H}_2^{(l-1)}$ exists, the following holds:

$$\|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 = \|\sigma(\hat{\mathbf{A}}_1 \mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)}) - \sigma(\hat{\mathbf{A}}_2 \mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)})\|_2 \leq C \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2$$

In particular, consider the case where $\mathbf{H}_1^{(l-1)} = \mathbf{H}_2^{(l-1)}$. Then $\|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 = 0$, so the right-hand side of above equation is zero. Thus, we must have $\sigma(\hat{\mathbf{A}}_1 \mathbf{H}_1^{(l-1)} \mathbf{W}^{(l)}) = \sigma(\hat{\mathbf{A}}_2 \mathbf{H}_2^{(l-1)} \mathbf{W}^{(l)})$. However, by Lemma C.1, there exists a suitable choice of $\mathbf{H}$ so that $\mathbf{H}\mathbf{W}^{(l)}$ corresponds to $\mathbf{Z}$ of the lemma, leading to $\sigma(\hat{\mathbf{A}}_1 \mathbf{H}\mathbf{W}^{(l)}) \neq \sigma(\hat{\mathbf{A}}_2 \mathbf{H}\mathbf{W}^{(l)})$. If $\mathbf{H}_1^{(l-1)} = \mathbf{H}_2^{(l-1)} = \mathbf{H}$, we have

$$\|\sigma(\hat{\mathbf{A}}_1 \mathbf{H}\mathbf{W}^{(l)}) - \sigma(\hat{\mathbf{A}}_2 \mathbf{H}\mathbf{W}^{(l)})\|_2 > 0 = C \|\mathbf{H} - \mathbf{H}\|_2$$

This contradiction shows that no such input-independent constant C can exist.

C.6. Proof of Theorem 3.9

In GNNs with row-normalized or symmetric-normalized aggregation, the propagation of representation discrepancy is bounded as:

$$\begin{aligned} \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 &\leq L_\sigma \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2)\|_2 \\ &\leq L_\sigma \|\mathbf{W}^{(l)}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + L_\sigma |V| \|\mathbf{W}^{(l)}\|_2 \|\hat{\mathbf{A}}_1 - \hat{\mathbf{A}}_2\|_2 \end{aligned}$$

where $\hat{\mathbf{A}}$ represents the normalized adjacency matrix and L_σ be the Lipschitz constant of the activation function. This result shows that adjacency matrix perturbations introduce an additional error term, which grows through layers.

D. Proofs in Section 4

D.1. Proof of Theorem 4.1

Note that once condition C2 is satisfied, then C1 automatically holds. Thus, it suffices to show

$$\|D_a^{-1} \mathbf{M}\|_F < \|\tilde{D}_a^{-1} \mathbf{M}\|_F, \quad \text{where } D_a = D + I, \quad \mathbf{M} = \mathbf{H}^{(0:l-1)} \mathbf{W}^{(l)}.$$

Since $\tilde{\mathbf{A}} \subset \mathbf{A}$ by edge removal, each diagonal entry $\tilde{d}_i$ of $\tilde{D}_a$ satisfies $0 < \tilde{d}_i \leq d_i$. Thus $\frac{1}{\tilde{d}_i} \geq \frac{1}{d_i}$ for all i .

Write $D_a^{-1} \mathbf{M}$ and $\tilde{D}_a^{-1} \mathbf{M}$ row by row. If $M_{i,\cdot}$ denotes the i -th row of $\mathbf{M}$, then

$$(D_a^{-1} \mathbf{M})_{i,\cdot} = \frac{1}{d_i} M_{i,\cdot}, \quad (\tilde{D}_a^{-1} \mathbf{M})_{i,\cdot} = \frac{1}{\tilde{d}_i} M_{i,\cdot}.$$

The Frobenius norm is

$$\|D_a^{-1} \mathbf{M}\|_F^2 = \sum_{i=1}^n \frac{1}{d_i^2} \|M_{i,\cdot}\|_2^2, \quad \|\tilde{D}_a^{-1} \mathbf{M}\|_F^2 = \sum_{i=1}^n \frac{1}{\tilde{d}_i^2} \|M_{i,\cdot}\|_2^2.$$

