

# Robust Graph Learning on the Web: Challenges, Methods, and Applications

Xiang Ao\*  
Institute of Computing  
Technology, Chinese  
Academy of Sciences  
Beijing, China  
aoxiang@ict.ac.cn

Yang Liu\*  
Institute of Computing  
Technology, Chinese  
Academy of Sciences  
Beijing, China  
liuyang2023@ict.ac.cn

Guansong Pang  
Singapore Management  
University  
Singapore  
gpsang@smu.edu.sg

Yuanhao Ding\*  
Institute of Computing  
Technology, Chinese  
Academy of Sciences  
Beijing, China  
dingyuanhao19@mails.ucas.ac.cn

Hezhe Qiao  
Singapore Management  
University  
Singapore  
hezheqiao@gmail.com

Dawei Cheng  
Tongji University  
Shanghai, China  
dcheng@tongji.edu.cn

Qing He\*  
Institute of Computing  
Technology, Chinese  
Academy of Sciences  
Beijing, China  
heqing@ict.ac.cn

## Abstract

Graph learning is transforming web intelligence, powering applications from recommender systems to anomaly detection. However, most existing approaches implicitly assume ideal conditions where training and testing data are accurate, complete, and free from manipulation. In reality, web environments rarely exhibit such stability. Dynamic user behavior, incomplete or outdated content, adversarial interference, and sudden distribution shifts can all erode the reliability of even state-of-the-art models, leading to biased or unsafe outcomes. This tutorial provides a comprehensive survey of emerging strategies for robust graph learning on the web. We first present a structured taxonomy of the principal robustness threats specific to web contexts. Next, we categorize current robust graph learning approaches, spanning data-level preprocessing to model-level adaptation and generalization, and discuss representative models in detail. We then showcase real-world case studies illustrating how robustness challenges emerge and how targeted methods can mitigate them in the web system. This tutorial offers researchers, engineers, and platform developers actionable strategies to safeguard graph-based AI in dynamic, high-impact web environments. The website is available at here<sup>1</sup>.

## Keywords

Robust learning, Graph learning, AI on the Web

### ACM Reference Format:

Xiang Ao, Yang Liu, Guansong Pang, Yuanhao Ding, Hezhe Qiao, Dawei Cheng, and Qing He. 2026. Robust Graph Learning on the Web: Challenges,

Methods, and Applications. In *Companion Proceedings of the ACM Web Conference 2026 (WWW Companion '26)*, April 13–17, 2026, Dubai, United Arab Emirates. ACM, New York, NY, USA, 4 pages. <https://doi.org/10.1145/3774905.3793923>

## 1 Introduction

Web data, characterized by diverse modalities and intricate interconnections, are naturally represented as graphs, which capture both the attributes of individual webpages or users and the complex relationships among them. Graph learning provides a principled framework for modeling such rich relational structures on the web, where entities such as users, websites, hyperlinks, and interactions naturally form interconnected graphs. In recommendation systems, user-item graphs capture behavioral dependencies that improve personalization beyond independent-sample models. In web search and ranking, hyperlink graphs enable relevance estimation with global structural awareness. Similar advantages arise in social network analysis through friendship or interaction graphs, misinformation detection via content-source networks, and knowledge extraction from heterogeneous web corpora. These examples demonstrate that graph learning is both theoretically grounded and practically effective in modeling complex web ecosystems.

Despite the remarkable success of graph learning in web-related applications, the inherent complexity, scale, and dynamics of the online environment often yield graphs that are heterogeneous, noisy, rapidly evolving, and susceptible to manipulation. Data irregularities arise from incomplete or outdated webpages, user behavior drift, and platform-specific content biases. Adversarial actors may create fake links, bots, or misleading associations, as seen in spam networks or coordinated misinformation campaigns. Out-of-distribution (OOD) shifts occur during major events or algorithmic updates, disrupting assumptions underlying pre-trained models.

Robust graph learning has emerged as a critical research frontier for ensuring the reliability of graph-based models in complex, large-scale web environments. Its core objective is to develop learning algorithms that maintain predictive accuracy and stability in the presence of structural noise, adversarial perturbations, and out-of-distribution (OOD) shifts, which are pervasive in real-world

\*State Key Lab of AI Safety, Institute of Computing Technology, Chinese Academy of Sciences (CAS), Beijing, China.

<sup>1</sup><https://qwer12191.github.io/Robust-Graph-Learning-on-Web>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

WWW Companion '26, Dubai, United Arab Emirates

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2308-7/2026/04

<https://doi.org/10.1145/3774905.3793923>

web systems. We introduce a hierarchical taxonomy of robust graph learning approaches, organized from data-level to model-level robustness methods. At the data level, methods such as graph anomaly detection [12, 20] identify and mitigate the impact of corrupted nodes, edges, or features—such as spam accounts, fake links, or misleading content—resulting in cleaner graph structures for downstream learning. At the model level, adversarial defense techniques [17, 18, 44] leverage graph augmentation, feature smoothing, and adversarial regularization to enhance resistance against targeted manipulations like link spamming or coordinated misinformation. For distributional robustness, graph invariant learning [19, 30] extracts stable, domain-agnostic substructures to support generalization under dynamic web conditions, such as platform updates or user behavior drift. Throughout the tutorial, we interleave these methodological advances with web-centric case studies, including recommendation systems that remain reliable under bot activity, misinformation detection pipelines that adapt to evolving narratives, and knowledge graphs that sustain performance despite incomplete or inconsistent relational data. These examples demonstrate both the practical necessity and the methodological richness of robust graph learning in the web domain.

