

Mitigating Message Imbalance in Fraud Detection with Dual-View Graph Representation Learning

Yudan Song^{1,2†}, Yuecen Wei^{3,4†}, Yuhang Lu^{1,2}, Qingyun Sun³, Minglai Shao⁵,
Li-e Wang^{1,2*}, Chunming Hu^{3,4}, Xianxian Li^{1,2}, Xingcheng Fu^{1,2*}

¹Key Lab of Education Blockchain and Intelligent Technology, Ministry of Education,
Guangxi Normal University, Guilin, China

²Guangxi Key Lab of Multi-Source Information Mining and Security,
Guangxi Normal University, Guilin, China

³School of Software, Beihang University, Beijing, China

⁴SKLCCSE, School of Computer Science and Engineering, Beihang University, China

⁵School of New Media and Communication, Tianjin University, China

{songyudan, lyh0620}@stu.gxnu.edu.cn, {wanglie, lixx, fuxc}@gxnu.edu.cn,
{weiyu, sunqy, hucm}@buaa.edu.cn, shaoml@tju.edu.cn

Abstract

Graph representation learning has become a mainstream method for fraud detection due to its strong expressive power, which focuses on enhancing node representations through improved neighborhood knowledge capture. However, the focus on local interactions leads to imbalanced transmission of global topological information and increased risk of node-specific information being overwhelmed during aggregation due to the imbalance between fraud and benign nodes. In this paper, we first summarize the impact of topology and class imbalance on downstream tasks in GNN-based fraud detection, as the problem of imbalanced supervisory messages is caused by fraudsters' topological behavior obfuscation and identity feature concealment. Based on statistical validation, we propose a novel dual-view graph representation learning method to mitigate **Message imbalance in Fraud Detection (MimbFD)**. Specifically, we design a topological message reachability module for high-quality node representation learning to penetrate fraudsters' camouflage and alleviate insufficient propagation. Then, we introduce a local confounding debiasing module to adjust node representations, enhancing the stable association between node representations and labels to balance the influence of different classes. Finally, we conducted experiments on three public fraud datasets, and the results demonstrate that MimbFD exhibits outstanding performance in fraud detection.

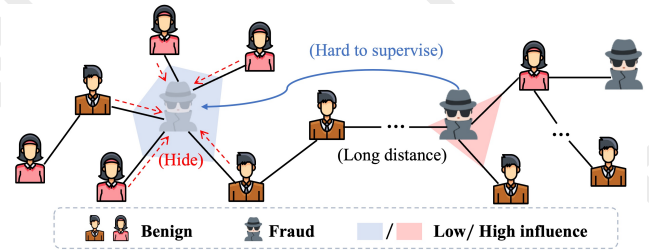


Figure 1: Fraudster distribution patterns. The concealment strategies of fraudsters in different topological spaces result in varying levels of influence and lead to unreachable identity supervisory messages.

1 Introduction

With the rapid development of intelligent technology, various applications facilitate people's lives, such as financial applications [Xu *et al.*, 2021], social media [Peng *et al.*, 2023; Hao *et al.*, 2024], and review systems [Dhawan *et al.*, 2019]. However, factors such as limited access to advanced knowledge by vulnerable groups (benign users) cause malicious fraudsters to exploit information gaps to execute successful fraudulent tactics [Song *et al.*, 2024]. This widespread phenomenon severely affects the order on various platforms, bringing significant attention to fraud detection technologies [Qiao *et al.*, 2024]. Thus, fraud detection serves as an effective means for early warning, timely intervention, and informed decision-making [Goswami *et al.*, 2017].

In recent years, thanks to the strong expressive power of Graph Neural Networks (GNNs) [Wei *et al.*, 2025; Wei *et al.*, 2024; Fu *et al.*, 2025], an increasing number of works [Qiao *et al.*, 2024] have adopted GNNs for fraud detection research. Some works focus on the fraudster's camouflage behavior, where fraudsters mimic the behavior patterns of benign users and interact with them, deceiving the model into classifying them as benign [Wang *et al.*, 2023; Hu *et al.*, 2023]. To uncover such camouflage, these methods emphasize differenti-

*Co-corresponding Authors.

ating fraudsters from benign users regarding node representation, thereby enhancing the model’s ability to identify anomalies [Dou *et al.*, 2020; Zhang *et al.*, 2021; Liu *et al.*, 2021; Chen *et al.*, 2024]. However, they often focus only on local interaction features, overlooking critical long-distance information and neglecting to consider the impact of the fusion of diverse supervised information when dealing with locally confounding messages. In such cases, some fraudulent activities may be overlooked due to the lack of a global perspective. Therefore, starting from the topological distances in the graph, capturing relationships and patterns between long-distance nodes is promising for a better understanding of fraudulent behavior. Next, we discuss the impact of class and topological imbalance on supervisory message transmission in fraud scenarios, as shown in Fig. 1. *Firstly*, fraudsters are dispersed across different communities and exhibit an imbalance in the transmission of topological information to influence their neighbors due to the diversity of the topological structure, either to expand the impact of fraud or to achieve camouflage. *Secondly*, at the data level, fraudsters are the minority, while benign users are the majority. This class imbalance leads to the supervisory message of the minority class being overwhelmed during the message-passing process in GNN, causing severe information aggregation drift.