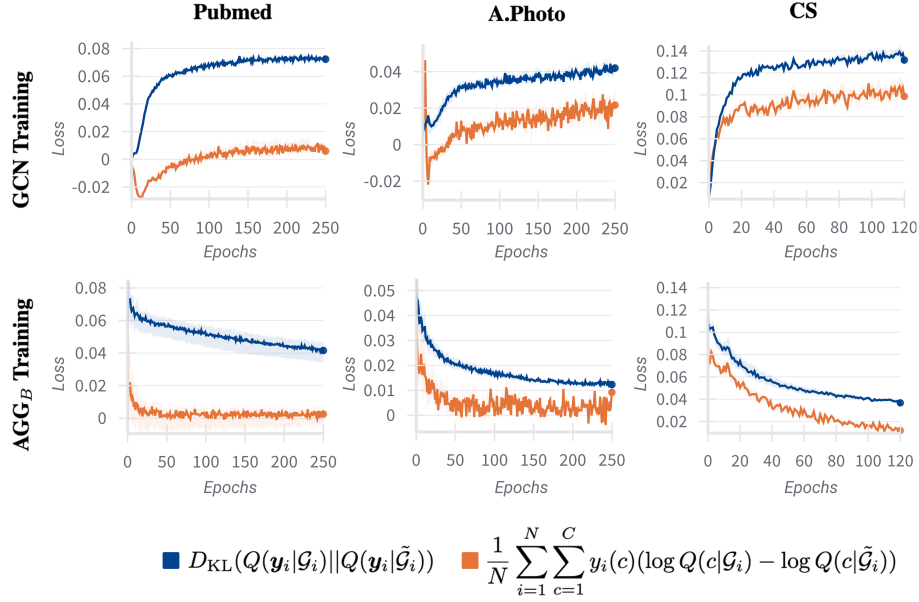


Figure 4. Changes in two different approximations of the robustness loss, $\mathbb{E}_P[\log Q(\mathbf{y}_i|\mathcal{G}_i) - \log Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)]$, during training of the base GCN (top row) and AGG_B (bottom row). Each curve represents the average over 10 independent runs, with shaded areas indicating the minimum and maximum values. Blue represents the robustness term in our proposed *robustness-controlled loss*, where P is approximated by Q . Orange represents the label-based approximation, where P is approximated using ground-truth labels. Both approximations exhibit similar trends: robustness loss gradually emerges during GCN training and is further optimized during AGG_B training.

Because $\frac{1}{\tilde{d}_i^2} \geq \frac{1}{d_i^2}$ whenever $\tilde{d}_i \leq d_i$, each term in the sum for $\tilde{D}^{-1}\mathbf{M}$ is greater or equal. Therefore

$$\|D_a^{-1}\mathbf{M}\|_F^2 \leq \|\tilde{D}_a^{-1}\mathbf{M}\|_F^2 \implies \|D_a^{-1}\mathbf{M}\|_F \leq \|\tilde{D}_a^{-1}\mathbf{M}\|_F.$$

Note that the equality holds if and only if $\|M_{j,\cdot}\|_2^2 = 0$ for every j 's such that $\tilde{d}_j < d_j$. Such a configuration occupies a measure-zero subset of whole space, and thus arises with probability zero in typical real-world scenarios.

Substituting back completes the proof:

$$\|\text{AGG}_B^{(l)}(\mathbf{A})\|_F = \|D_a^{-1}\mathbf{M}\|_F < \|\tilde{D}_a^{-1}\mathbf{M}\|_F = \|\text{AGG}_B^{(l)}(\tilde{\mathbf{A}})\|_F.$$

E. Validity of the Approximation on Robustness Loss

Between equation 3 and equation 4, we approximate the robustness term in the shifted objective under DropEdge. Specifically, the expectation with respect to the true distribution P is approximated using the model's predictive distribution Q as follows:

$$\mathbb{E}_P[\log Q(\mathbf{y}_i|\mathcal{G}_i) - \log Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)] \approx D_{\text{KL}}(Q(\mathbf{y}_i|\mathcal{G}_i)||Q(\mathbf{y}_i|\tilde{\mathcal{G}}_i)).$$

This approximation is based on the assumption $Q \approx P$. Since the true distribution P is inaccessible during training, this assumption allows the term to be computed in practice.

Although the assumption $Q \approx P$ may not strictly hold—particularly in the early stages of training—it becomes increasingly valid as training progresses. Since the model is trained using cross-entropy loss, it explicitly minimizes the KL divergence $D_{\text{KL}}(P(y_i|\mathcal{G}_i)||Q(y_i|\mathcal{G}_i))$, gradually aligning Q with P on the training distribution. Moreover, our framework employs a two-step training procedure, in which this approximation is utilized only after the base GCN has been trained. This staged design ensures that the approximation is applied under more reliable conditions, promoting stable and effective optimization of the proposed *robustness-controlled loss*, $\mathcal{L}_{\text{RC}}$.

Table 10. Test accuracy under varying levels of random edge removal (%) across 12 datasets. A value of 100% indicates that no edges are removed, whereas 0% indicates complete edge removal. Bold entries denote the highest performance for each setting. AGG_B significantly enhances robustness, outperforming the baselines (GCN, DropEdge) in 56 out of 60 cases.