Challenges. Although GNN-based fraud detection has achieved nice performance, it still faces the following two challenges in solving the problem of inadequate representation of the imbalanced supervisory message: (1) Graph representation learning (GRL) often distorts when aggregating supervisory messages from distant topological locations. Furthermore, due to the topological message imbalance brought by fraudsters’ preferences, nodes fail to learn key useful representations. Therefore, *the challenge is to fully acquire the messages from the global topology to mitigate the conflict between the intrinsic features of nodes and the representations of local neighborhoods.* (2) Existing methods tend to adjust the local neighborhood to avoid unnecessary noise. Due to the extreme class imbalance between fraudsters and benign users, such methods magnify information bias and provide favorable conditions for fraudsters to camouflage. *The challenge lies in how to establish an essential correlation between node representation and labels and enhance their connection among massive messages.*

Contributions. We first perform statistical analysis on fraudsters’ topological behavior obfuscation and identity feature concealment intentions, uncovering the supervisory message imbalance phenomenon in GNN-based fraud detection. To address these challenges, we propose a novel dual-view graph representation learning method to mitigate **Message imbalance in Fraud Detection**, named MimbFD. Specifically, we uncover fraudsters’ camouflage by capturing information that is strongly correlated with the network topology and offer guidance for mitigating topological imbalance. We then introduce a mechanism to de-correlate confounding associations in local messages, aiming to emphasize representations that reflect nodes’ intrinsic information and their association with labels. Finally, empirical results on three datasets confirm the method’s effectiveness in alleviating supervisory message imbalance. In conclusion, our contributions are

summarized as follows:

- We are the first to categorize the imbalance problem in fraud detection from a dual view as an imbalance in supervisory messages in GRL, highlighting the importance of mitigating informational gaps.
- We propose a dual-view graph representation method for fraud detection, adjusting the propagation of supervisory messages while enhancing the connection between critical information that determines node identity and labels.
- Comprehensive experiments demonstrate that MimbFD significantly mitigates the effect of imbalance and improves fraud detection performance.

2 Related Work

2.1 Imbalanced Learning On Graphs

The direct application of GNNs on imbalanced datasets results in outcomes that are biased toward the majority class, which is why most solutions [Qin *et al.*, 2025; Fu *et al.*, 2023] to unbalanced problems in graph learning focus on addressing class imbalance. Common methods include generating nodes and reweighting techniques. Generative node methods primarily focus on generating high-quality minority class samples to enhance the learning of the minority class [Park *et al.*, 2022; Zhao *et al.*, 2021; Li *et al.*, 2023; Qu *et al.*, 2021]. Reweighting [Song *et al.*, 2022], on the other hand, adjusts the classification boundaries of individual classes by designing specialized loss functions. Another significant area in graph imbalance learning is topological imbalance. [Ju *et al.*, 2023] focuses on the local topological imbalance of nodes and achieves enhancement of nodes by repairing low-degree nodes. ReNode [Chen *et al.*, 2021] is the first work to address the problem of global topological imbalance, addressing the issue of information conflict among nodes through adaptive reweighting. Similarly, [Sun *et al.*, 2022] advances the exploration of topological imbalance by optimizing information propagation paths via position-aware augmentation. However, these works lack an in-depth analysis of fraudster behavior. As a result, they are limited against camouflage. This limitation makes them difficult to apply directly to fraud detection.

2.2 Fraud Detection Based on GNNs

The powerful expressive capability of GNNs has garnered significant research interest. In recent years, numerous graph-based fraud detection methods have been proposed, which can generally be categorized into two main types based on their approach to adjusting neighborhood structures: local adjustment and high-order capture. **Local adjustment:** CARE-GNN [Dou *et al.*, 2020] leverages reinforcement learning to guide the identification of similar nodes within a neighborhood. FRAUDRE [Zhang *et al.*, 2021] emphasizes the differences between normal users and fraudsters in node pairs to strengthen the model’s ability to detect fraudulent behavior. However, these works mainly focus on the attributes of local neighborhoods and tend to ignore the critical information from distant nodes. **High-order capturing:** PC-GNN [Liu *et al.*, 2021] introduces a node sampler and a label-aware

neighbor selector, and additionally connects the fraudsters that are not directly linked to one another. GAGA [Wang *et al.*, 2023] employs group aggregation for both first-order and second-order neighborhoods to enhance node representation. COFD [Hu *et al.*, 2023] adjusts node neighborhoods by comparing first-order neighborhoods with second-order neighborhoods. While these methods consider the capture of high-order neighborhoods, they remain limited to small-scale mining of homophily to enhance node representations. The supervisory message available in the graph topology is not sufficiently leveraged in these works.

3 Preliminary

3.1 Definition

Fraud detection on multi-relation graph Given a graph $G = \{\mathcal{V}, \mathbf{X}, \{\mathcal{E}_r\}_{r=1}^R, \mathcal{Y}\}$, which $v_i \in \mathcal{V}$ denotes the node, $\mathbf{X}$ denotes the features of the node, $\mathcal{R}$ denotes the kind of relationship, $e_{i,j,r} \in \mathcal{E}$ denotes the two nodes v_i and v_j are connected under the relationship r . Graph-based fraud detection is a semi-supervised binary classification task. Only a small number of nodes are labeled (denoted by $\mathcal{Y}$) on the graph G . Each node has a binary label. $y_i = 0$ denotes a benign user, $y_i = 1$ denotes a fraudster.

3.2 Problem Analysis

To verify the existence of topological imbalance, we conducted data statistics on three public datasets. This helped us investigate the potential causes of the topological imbalance issue in fraud detection.

Topological Behavior Obfuscation Closeness centrality (CC) [Dong *et al.*, 2024] measures the average shortest path distance from a given node to all other nodes in the network, reflecting the speed at which the message spreads from that node to others. As illustrated in Fig. 2, we analyze the distribution of fraudulent and benign neighbors for benign users with different CC values. A higher CC value indicates faster user information exchange in the social network.

On YelpChi and Comp, benign users with lower CC tend to have more fraudster neighbors. On Amazon, as CC increases, the number of both types of neighbors also increases, but the difference between their distributions becomes larger. It suggests that benign users with lower CC are more likely to be surrounded by fraudsters, limiting the spread of fraud-related signals to a small area. In contrast, benign users with higher CC have more benign neighbors, and the number of fraudster neighbors grows more slowly, making it harder for fraud-related signals to spread. Overall, fraudsters avoid connecting too much with core benign users to prevent clear fraud patterns in the network structure. As a result, their activities are often limited to areas where information spreads slowly. This makes it hard for fraud-related signals to reach wider regions, contributing to topological imbalance.

Identity Feature Concealment Degree centrality (DC) [Esfandiari and Fakhrahmad, 2024] quantifies the importance of a node in the network based on the number of direct connections it has with other nodes, indicating its direct influence and connectivity. As shown in Fig. 3, we analyze the distribution of fraudulent neighbors and benign neighbors of benign

users with different DC values. A higher DC indicates more user connections within the network.