	Cora			Citeseer			Pubmed			Wiki-CS		
	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B
100%	83.44±1.44	83.27±1.55	84.84±1.39	72.45±0.80	72.29±0.60	73.32±0.85	86.48±0.17	86.47±0.21	87.56±0.27	80.26±0.34	80.22±0.55	80.75±0.42
75%	81.85±1.28	81.78±1.53	84.53±1.38	71.66±0.69	71.54±0.47	73.16±1.01	85.99±0.20	86.01±0.12	87.52±0.33	79.28±0.41	79.34±0.55	80.17±0.47
50%	79.63±1.65	79.09±1.51	84.31±1.37	70.76±0.64	70.33±0.40	72.90±1.23	85.48±0.22	85.69±0.22	87.51±0.26	77.85±0.22	78.09±0.47	79.20±0.38
25%	76.28±1.67	76.48±1.06	84.44±1.59	69.35±0.69	68.94±0.46	72.66±1.39	84.80±0.22	84.90±0.24	87.43±0.26	75.50±0.45	76.51±0.45	77.83±0.72
0%	72.63±1.82	72.33±1.70	84.17±1.50	68.12±0.54	67.54±0.49	72.23±1.71	83.86±0.38	84.18±0.39	86.81±0.31	69.60±0.96	72.54±0.82	72.86±1.50

	A.Photo			A.Computer			CS			Physics		
	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B
100%	92.21±1.36	92.14±1.42	92.44±1.42	88.24±0.63	88.08±1.08	88.76±0.65	91.85±0.29	91.91±0.16	93.54±0.37	95.18±0.17	95.13±0.16	95.79±0.17
75%	92.08±1.31	92.02±1.42	92.38±1.35	88.01±0.59	87.83±0.99	88.71±0.66	91.41±0.28	91.43±0.22	93.59±0.38	94.88±0.19	94.87±0.16	95.77±0.16
50%	91.67±1.41	91.77±1.39	92.29±1.40	87.41±0.57	87.38±1.04	88.54±0.55	90.77±0.25	90.71±0.22	93.56±0.42	94.53±0.17	94.54±0.20	95.76±0.17
25%	90.79±1.72	90.92±1.51	91.90±1.51	86.20±0.54	86.28±0.88	88.02±0.55	89.91±0.18	89.95±0.21	93.53±0.53	93.99±0.18	93.96±0.18	95.72±0.17
0%	84.88±1.81	85.99±1.66	86.11±3.35	76.77±1.48	78.17±1.32	82.50±1.19	93.10±0.31	93.17±0.23	93.18±0.74	94.33±0.51	94.63±0.43	95.44±0.20

	Arxiv			Actor			Squirrel			Chameleon		
	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B	GCN	DropEdge	GCN _B
100%	71.80±0.10	71.73±0.21	72.43±0.16	30.16±0.73	29.86±0.82	30.56±0.84	41.67±2.42	41.66±2.11	42.39±2.19	40.19±4.29	40.23±4.34	40.96±4.83
75%	70.41±0.17	70.42±0.22	71.57±0.17	30.72±0.98	30.33±0.95	30.70±0.92	40.64±2.89	40.71±2.71	41.45±2.54	39.52±4.37	39.51±4.21	40.16±4.86
50%	67.97±0.18	68.13±0.30	69.95±0.14	31.09±1.21	30.49±0.99	31.26±1.08	39.88±2.19	39.63±2.36	40.86±2.21	39.93±4.22	40.13±5.18	39.83±4.08
25%	62.81±0.32	63.11±0.42	66.28±0.14	31.24±0.91	31.20±1.42	31.00±1.61	37.47±2.52	37.33±2.30	40.06±2.66	38.41±3.20	38.46±3.68	38.57±4.39
0%	44.09±0.64	44.55±0.95	53.57±0.35	33.68±1.38	32.80±1.66	30.76±1.76	32.18±2.51	32.40±1.67	35.40±2.43	33.24±3.13	33.22±3.10	36.29±4.64

To empirically evaluate the validity of this approximation, we estimate the expectation under P using ground-truth labels as a proxy. Specifically, we computed the following quantity:

$$\frac{1}{N} \sum_{i=1}^N \sum_{c=1}^C y_i(c) (\log Q(c|\mathcal{G}_i) - \log Q(c|\tilde{\mathcal{G}}_i)),$$

where $y_i(c)$ denotes the one-hot encoded ground-truth label for node i . While this proxy samples only one class per node and may be affected by label noise, it still offers a practical estimate for validating the approximation. As shown in Figure 4, the trend of this label-based quantity closely mirrors that of the approximated KL divergence, indicating that our approximation effectively captures the underlying behavior. Furthermore, AGG_B exhibits robust optimization behavior even under this label-based approximation, demonstrating its effectiveness in terms of robustness optimization.

F. Assessing Edge-Robustness via Random Edge Removal at Test Time

While we previously demonstrated the edge-robustness benefits of AGG_B through improvements in degree bias and structural disparity, we now provide a more direct evaluation by measuring model performance under random edge removal during inference. Specifically, we assess how test accuracy degrades as edges are randomly removed from the input graph. We compare three models: (1) a standard GCN trained normally, (2) a GCN trained with DropEdge, and (3) GCN_B, which incorporates AGG_B into a pre-trained GCN. The results are presented in Table 10.

GCN_B significantly outperforms both DropEdge and standard GCN in 56 out of 60 cases, indicating that AGG_B enables GCNs to generate more consistent representations under structural perturbations, thereby exhibiting superior edge-robustness.

Interestingly, in 3 of the 4 cases where GCN_B does not outperform the baselines, the performance of the standard GCN improves as edges are removed—specifically on the Actor dataset. This aligns with our observation in Table 1 that an MLP outperforms GCN on this dataset, suggesting that leveraging edge information may not be beneficial. These findings imply that the edges in Actor are likely too noisy or uninformative. Nevertheless, even on Actor, GCN_B maintains stable accuracy under edge removal, highlighting that AGG_B still contributes to enhanced edge-robustness.