Fraudsters are mostly found near benign users with low DC, while benign users with high DC are surrounded by more benign neighbors. Even when a fraudster connects to the benign users with high DC, they can still hide their identity effectively, as their information become diluted. As a result, the spread of fraud-related information remains limited in scope, contributing to topological imbalance.

4 Methodology

Our goal is to address the challenges in supervised message passing for fraud detection, which are impacted by both topological imbalance and class imbalance. For topological imbalance, we introduce the topological message reachability (TMR) module, which enables global topology sensing to adjust the impact of long-distance messages received by nodes fully. To tackle class imbalance, we propose the local confounding debiasing module (LCD), which amplifies the weights of label-related variables while mitigating the discrimination bias introduced by the neighborhood. The overall architecture of MimbFD is illustrated in Fig. 4.

4.1 Topological Message Reachability

To adaptively adjust the information received by nodes regarding the supervisory message of different classes of identity nodes at a distance. This adjustment enhances node connections within the same class and improves the ability to discriminate between benign users and fraudsters. To enable the unknown labeled nodes to adjust the supervisory information from two distinct classes of nodes, we first capture the diffuse impact of the supervisory message from the labeled nodes on the global topology. Thus, we introduce Group PageRank [Chen *et al.*, 2020] to measure the propagation influence of labeled nodes on the topology. Specifically, under each relationship graph, the impact matrix corresponding to each class is calculated as follows:

$$g_{gpr}^r(c) = (1 - \alpha^r) A_r' g_{gpr}^r(c) + \alpha^r I_c^r, \quad (1)$$

where c represents the labeled categories 0 and 1 of the labeled nodes, α^r denotes the probability that class c is restarted by a randomized wandering under a relational subgraph; $A_r' = A_r D^{-1}$, A denotes the adjacency matrix of the subgraph of relation r and D is the degree matrix. Additionally, $I_c^r \in R^n$ is the transmission vector, defined as follows:

$$I_c^r = \begin{cases} \frac{1}{|\mathcal{V}_L^c|}, & \text{if } y_i = c \\ 0, & \text{otherwise} \end{cases} \quad (2)$$

where $\mathcal{V}_L^c$ is the number of nodes that have labeled class c . I_c^r denotes the supervised dissemination degree for benign ($c = 0$) and fraudulent ($c = 1$).

After calculating the corresponding information propagation matrix for each class, we connect them. The final impact propagation matrix is as follows:

$$G^r = \alpha^r \left(E - \left(1 - \alpha^r A_r' \right) \right)^{-1} I_r^*, \quad (3)$$

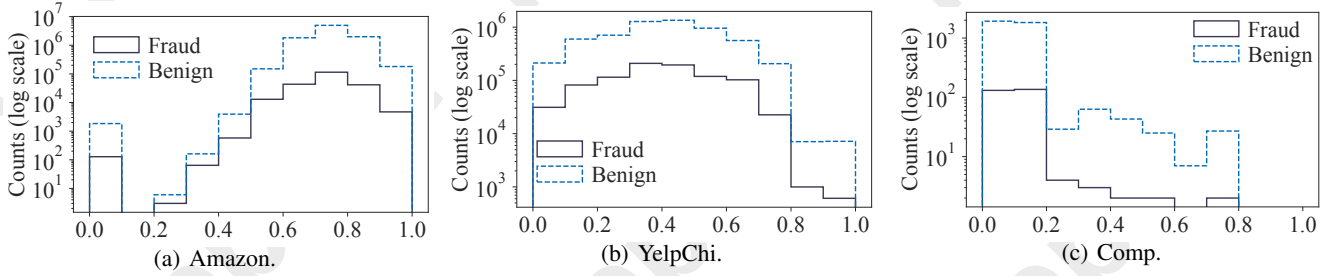


Figure 2: Neighborhood distributions under datasets with different closeness centrality.

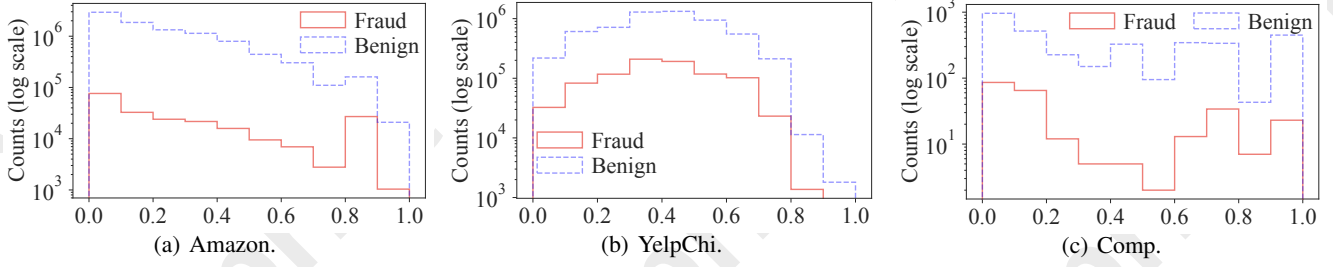


Figure 3: Neighborhood distributions under datasets with different degree centrality.

where E is the unit matrix and I_r^* is the composition of the collocation of all classes I_c^r . To some extent, the values in the matrix reflect the influence of the labeled nodes in the global topology. A higher value indicates a greater ability of the node to disseminate information throughout the topology.

To highlight the importance of nodes in the global topology message, we capture the neighborhood of the central node at a fine-grained level. For labeled neighbor nodes, weighted using normalized group PageRank values:

$$\mathbf{p}^r = \text{softmax}(G^r), \quad (4)$$

$$\mathbf{h}_{be,j}^{(l),r} = \sum_{v_j \in N_{be}(v_i)} \mathbf{p}_{be}^r \mathbf{h}_j^{(l-1),r}, \quad (5)$$

$$\mathbf{h}_{fr,j}^{(l),r} = \sum_{v_j \in N_{fr}(v_i)} \mathbf{p}_{fr}^r \mathbf{h}_j^{(l-1),r}, \quad (6)$$

where $\mathbf{h}_{be,j}^{(l),r}$ and $\mathbf{h}_{fr,j}^{(l),r}$ are weighted representations of benign users and fraudsters, respectively.

For nodes in the neighborhood with unknown labels, we must address the issue of biasing them towards receiving information from any specific class of nodes. After all, information propagation in a topology inevitably involves some conflicts. Therefore, we aim to distinguish the relationship between the unknown labeled nodes and the two types of known labeled nodes. Influence representations are assigned to unlabeled nodes as discriminative guidance, as follows:

$$\mathbf{h}_{un,be,j}^{(l),r} = \sum_{v_j \in N_{un}(v_i)} \mathbf{p}_{be}^r \mathbf{h}_j^{(l-1),r}, \quad (7)$$

$$\mathbf{h}_{un,fr,j}^{(l),r} = \sum_{v_j \in N_{un}(v_i)} \mathbf{p}_{fr}^r \mathbf{h}_j^{(l-1),r}, \quad (8)$$

where $\mathbf{h}_{un,be,j}^{(l),r}$ and $\mathbf{h}_{un,fr,j}^{(l),r}$ denote the influence representations of benign users and fraudsters on unlabeled nodes.

We then introduce a learnable adjustment factor to help the unknown labeled nodes adaptively combine these two representations in relation to themselves. The node’s own representation of these two identities carries selective associations. To restore the association between the nodes, the adjustment factor is computed as follows:

$$\beta = \sigma(\mathbf{h}_i^{(l-1),r}), \quad (9)$$

where $\sigma(\cdot)$ is the activation function. The final neighborhood representation of the unknown labeled node is obtained by a weighted combination of the two representations using the adjustment factor, as follows:

$$\mathbf{h}_{un,j}^{(l),r} = \beta \mathbf{h}_{un,fr,j}^{(l),r} + (1 - \beta) \mathbf{h}_{un,be,j}^{(l),r}. \quad (10)$$

After obtaining three different types of node representations in the neighborhood, we integrate them with the central node’s own representation. Simultaneously, the integration of the representations within each relationship subgraph is carried out. The final node representation is as follows:

$$\mathbf{h}_i^{(l)} = AGG_{all}^l \left[MLP(\mathbf{h}_i^{(l-1),r}), \mathbf{h}_{be,j}^{(l),r}, \mathbf{h}_{fr,j}^{(l),r}, \mathbf{h}_{un,j}^{(l),r} \right]_{r=1}^R. \quad (11)$$

4.2 Local Confounding Debiasing

Due to class imbalance, the node representations obtained through message passing contain a significant amount of information about normal users. Specifically, the node representation contains excessive confounding information, making it difficult for the model to identify key factors for determining identity. Thus, the model may favor the assumption

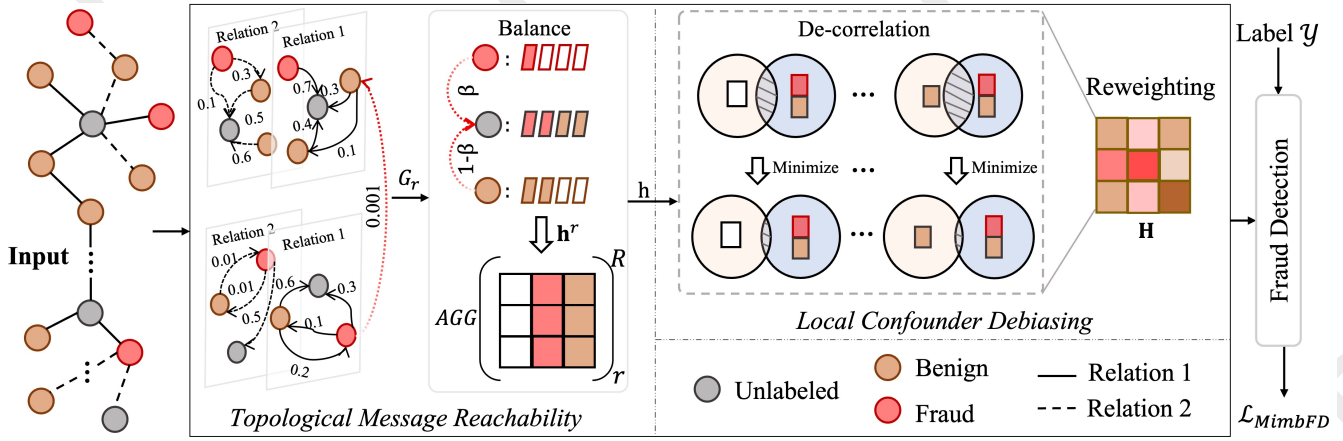


Figure 4: MimbFD consists of two modules: (1) TMR: Utilizing the balancing factor to assist nodes in adapting to the effects of the propagation of labeled nodes. (2) LCD: Analyze the variables in node representations to amplify the key message associated with the labels.

that the node is a “benign user.” According to stable learning [Kuang *et al.*, 2020], node representations can be categorized into stable and unstable variables. Stable variables are those key factors that determine the labeling categories, regardless of the environment. Therefore, we aim to perform a de-correlation between variables in the node representation, identifying stabilizing variables that will strengthen the connection between the nodes and their corresponding labels.

For the node representation, we employ the de-correlation term [Fan *et al.*, 2024] to ensure that the node variables are as independent from each other as possible. This helps alleviate spurious correlations between variables and labels. It is defined as follows:

$$\begin{aligned} \mathcal{L}_{LCD} = & \sum_{j=1}^p (\gamma^T \cdot \mathbf{h}_{i,m}^T |\Lambda_{\mathbf{w}} \mathbf{h}_{i,-m} / n \\ & - \mathbf{w} / n \cdot \mathbf{h}_{i,-m}^T \mathbf{w} / n|)^2 \\ & + \frac{\lambda_1}{n} \sum_{i=1}^n \mathbf{w}_i^2 + \lambda_2 \left(\frac{1}{n} \sum_{i=1}^n \mathbf{w}_i - 1 \right)^2, \end{aligned} \quad (12)$$

s.t. $\mathbf{w} \geq 0$

where $\mathbf{h}_{i,m}$ denotes the node representation m of the variables in $\mathbf{h}_i$, and $\mathbf{h}_{i,-m}$ is the node representation $\mathbf{h}_i$ of all the variables in $\mathbf{h}_i$ except m . The first term represents the computation of the differences between different variables in the node representation and applies weights γ to these variables to amplify their association with the true identity of the node. The second term indicates stability by enforcing regularization. The third serves to prevent the sample weights from going to zero. The s.t. $\mathbf{w} \geq 0$ is a non-negative constraint. To reduce computational complexity, this regularization is used at the last layer of the linear transformation.