In contrast, models trained with DropEdge often show marginal improvements or even performance degradation compared to standard GCNs. This supports our claim that DropEdge alone is insufficient for achieving edge-robustness, due to the inherent inductive bias of GNNs.

Table 11. Accuracy across different layer architectures used as AGG_B . Blue indicates no improvement over the base GCN, and bold text highlights the best-performing architecture per dataset. The rightmost column reports the average rank. Our proposed design consistently improves performance and achieves the best overall ranking.

Method	Cora	Citeseer	PubMed	Wiki-CS	A.Photo	A.Computer	CS	Physics	Arxiv	Actor	Squirrel	Chameleon	Rank
GCN	83.44 $\pm$ 1.44	72.45 $\pm$ 0.80	86.48 $\pm$ 0.17	80.26 $\pm$ 0.34	92.21 $\pm$ 1.36	88.24 $\pm$ 0.63	91.85 $\pm$ 0.29	95.18 $\pm$ 0.17	71.80 $\pm$ 0.10	30.16 $\pm$ 0.73	41.67 $\pm$ 2.42	40.19 $\pm$ 4.29	4.75
AGG	84.01 $\pm$ 1.58	72.70 $\pm$ 0.82	86.82 $\pm$ 0.55	79.54 $\pm$ 0.94	91.74 $\pm$ 1.76	88.03 $\pm$ 0.74	91.63 $\pm$ 0.28	95.02 $\pm$ 0.31	72.27 $\pm$ 0.12	29.29 $\pm$ 1.01	40.18 $\pm$ 4.48	40.69 $\pm$ 2.69	4.92
Residual	83.29 $\pm$ 2.50	72.71 $\pm$ 0.83	87.46 $\pm$ 0.24	80.80 $\pm$ 1.08	92.40 $\pm$ 1.64	88.95 $\pm$ 0.44	92.05 $\pm$ 0.28	95.27 $\pm$ 0.38	72.29 $\pm$ 0.12	30.57 $\pm$ 0.53	41.32 $\pm$ 2.95	39.77 $\pm$ 4.57	3.00
JKNet-style	83.35 $\pm$ 2.22	72.76 $\pm$ 0.78	87.45 $\pm$ 0.25	80.78 $\pm$ 0.62	92.08 $\pm$ 1.61	87.97 $\pm$ 0.77	93.36 $\pm$ 0.56	95.85 $\pm$ 0.23	72.19 $\pm$ 0.18	30.70 $\pm$ 1.21	40.70 $\pm$ 2.30	40.29 $\pm$ 4.68	3.42
AGG _B (single)	84.28 $\pm$ 1.57	72.78 $\pm$ 0.52	87.58 $\pm$ 0.38	80.70 $\pm$ 0.00	92.00 $\pm$ 1.51	88.50 $\pm$ 0.73	92.01 $\pm$ 0.37	95.23 $\pm$ 0.27	72.09 $\pm$ 0.11	30.13 $\pm$ 0.73	41.91 $\pm$ 2.55	40.64 $\pm$ 4.83	3.33
AGG _B (ours)	84.84 $\pm$ 1.39	73.32 $\pm$ 0.85	87.56 $\pm$ 0.27	80.75 $\pm$ 0.42	92.44 $\pm$ 1.42	88.76 $\pm$ 0.65	93.54 $\pm$ 0.37	95.79 $\pm$ 0.17	72.43 $\pm$ 0.16	30.56 $\pm$ 0.84	42.39 $\pm$ 2.19	40.96 $\pm$ 4.83	1.58

G. Extensive Ablation Study of Alternative Layer Architectures

In this section, we extend the results presented in Table 4 to all 12 datasets used in our experiments. As shown in Table 11, although several alternative layer architectures provide performance gains in specific cases under our training scheme and loss, only our original AGG_B design consistently and significantly improves performance across all datasets.

We also evaluate a simplified, single-layer variant of AGG_B that restricts the usable representation to only the immediate previous layer, $\mathbf{H}^{(l-1)}$, and is formulated as:

$$(\mathbf{D} + \mathbf{I})^{-1} \mathbf{H}^{(l-1)} \mathbf{W}^{(l)}.$$

This variant satisfies the same theoretical conditions—C1 (edge-awareness) and C2 (stability)—outlined in Section 4.2, just like our proposed architecture. While it improves performance on 10 out of 12 datasets—making it the most competitive alternative—the improvements are relatively marginal compared to those achieved by AGG_B.

The motivation for integrating representations from all preceding layers, rather than relying on a single layer, is to mitigate the risk of accumulating structural discrepancies. Ideally, AGG_B could fully correct such inconsistencies at each layer. In practice, however, residual discrepancies may persist in intermediate layers and propagate through the network, ultimately affecting the final output. Relying solely on $\mathbf{H}^{(l-1)}$ risks amplifying these unresolved issues, whereas aggregating $\mathbf{H}^{(0:l-1)}$ enables AGG_B to leverage earlier, potentially less corrupted representations, leading to more robust corrections.