Overall, the significance of de-correlation lies in analyzing the relationships between pairs of variables and reweighting each variable until the objective is minimized. This process results in better tuning of the node representation, strengthening its intrinsic connection with the labels.

4.3 Training

After iterative multiple layers, we take the output of the last layer as the final representation of the node $\mathbf{H}_i = \mathbf{h}_i^L$. We adopt the cross-entropy loss as the loss function for the classification result, which is defined as follows:

$$\mathcal{L}_{GNN} = \sum_{v \in V} -\log(y_v \cdot \sigma(\text{MLP}(\mathbf{H}_i))). \quad (13)$$

Combined with the loss function in the local confounding debiasing module, we define the loss of MimbFD as:

$$\mathcal{L}_{MimbFD} = \mathcal{L}_{GNN} + \eta \mathcal{L}_{LCD}, \quad (14)$$

where η is the balance parameter.

5 Experiments

In this section, we conduct experiments to verify the validity of MimbFD. Specifically, we aim to answer the following research questions:

RQ1: How does our model performance compare to existing state-of-the-art baselines? **RQ2:** How do the modules of TMR and LCD benefit the prediction? **RQ3:** How does the MimbFD perform with different hyperparameters? **RQ4:** Is the MimbFD able to find fraudsters while still maintaining close ties within the two groups? **RQ5:** Can the MimbFD capture supervisory messages about fraudsters across different levels of imbalance settings?

5.1 Experimental Setup

Dataset. Three multi-relation graph fraud datasets, YelpChi (Rayana and Akoglu 2015), Amazon (McAuley and Leskovec 2013), and Comp [Wu *et al.*, 2023] are used.

Baseline. We compare with several state-of-the-art GNN-based methods: (1) Traditional GNNs: GCN [Kipf and Welling, 2017], GAT [Velickovic *et al.*, 2017], and GraphSAGE [Hamilton *et al.*, 2017]. (2) Topological imbalance models: ReNode [Chen *et al.*, 2021], and TAM [Song *et al.*, 2022]. (3) General graph-based FD: CARE-GNN [Dou *et al.*, 2020], GAGA [Wang *et al.*, 2023], and COFD [Hu *et al.*,

Model	Amazon			YelpChi			Comp		
	AUC	Rec	F1	AUC	Rec	F1	AUC	Rec	F1
GCN	85.83±0.06	75.15±0.20	55.26±0.96	52.23±0.28	54.20±1.37	55.96±0.31	55.03±0.11	50.00±0.0	47.24±0.0
GAT	81.02±0.49	54.88±0.31	64.64±0.38	55.40±1.04	48.00±0.0	46.14±0.0	55.03±0.0	50.00±0.0	47.24±0.0
GraphSAGE	73.67±0.17	68.23±0.23	56.96±0.82	50.82±1.13	50.08±0.30	51.86±0.88	53.81±0.16	45.69±0.56	46.65±0.51
ReNode	40.52±27.39	25.68±11.2	59.7±3.9	-	-	-	51.18±3.62	43.42±16.66	43.2±6.3
TAM	50.19±0.88	59.83±0.8	61.60±0.75	57.28±0.24	53.75±0.24	53.98±0.35	53.08±0.23	50.37±0.14	48.92±0.47
CARE-GNN	93.49±0.36	89.13±0.10	88.83±0.22	76.85±0.07	70.36±0.03	60.99±0.10	54.16±0.52	46.84±0.4	46.78±0.19
GAGA	95.27±0.24	87.95±0.70	91.32±0.60	92.99±0.42	80.95±1.18	80.15±0.83	59.87±0.69	<u>55.42±4.25</u>	45.32±2.02
COFD	94.84±1.10	<u>89.52±1.30</u>	88.64±1.42	81.95±5.45	71.28±1.03	70.62±2.49	53.18±2.00	50.86±0.80	49.51±1.44
PC-GNN	94.92±0.45	83.36±0.39	86.26±0.47	80.90±0.09	78.11±0.27	67.36±0.43	56.66±0.35	44.03±3.88	46.80±1.55
FRAUDER	93.27±0.62	88.21±1.36	89.24±0.35	72.84±2.78	62.09±4.56	59.32±2.18	51.37±1.22	52.09±2.09	48.16±0.93
ConsisGAD	93.20±0.32	75.29±0.69	89.96±0.69	90.54±0.47	62.46±1.64	76.36±0.15	53.79±0.84	15.00±4.21	<u>51.42±0.20</u>
MimbFD	96.62±0.04	91.83±0.01	91.77±0.51	<u>92.41±0.19</u>	81.55±1.05	<u>80.08±0.31</u>	65.12±0.17	61.18±0.35	52.34±1.13

Table 1: Performance Comparison on Amazon, YelpChi, and Comp. “-”: out of memory; Bold: the best of baselines; Underline: runner-up.

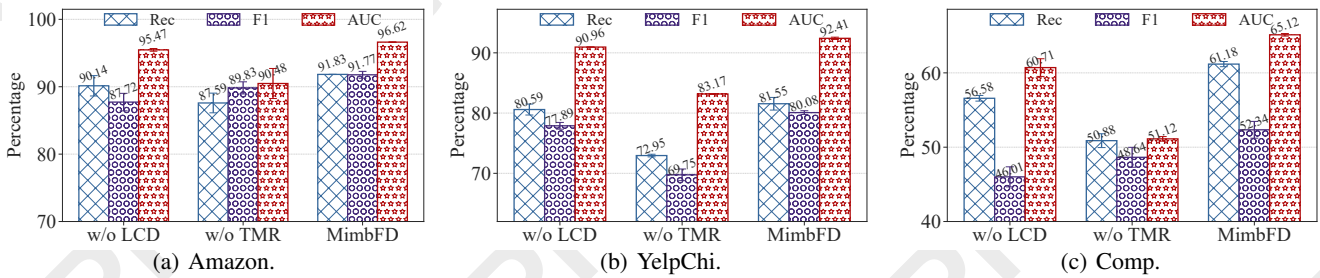


Figure 5: The ablation analysis on Amazon, YelpChi, and Comp.

2023]. (4) Class imbalance FD: PC-GNN [Liu *et al.*, 2021], FRAUDER [Zhang *et al.*, 2021], and ConsisGAD [Chen *et al.*, 2024].

Experimental Setting. In our experiments, Adam is chosen as the optimizer. We implement our method through Pytorch and DGL. For the dataset split, we divide it into three parts: training set, validation set, and test set, with ratios of 4:2:4, respectively, following common settings. For GCN, GAT, GraphSage, and ReNode, we convert multi-relational graphs to isomorphic graphs as input. For fraud detection methods, we use publicly available source code and input multi-relational graphs.

Evaluation Metric. Given the imbalance in the dataset, three widely used metrics, AUC, Recall, and F1-score are used.

5.2 Performance (RQ1)

To answer RQ1, we compare our method with the state-of-the-art approaches. The results are shown in Tab. 1. Due to the homophily assumption, the topological methods poorly handle fraud supervisory signals, hindering balanced information propagation. The relatively lower performance of FRAUDER and ConsisGAD, which focus on class imbalance, as well as CARE-GNN, a general graph-based fraud detector, can be attributed to their emphasis on local tuning. This highlights the fact that learning node representations using only first-order neighborhood information is insufficient to support effective model judgments. In contrast, fraud detectors like PC-GNN, which address class imbalance, and

general graph-based fraud detectors like GAGA and COFD, are more competitive, suggesting that the consideration of long-distance information is crucial.

Our method almost outperforms the baseline performance in terms of overall performance. Its strong performance on Comp, which contains more conflicting supervisory signals, highlights its effectiveness in balancing supervision coverage and amplifying fraud-related information.

5.3 Ablation Analysis (RQ2)

To answer RQ2, we conduct a series of ablation studies on certain modules in MimbFD on three datasets, as shown in Fig. 5. The variants tested are: w/o LCD and w/o TMR. The w/o TMR module focuses on distinguishing between the information in node representations and can capture attempts by fraudsters to obfuscate messages that reveal their identity. This helps alleviate the class imbalance problem in fraud detection to some extent. On the other hand, the w/o LCD module, which captures global supervisory information across the complex network structure, enables adaptive learning of supervisory messages by the nodes. It has been shown to mitigate the issue of uneven signal propagation caused by the fraudster’s long-distance camouflage.

MimbFD outperforms both variants across all metrics. The performance of both variants on Comp confirms that long-distance supervisory information provides critical signals for effective identification.

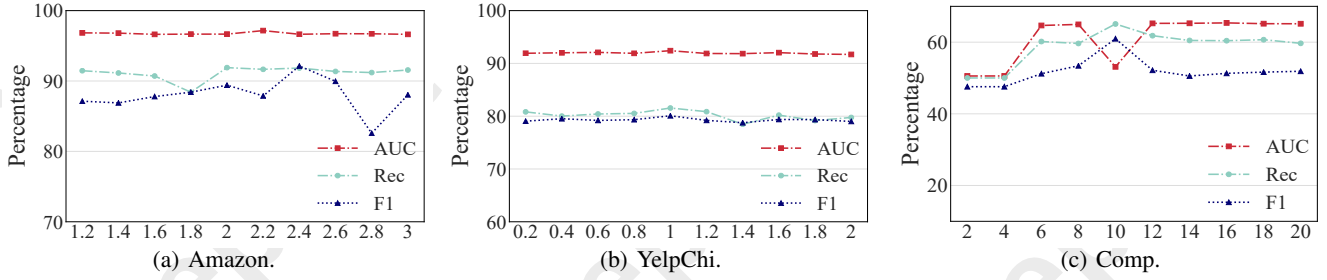


Figure 6: The hyperparameter analysis on Amazon, YelpChi, and Comp.

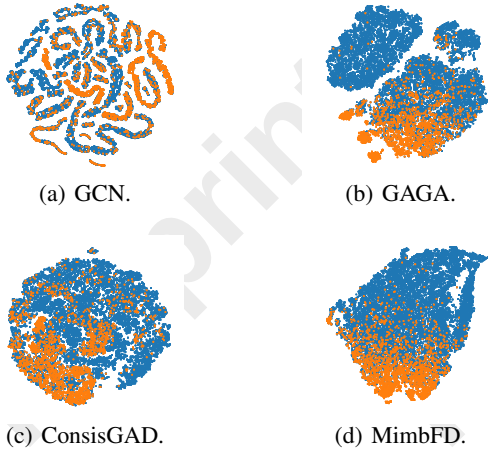


Figure 7: Visualization analysis on YelpChi.

5.4 Hyperparameter Analysis (RQ3)

To answer RQ3, we evaluate the impact of the hyperparameter η , which balances the interaction between the TMR and LCD modules, on the performance of MimbFD across different datasets, as shown in Fig. 6. The results show that all metrics are relatively stable on the Amazon and YelpChi datasets, whereas the metrics fluctuate more on the Comp dataset. This is due to the smaller data volume of Comp, where the supervisory information received by the nodes leads to significant conflicts, thus requiring more training for the LCD module.

Overall, the trends are roughly consistent across the datasets. When η is small, the model struggles to capture key signals in node representations, resulting in lower performance. As η increases, the model trains more effectively, leading to improved judgment. However, when η becomes too large, excessive bias correction occurs, as the model places too much emphasis on addressing weak signals while neglecting the connection between the true key information and the node identity.