Importantly, the performance gains from AGG_B are not solely attributed to multi-layer integration. We also compare it with a JKNet-style block, which similarly aggregates outputs from all previous layers and is formulated as $\mathbf{H}^{(0:l-1)} \mathbf{W}^{(l)}$. AGG_B outperforms this JKNet-style design on 9 out of 12 datasets, while the JKNet-style variant even degrades performance on 4 datasets. This result suggests that the inclusion of degree normalization, $(\mathbf{D} + \mathbf{I})^{-1}$ —a key component of AGG_B that ensures satisfaction of the conditions outlined in Section 4.2 (i.e., (1) edge-awareness and (2) stability)—is crucial for achieving consistent performance improvements across diverse datasets.

Although AGG_B performs robustly in our experiments and ablations, we acknowledge that concatenating all preceding representations can introduce information redundancy and noise, particularly as GNN depth increases. However, this is not currently a critical issue, as GNNs typically achieve optimal performance at relatively shallow depths (e.g., two layers) due to over-smoothing. That said, as deeper GNNs become more effective in future research, developing more streamlined integration mechanisms that reduce redundancy and noise presents a promising direction for extending this work.

H. Extensive Ablation Study on Deeper GCNs

In this section, we evaluate the effectiveness of AGG_B when applied to deeper GCN, using 4, 8, 12, 16, and 20-layer models across 6 datasets. In addition to the standard GCN and GCN_B, we include two more variants: (1) GCN trained with DropEdge, and (2) DropEdge_B, where AGG_B is integrated into a GCN trained with DropEdge.

These variants are included to assess whether AGG_B can provide further improvements beyond what DropEdge achieves, particularly in deep architectures. This is motivated by the original DropEdge paper (Rong et al., 2020), which highlights its effectiveness in alleviating oversmoothing and demonstrates more substantial performance gains in deeper GNNs. The results are presented in Table 12.

AGG_B improves performance in 28 out of 30 configurations, demonstrating that its effectiveness is robust to architectural depth. Notably, performance gains tend to increase with depth, suggesting that deeper GNNs are more susceptible

Table 12. Accuracy with varying GCN depths. DropEdge_B denotes a GCN trained with DropEdge followed by integration of AGG_B. All experiments use a hidden dimension of 256 and a learning rate of 0.001. For fair comparison, the edge dropping ratio is fixed at 0.5 for both DropEdge and AGG_B, and the hyperparameter λ is fixed at 1.0. Integrating AGG_B improves performance in 56 out of 60 configurations.

	Cora				Pubmed				A.Computer			
	GCN	GCN _B	DropEdge	DropEdge _B	GCN	GCN _B	DropEdge	DropEdge _B	GCN	GCN _B	DropEdge	DropEdge _B
4	81.98 $\pm$ 1.45	82.50$\pm$1.65	82.51 $\pm$ 1.76	82.96$\pm$1.83	83.73 $\pm$ 0.20	86.87$\pm$0.55	84.13 $\pm$ 0.25	87.14$\pm$0.35	86.66 $\pm$ 0.53	86.82$\pm$0.49	86.23 $\pm$ 1.11	86.34$\pm$1.14
8	77.54 $\pm$ 2.05	79.99$\pm$1.45	79.64 $\pm$ 1.89	80.86$\pm$1.44	83.07 $\pm$ 0.16	85.90$\pm$0.48	83.19 $\pm$ 0.27	85.49$\pm$0.58	80.60 $\pm$ 2.22	81.07$\pm$2.33	79.78 $\pm$ 1.86	80.21$\pm$1.92
12	73.42 $\pm$ 2.08	77.70$\pm$1.99	78.41 $\pm$ 1.71	79.65$\pm$2.01	82.44 $\pm$ 0.29	85.48$\pm$0.68	82.53 $\pm$ 0.35	84.64$\pm$0.52	75.08 $\pm$ 2.24	76.45$\pm$2.06	76.17 $\pm$ 3.65	77.16$\pm$3.75
16	69.78 $\pm$ 3.45	75.67$\pm$3.32	77.92 $\pm$ 1.92	79.39$\pm$1.72	82.08 $\pm$ 0.45	85.21$\pm$0.64	81.76 $\pm$ 0.47	83.78$\pm$0.68	73.50 $\pm$ 4.52	74.83$\pm$4.73	74.89 $\pm$ 1.42	76.20$\pm$1.41
20	64.15 $\pm$ 5.08	72.23$\pm$3.45	73.89 $\pm$ 2.99	76.31$\pm$3.41	81.75 $\pm$ 0.55	85.23$\pm$0.75	81.17 $\pm$ 0.36	83.20$\pm$0.51	67.99 $\pm$ 3.98	68.83$\pm$4.42	72.14 $\pm$ 2.88	73.69$\pm$3.05