5.5 Visualization (RQ4)

To answer RQ4, we visualize the node embeddings of different methods on the YelpChi dataset. The results are shown in Fig. 7 and the appendix. GCN only considers the homophily problem in node classification, which is insufficient for fraud

ρ	GCN	GAGA	ConsisGAD	MimbFD
5	79.12 $\pm$ 3.42	84.39 $\pm$ 6.86	75.52 $\pm$ 1.15	95.90$\pm$0.38
10	82.17 $\pm$ 0.65	87.95 $\pm$ 0.70	76.57 $\pm$ 1.11	91.83$\pm$0.01
20	81.01 $\pm$ 0.54	84.62 $\pm$ 1.48	76.16 $\pm$ 0.55	92.13$\pm$0.66

Table 2: The case study on Amazon. Bold: the best of the baselines.

detection, leading to poor separation results. GAGA is more competitive, achieving significant inter-class separation. This confirms that leveraging distant neighbors to exploit high homophily to resist camouflage is effective. ConsisGAD addresses the class imbalance problem in fraud detection, but does not fully utilize information from various types of nodes.

In contrast, our method generates more distinct separation effects and better intra-class cohesion. This demonstrates that MimbFD effectively adjusts the supervisory information within the global topology and amplifies the capture of key signals from fraudsters, allowing it to uncover the fraudster’s camouflaged behavior and identity more efficiently.

5.6 Case Study (RQ5)

To answer RQ5, we use Recall as the evaluation metric, as shown in Tab. 2. We set up three imbalance scenarios based on the original Amazon dataset, where about 10% are fraudsters. We adjusted the ratio of fraudsters to benign users, ρ , to 1 : 5 and 1 : 20 ($\rho = 5$ and $\rho = 20$) to represent moderately to extremely imbalanced data.

The results show that modelname performs well across all imbalance settings. At $\rho = 5$, where the two classes are easily confused, modelname achieves the best performance, highlighting its strength in mitigating the bottleneck of limited supervision transmission. Other baselines perform poorly at $\rho = 20$ because they cannot enhance their representations by obtaining long-distance fraud information.

6 Conclusions

We propose MimbFD to tackle fraud information imbalance from a dual perspective. For topological imbalance, we balance supervisory message reception to maximize information utilization. For class imbalance, we enhance identity gain by extracting key signals from node representations. Extensive experiments demonstrate that our model effectively mitigates information imbalance.

Acknowledgments

The corresponding authors are Li-e Wang and Xingcheng Fu. This work is supported in part by the National Natural Science Foundation of China (Nos. 62262003, 62462007 and U21A20474), the Guangxi Science and technology project (Guike AA22068070), the Guangxi Bagui Youth Talent Training Program, Guangxi Collaborative Innovation Center of Multisource Information Integration and Intelligent Processing and the Key Lab of Education Blockchain and Intelligent Technology, Ministry of Education (EBME24-01).

Contribution Statement

Co-first authors are Yudan Song and Yuecen Wei, marked with †.

References

- [Chen *et al.*, 2020] Deli Chen, Yanyai Lin, Lei Li, Xuancheng Ren, Peng Li, Jie Zhou, and Xu Sun. Distance-wise graph contrastive learning, 2020.
- [Chen *et al.*, 2021] Deli Chen, Yankai Lin, Guangxiang Zhao, Xuancheng Ren, Peng Li, Jie Zhou, and Xu Sun. Topology-imbalance learning for semi-supervised node classification. In *NeurIPS*, pages 29885–29897, 2021.
- [Chen *et al.*, 2024] Nan Chen, Zemin Liu, Bryan Hooi, Bingsheng He, Rizal Fathony, Jun Hu, and Jia Chen. Consistency training with learnable data augmentation for graph anomaly detection with limited supervision. In *ICLR*, 2024.
- [Dhawan *et al.*, 2019] Sarthika Dhawan, Siva Charan Reddy Gangireddy, Shiv Kumar, and Tanmoy Chakraborty. Spotting collective behaviour of online frauds in customer reviews. In *IJCAI*, pages 245–251, 2019.
- [Dong *et al.*, 2024] Zhao Dong, Yuanzhi Duan, Yue Zhou, Shukai Duan, and Xiaofang Hu. Weight-adaptive channel pruning for cnns based on closeness-centrality modeling. *Appl. Intell.*, 54(11-12):201–215, 2024.
- [Dou *et al.*, 2020] Yingdong Dou, Zhiwei Liu, Li Sun, Yutong Deng, Hao Peng, and Philip S. Yu. Enhancing graph neural network-based fraud detectors against camouflaged fraudsters. In *CIKM*, pages 315–324, 2020.
- [Esfandiari and Fakhrahmad, 2024] Shima Esfandiari and Seyed Mostafa Fakhrahmad. Predicting node influence in complex networks by the k-shell entropy and degree centrality. In *WWW*, pages 629–632, 2024.
- [Fan *et al.*, 2024] Shaohua Fan, Xiao Wang, Chuan Shi, Kun Kuang, Nian Liu, and Bai Wang. Debiased graph neural networks with agnostic label selection bias. *IEEE Trans. Neural Networks Learn. Syst.*, 35(4):4411–4422, 2024.
- [Fu *et al.*, 2023] Xingcheng Fu, Yuecen Wei, Qingyun Sun, Haonan Yuan, Jia Wu, Hao Peng, and Jianxin Li. Hyperbolic geometric graph representation learning for hierarchy-imbalance node classification. In *WWW*, pages 460–468. ACM, 2023.
- [Fu *et al.*, 2025] Xingcheng Fu, Jian Wang, Yisen Gao, Qingyun Sun, Haonan Yuan, Jianxin Li, and Xianxian Li. Discrete curvature graph information bottleneck. In *AAAI*, pages 16666–16673, 2025.
- [Goswami *et al.*, 2017] Kunal Goswami, Younghee Park, and Chungsik Song. Impact of reviewer social interaction on online consumer review fraud detection. *J. Big Data*, 4:15, 2017.
- [Hamilton *et al.*, 2017] William L. Hamilton, Zhitaoying, and Jure Leskovec. Inductive representation learning on large graphs. In *NIPS*, pages 1024–1034, 2017.
- [Hao *et al.*, 2024] Xiaorong Hao, Bo Liu, Xinyan Yang, Xianguo Sun, Qing Meng, and Jiuxin Cao. Multi-stage dynamic disinformation detection with graph entropy guidance. *World Wide Web*, 27(2):1–21, 2024.
- [Hu *et al.*, 2023] Jinzhang Hu, Ruimin Hu, Zheng Wang, Dengshi Li, Junhang Wu, Lingfei Ren, Yilong Zang, Zijun Huang, and Mei Wang. Collaborative fraud detection: How collaboration impacts fraud detection. In *ACM Multimedia*, pages 8891–8899, 2023.
- [Ju *et al.*, 2023] Mingxuan Ju, Tong Zhao, Wenhao Yu, Neil Shah, and Yanfang Ye. Graphpatcher: Mitigating degree bias for graph neural networks via test-time augmentation. In *NeurIPS*, 2023.
- [Kipf and Welling, 2017] Thomas N. Kipf and Max Welling. Semi-supervised classification with graph convolutional networks. In *ICLR (Poster)*, 2017.
- [Kuang *et al.*, 2020] Kun Kuang, Ruoxuan Xiong, Peng Cui, Susan Athey, and Bo Li. Stable prediction with model misspecification and agnostic distribution shift. In *AAAI*, pages 4485–4492, 2020.
- [Li *et al.*, 2023] Wen-Zhi Li, Chang-Dong Wang, Hui Xiong, and Jian-Huang Lai. Graphsha: Synthesizing harder samples for class-imbalanced node classification. In *KDD*, pages 1328–1340, 2023.
- [Liu *et al.*, 2021] Yang Liu, Xiang Ao, Zidi Qin, Jianfeng Chi, Jinghua Feng, Hao Yang, and Qing He. Pick and choose: A gnn-based imbalanced learning approach for fraud detection. In *WWW*, pages 3168–3177, 2021.
- [Park *et al.*, 2022] Joonhyung Park, Jaeyun Song, and Eunho Yang. Graphens: Neighbor-aware ego network synthesis for class-imbalanced node classification. In *ICLR*, 2022.
- [Peng *et al.*, 2023] Hao Peng, Ruitong Zhang, Shaoning Li, Yuwei Cao, Shirui Pan, and Philip S. Yu. Reinforced, incremental and cross-lingual event detection from social messages. *IEEE Trans. Pattern Anal. Mach. Intell.*, 45(1):980–998, 2023.
- [Qiao *et al.*, 2024] Hezhe Qiao, Hanghang Tong, Bo An, Irwin King, Charu C. Aggarwal, and Guansong Pang. Deep graph anomaly detection: A survey and new perspectives, 2024.
- [Qin *et al.*, 2025] Jiawen Qin, Haonan Yuan, Qingyun Sun, Lyujin Xu, Jiaqi Yuan, Pengfeng Huang, Zhaonan Wang, Xingcheng Fu, Hao Peng, Jianxin Li, and Philip S. Yu.

- Igl-bench: Establishing the comprehensive benchmark for imbalanced graph learning. In *ICLR*, 2025.
- [Qu *et al.*, 2021] Liang Qu, Huaisheng Zhu, Ruiqi Zheng, Yuhui Shi, and Hongzhi Yin. Imgagn: Imbalanced network embedding via generative adversarial graph networks. In *KDD*, pages 1390–1398, 2021.
- [Song *et al.*, 2022] Jaeyun Song, Joonhyung Park, and Eunho Yang. TAM: topology-aware margin loss for class-imbalanced node classification. In *ICML*, volume 162 of *Proceedings of Machine Learning Research*, pages 20369–20383, 2022.
- [Song *et al.*, 2024] Yudan Song, Yuecen Wei, Haonan Yuan, Qingyun Sun, Xingcheng Fu, Li-e Wang, and Xianxian Li. Causalfd: causal invariance-based fraud detection against camouflaged preference. *Int. J. Mach. Learn. Cybern.*, 15(11):5053–5070, 2024.
- [Sun *et al.*, 2022] Qingyun Sun, Jianxin Li, Haonan Yuan, Xingcheng Fu, Hao Peng, Cheng Ji, Qian Li, and Philip S. Yu. Position-aware structure learning for graph topology-imbalance by relieving under-reaching and over-squashing. In *CIKM*, pages 1848–1857, 2022.
- [Velickovic *et al.*, 2017] Petar Velickovic, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph attention networks, 2017.
- [Wang *et al.*, 2023] Yuchen Wang, Jinghui Zhang, Zhengjie Huang, Weibin Li, Shikun Feng, Ziheng Ma, Yu Sun, Dianhai Yu, Fang Dong, Jiahui Jin, Beilun Wang, and Junzhou Luo. Label information enhanced fraud detection against low homophily in graphs. In *WWW*, pages 406–416, 2023.
- [Wei *et al.*, 2024] Yuecen Wei, Haonan Yuan, Xingcheng Fu, Qingyun Sun, Hao Peng, Xianxian Li, and Chunming Hu. Poincaré differential privacy for hierarchy-aware graph embedding. In *AAAI*, pages 9160–9168, 2024.
- [Wei *et al.*, 2025] Yuecen Wei, Xingcheng Fu, Lingyun Liu, Qingyun Sun, Hao Peng, and Chunming Hu. Prompt-based unifying inference attack on graph neural networks. In *AAAI*, pages 12836–12844, 2025.
- [Wu *et al.*, 2023] Bin Wu, Xinyu Yao, Boyan Zhang, Kuo-Ming Chao, and Yinsheng Li. Splitgcn: Spectral graph neural network for fraud detection against heterophily. In *CIKM*, pages 2737–2746, 2023.
- [Xu *et al.*, 2021] Bingbing Xu, Huawei Shen, Bing-Jie Sun, Rong An, Qi Cao, and Xueqi Cheng. Towards consumer loan fraud detection: Graph neural networks with role-constrained conditional random field. In *AAAI*, pages 4537–4545, 2021.
- [Zhang *et al.*, 2021] Ge Zhang, Jia Wu, Jian Yang, Amin Beheshti, Shan Xue, Chuan Zhou, and Quan Z. Sheng. FRAUDRE: fraud detection dual-resistant to graph inconsistency and imbalance. In *ICDM*, pages 867–876, 2021.
- [Zhao *et al.*, 2021] Tianxiang Zhao, Xiang Zhang, and Suhang Wang. Graphsmote: Imbalanced node classification on graphs with graph neural networks. In *WSDM*, pages 833–841, 2021.