	CS				Squirrel				Chameleon			
	GCN	GCN _B	DropEdge	DropEdge _B	GCN	GCN _B	DropEdge	DropEdge _B	GCN	GCN _B	DropEdge	DropEdge _B
4	89.63 $\pm$ 0.40	91.61$\pm$0.34	89.70 $\pm$ 0.37	91.51$\pm$0.43	33.96 $\pm$ 1.01	34.37$\pm$1.31	33.76 $\pm$ 1.42	33.82$\pm$1.38	39.46 $\pm$ 3.17	40.00$\pm$3.05	37.95 $\pm$ 4.27	38.96$\pm$4.17
8	88.44 $\pm$ 0.38	91.33$\pm$0.32	88.58 $\pm$ 0.22	91.24$\pm$0.18	34.06 $\pm$ 1.28	34.41$\pm$1.96	34.14 $\pm$ 1.47	34.41$\pm$1.38	37.55$\pm$2.51	37.33 $\pm$ 2.86	37.99$\pm$3.98	37.06 $\pm$ 4.31
12	86.84 $\pm$ 0.48	90.57$\pm$0.28	87.89 $\pm$ 0.26	91.03$\pm$0.33	34.78$\pm$1.67	34.74 $\pm$ 1.69	34.91 $\pm$ 1.74	35.05$\pm$1.65	37.18 $\pm$ 4.83	37.32$\pm$5.14	36.11$\pm$4.43	35.25 $\pm$ 3.76
16	85.04 $\pm$ 0.81	89.98$\pm$0.52	87.41 $\pm$ 0.44	90.76$\pm$0.53	34.49 $\pm$ 1.66	34.70$\pm$1.96	34.77 $\pm$ 1.39	34.90$\pm$1.28	36.35 $\pm$ 3.37	36.59$\pm$2.50	34.92 $\pm$ 5.12	35.88$\pm$3.57
20	82.34 $\pm$ 1.86	88.59$\pm$1.47	85.55 $\pm$ 0.99	88.77$\pm$1.08	34.23 $\pm$ 1.24	34.80$\pm$1.61	34.83 $\pm$ 1.33	34.92$\pm$1.58	35.63 $\pm$ 4.37	37.30$\pm$4.52	34.57 $\pm$ 4.63	34.66$\pm$3.92

Table 13. Accuracy from 5 independent runs of a 2-layer GCN (hidden dimension: 256), and after integration of AGG_B (GCN_B), evaluated on the public split of larger datasets: Flickr, Ogbn-arxiv, and Reddit. The GCN is trained using fixed hyperparameters (learning rate: 0.001, dropout: 0.5), while AGG_B is trained with fixed parameters ($\lambda = 1.0$, DropEdge rate: 0.5). Integrating AGG_B consistently improves overall performance, with the largest gains observed in low-degree, heterophilic nodes.

	Flickr		Arxiv		Reddit	
	GCN	GCN _B	GCN	GCN _B	GCN	GCN _B
Overall Accuracy	52.50 $\pm$ 0.15	52.84$\pm$0.08	71.06 $\pm$ 0.10	71.37$\pm$0.10	94.61 $\pm$ 0.01	94.89$\pm$0.01
High-degree Nodes	49.66 $\pm$ 0.26	49.87$\pm$0.18	80.06$\pm$0.11	79.95 $\pm$ 0.14	98.81 $\pm$ 0.01	98.84$\pm$0.01
Low-degree Nodes	53.93 $\pm$ 0.28	54.58$\pm$0.16	62.32 $\pm$ 0.05	63.16$\pm$0.04	88.06 $\pm$ 0.01	88.84$\pm$0.03
Homophilous Nodes	80.58$\pm$0.49	80.44 $\pm$ 0.10	94.87$\pm$0.02	94.60 $\pm$ 0.03	99.74$\pm$0.01	99.66 $\pm$ 0.01
Heterophilous Nodes	18.00 $\pm$ 0.07	18.18$\pm$0.07	32.30 $\pm$ 0.09	33.78$\pm$0.08	84.10 $\pm$ 0.01	85.01$\pm$0.02

to structural inconsistencies as representations undergo repeated aggregation—thus creating greater opportunities for improvement via enhanced edge-robustness.

As expected, DropEdge yields more substantial improvements in deeper architectures, while its effects remain marginal in shallow ones. Importantly, integrating AGG_B into DropEdge-trained models significantly boosts performance in 28 out of 30 settings. This demonstrates that AGG_B provides a distinct benefit—specifically, enhanced edge-robustness. These results reinforce our claim that DropEdge alone is insufficient for addressing edge-robustness, regardless of model depth, and that AGG_B offers a principled approach to mitigating structural inconsistencies in deep GNNs.

I. Additional Experiments on Larger Datasets

To further demonstrate the broad applicability of AGG_B, we include results on three larger datasets: Arxiv (Hu et al., 2020a), Reddit (Hamilton et al., 2017), and Flickr (Zeng et al., 2020), all of which are loaded from the Deep Graph Library (DGL). As shown in Table 13, AGG_B consistently improves performance across all three datasets, in line with earlier findings. In all cases, the performance gains primarily stem from improvements on low-degree and heterophilous nodes, highlighting that the observed benefits are indeed driven by enhanced edge-robustness. It is also worth noting that these results are obtained without any hyperparameter tuning. This suggests that further improvements are possible with tuning—as the larger performance gain observed on Arxiv in Table 1.

Additionally, we conduct the edge removal experiments described in Appendix F. The performance degradation from random edge removal is significantly reduced when using AGG_B, further validating its effectiveness on larger-scale datasets.

Table 14. Accuracy under random edge removal (%) in test, using the same experimental settings as Table 13. AGG_B consistently improves performance across all edge removal ratios and datasets, with greater gains at higher removal ratios.

	Flickr		Arxiv		Reddit	
	GCN	GCN_B	GCN	GCN_B	GCN	GCN_B
100%	52.50 $\pm$ 0.15	52.84$\pm$0.08	71.06 $\pm$ 0.10	71.37$\pm$0.10	94.61 $\pm$ 0.01	94.89$\pm$0.01
75%	50.95 $\pm$ 0.12	51.56$\pm$0.11	69.58 $\pm$ 0.08	70.59$\pm$0.07	94.47 $\pm$ 0.01	94.87$\pm$0.01
50%	47.49 $\pm$ 0.52	49.66$\pm$0.21	67.44 $\pm$ 0.11	69.00$\pm$0.06	94.18 $\pm$ 0.01	94.82$\pm$0.01
25%	41.31 $\pm$ 1.06	45.82$\pm$0.34	62.50 $\pm$ 0.08	65.50$\pm$0.06	93.52 $\pm$ 0.01	94.38$\pm$0.03
0%	31.73 $\pm$ 1.33	39.57$\pm$0.70	45.36 $\pm$ 0.19	54.56$\pm$0.04	38.91 $\pm$ 0.01	45.48$\pm$0.27

J. Generalizing Theorem 3.9 to Other GNN Architectures

In this section, we extend our discrepancy analysis—Theorem 3.9—beyond GCN to a broader class of GNN architectures. We provide proofs for three representative models: GraphSAGE, GIN, and GAT, which are also used in our experiments to assess the generalizability of AGG_B , as presented in Section 6.4. These results theoretically demonstrate that the issue of non-optimizable edge-robustness is not specific to GCN, but is a fundamental limitation shared across various GNN architectures—one that AGG_B is designed to address. We omit SGC from this analysis, as it can be regarded as a linearized variant of GCN and is therefore already covered by the proof in Appendix C.

J.1. GraphSAGE

The GraphSAGE layer is formulated as:

$$\mathbf{H}^{(l)} = \sigma(\text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}) + \mathbf{H}^{(l-1)}\mathbf{W}_2^{(l)}) = \sigma(\hat{\mathbf{A}}\mathbf{H}^{(l-1)}\mathbf{W}_1^{(l)} + \mathbf{H}^{(l-1)}\mathbf{W}_2^{(l)}),$$

where $\hat{\mathbf{A}} = \mathbf{D}^{-1}\mathbf{A}$ denotes the normalized adjacency matrix. Then, the discrepancy at layer l satisfies:

$$\begin{aligned} \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 &\leq L_\sigma \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2) + \mathbf{H}_1^{(l-1)}\mathbf{W}_2^{(l)} - \mathbf{H}_2^{(l-1)}\mathbf{W}_2^{(l)}\|_2 \\ &\leq L_\sigma \|\hat{\mathbf{A}}_1\mathbf{H}_1^{(l-1)}\mathbf{W}_1^{(l)} - \hat{\mathbf{A}}_2\mathbf{H}_2^{(l-1)}\mathbf{W}_1^{(l)}\|_2 + L_\sigma \|\mathbf{W}_2^{(l)}\|_2 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 \\ &\leq L_\sigma (\|\mathbf{W}_1^{(l)}\|_2 + \|\mathbf{W}_2^{(l)}\|_2) \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + L_\sigma |V| \|\mathbf{W}_1^{(l)}\|_2 \|\hat{\mathbf{A}}_1 - \hat{\mathbf{A}}_2\|_2 \\ &\leq C_1 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + C_2 \end{aligned}$$

where $C_1 = L_\sigma (\|\mathbf{W}_1^{(l)}\|_2 + \|\mathbf{W}_2^{(l)}\|_2)$, $C_2 = L_\sigma |V| \|\mathbf{W}_1^{(l)}\|_2 \|\hat{\mathbf{A}}_1 - \hat{\mathbf{A}}_2\|_2$.

J.2. GIN

The GIN layer is formulated as:

$$\mathbf{H}^{(l)} = \text{MLP}^{(l)}(\text{AGG}^{(l)}(\mathbf{H}^{(l-1)}, \mathbf{A}) + (1 + \epsilon^{(l)})\mathbf{H}^{(l-1)}) = \text{MLP}^{(l)}(\mathbf{A}\mathbf{H}^{(l-1)} + (1 + \epsilon^{(l)})\mathbf{H}^{(l-1)}),$$

where $\epsilon^{(l)}$ is a learnable scalar at layer l . Then, the discrepancy at layer l satisfies:

$$\begin{aligned} \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 &\leq C \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2) + (1 + \epsilon^{(l)})\mathbf{H}_1^{(l-1)} - (1 + \epsilon^{(l)})\mathbf{H}_2^{(l-1)}\|_2 \\ &\leq C \|\mathbf{A}_1\mathbf{H}_1^{(l-1)} - \mathbf{A}_2\mathbf{H}_2^{(l-1)}\|_2 + C |1 + \epsilon^{(l)}| \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 \\ &\leq C (\|\mathbf{A}_1\|_2 + |1 + \epsilon^{(l)}|) \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + C |V| \|\mathbf{A}_1 - \mathbf{A}_2\|_2 \\ &\leq C_1 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + C_2 \end{aligned}$$

where C is the discrepancy bound of $\text{MLP}^{(l)}$, $C_1 = C (\|\mathbf{A}_1\|_2 + |1 + \epsilon^{(l)}|)$, $C_2 = C |V| \|\mathbf{A}_1 - \mathbf{A}_2\|_2$.

J.3. GAT

The GAT layer is defined as:

$$\begin{aligned} \mathbf{h}_i^{(l)} &= \sigma \left(\sum_{j \in \mathcal{N}_i} \alpha_{ij} \mathbf{W}_1^{(l)} \mathbf{h}_j^{(l-1)} \right), \\ \alpha_{ij} &= \frac{\exp(\text{LeakyReLU}(\mathbf{a}^\top [\mathbf{W}' \mathbf{h}_i^{(l-1)} \parallel \mathbf{W}' \mathbf{h}_j^{(l-1)}]))}{\sum_{k \in \mathcal{N}_i} \exp(\text{LeakyReLU}(\mathbf{a}^\top [\mathbf{W}' \mathbf{h}_i^{(l-1)} \parallel \mathbf{W}' \mathbf{h}_k^{(l-1)}]))} \\ &= \frac{\exp(\text{LeakyReLU}(\mathbf{p}^\top \mathbf{W}' \mathbf{h}_i^{(l-1)} + \mathbf{q}^\top \mathbf{W}' \mathbf{h}_j^{(l-1)}))}{\sum_{k \in \mathcal{N}_i} \exp(\text{LeakyReLU}(\mathbf{p}^\top \mathbf{W}' \mathbf{h}_i^{(l-1)} + \mathbf{q}^\top \mathbf{W}' \mathbf{h}_k^{(l-1)}))}, \end{aligned}$$

where $\mathbf{a} \in \mathbb{R}^{2F'}$ is the original attention weight vector, and $\mathbf{p}, \mathbf{q} \in \mathbb{R}^{F'}$ are its components such that $\mathbf{a}^\top [\cdot \parallel \cdot] = \mathbf{p}^\top (\cdot) + \mathbf{q}^\top (\cdot)$. The induced attention matrix can be interpreted as:

$$\mathbf{A}^* = \text{RowNorm}(\exp(\text{LeakyReLU}(\text{diag}(\mathbf{p}^\top \mathbf{W}' \mathbf{H}^{(l-1)})^\top) \mathbf{A} + \mathbf{A} \text{diag}(\mathbf{q}^\top \mathbf{W}' \mathbf{H}^{(l-1)})^\top))).$$

Since $\mathbf{A}^*$ is row-stochastic, we have $\|\mathbf{A}^*\|_2 = 1$. The discrepancy is thus bounded by:

$$\begin{aligned} \|\mathbf{H}_1^{(l)} - \mathbf{H}_2^{(l)}\|_2 &\leq L_\sigma \|\text{AGG}^{(l)}(\mathbf{H}_1^{(l-1)}, \mathbf{A}_1) - \text{AGG}^{(l)}(\mathbf{H}_2^{(l-1)}, \mathbf{A}_2)\|_2 \\ &\leq L_\sigma \|\mathbf{A}_1^* \mathbf{H}_1^{(l-1)} \mathbf{W}_1^{(l)} - \mathbf{A}_2^* \mathbf{H}_2^{(l-1)} \mathbf{W}_1^{(l)}\|_2 \\ &\leq L_\sigma \|\mathbf{W}_1^{(l)}\| \|\mathbf{A}_1^* \mathbf{H}_1^{(l-1)} - \mathbf{A}_1^* \mathbf{H}_2^{(l-1)} + \mathbf{A}_1^* \mathbf{H}_2^{(l-1)} - \mathbf{A}_2^* \mathbf{H}_2^{(l-1)}\|_2 \\ &\leq L_\sigma \|\mathbf{W}_1^{(l)}\| \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + L_\sigma \|\mathbf{W}_1^{(l)}\|_2 \|\mathbf{H}_2^{(l-1)}\|_2 \|\mathbf{A}_1^* - \mathbf{A}_2^*\|_2 \\ &\leq C_1 \|\mathbf{H}_1^{(l-1)} - \mathbf{H}_2^{(l-1)}\|_2 + C_2, \end{aligned}$$

where $C_1 = L_\sigma \|\mathbf{W}_1^{(l)}\|_2$, $C_2 = C_1 V \|\mathbf{A}_1^* - \mathbf{A}_2^*\|_2$.