This tutorial provides a comprehensive review of robust graph learning methods, emphasizing their adaptation to web data and deployment in practice. We target researchers and practitioners in web intelligence, information retrieval, social computing, and AI-driven content moderation who seek robustness-aware solutions to complex online challenges. Participants will gain actionable insights into algorithmic design, defense strategies, and real-world case studies, fostering cross-disciplinary collaboration toward trustworthy, high-impact graph intelligence for the open web.

## 2 Organizing Committee

**Dr. Xiang Ao** is a Professor in the State Key Laboratory of AI Safety, Institute of Computing Technology, Chinese Academy of Sciences (ICT, CAS). His research interests include algorithms and models for AI finance tasks, e.g. graph learning for financial applications and NLP for financial applications. His research works were honored with the Best Short Paper Honorable Mention at SIGIR2021 and CIKM 2022, and he was the first to pioneer the Event-Level Financial Sentiment Analysis (EFSA) task. He has authored more than 100 referred publications at prestigious international conferences and journals and has served as SPC or PC members over top tier international conferences such as KDD, WWW, IJCAI, AAAI, ACL, NeurIPS, ICML, etc.

**Dr. Yang Liu** is an Assistant Professor in the State Key Laboratory of AI Safety, Institute of Computing Technology, Chinese Academy of Sciences. His research interests include Financial Data Mining and Trustworthy Graph Machine Learning, among others. He has published over 20 referred papers on graph-based financial data mining at prestigious international conferences and journals and has served as PC members over top tier international conferences such as KDD, WWW, AAAI, NeurIPS, etc.

**Dr. Guansong Pang** is a tenure-track Assistant Professor of Computer Science and Lee Kong Chian Fellow at the School of Computing and Information Systems, Singapore Management University (SMU), where he leads the Machine Learning & Applications

(MaLA) Lab. His research interests include machine learning, data mining, and computer vision, with a research theme focused on recognizing and generalizing to abnormal/unknown/unseen data for creating trustworthy AI systems. His research has attracted 10,000+ citations and received multiple global recognition/awards, e.g., the prestigious 2020 UTS Chancellor's Award List, the World's Top 2% Scientists in 2022-2024, DSAA 2023 Best Paper Award, and Most Influential KDD 2023 Paper.

**Yuanhao Ding** is a Ph.D. candidate at the State Key Laboratory of AI Safety, Institute of Computing Technology, Chinese Academy of Sciences (ICT, CAS), under the supervision of Prof. Xiang Ao. He received B.S. degree in Computer Science and Technology from University of Chinese Academy of Sciences in 2023. His research interests broadly encompass Safety and Security of Machine Learning, with a particular focus on Graph Adversarial Learning.

**Hezhe Qiao** is a Ph.D. candidate at Singapore Management University. His research interests include graph representation learning, graph anomaly detection, and the combination of anomaly detection and large language models. He has published multiple articles in refereed top-tier venues, including NeurIPS, ICML, KDD, IJCAI, and TKDE. He has received several awards, such as the Mark Weiser Best Paper Award in PerCom (2024) and SMU Presidential Doctoral Fellowship Award (2024,2025), SMU Dean List (2025).

**Dr. Dawei Cheng** is an Associate Professor at the School of Computer Science and Technology, Tongji University, and serves as Assistant Dean of the National Collaborative Innovation Center for Internet Finance Security. His research focuses on graph learning, financial big data, deep learning on complex graphs, generative AI, and large-scale graph analytics. He has published over 100 papers in leading venues and serves as SPC/PC member for international conferences such as KDD, ICML, ICLR, NeurIPS, and WWW. He has received several honors, including the ACM Shanghai Rising Star Award and the WAIC Outstanding Paper Award.

**Dr. Qing He** is a Full Professor at the Institute of Computing Technology, Chinese Academy of Sciences(CAS), and he is a Professor of University of Chinese Academy of Sciences (UCAS). He is also the Vice Secretary of Chinese Association for Artificial Intelligence, the Member of China Computer Federation Artificial Intelligence and Pattern Recognition Committee, the Member of Chinese Institute of Electronics and Clouding Computing and Big Data Experts Committee. His interests include data mining, machine learning, classification, fuzzy clustering, cloud computing, and big data. More than 100 papers have been published in journals, 40 of which are SCI Indexed, 66 of which are EI Indexed.

## 3 Type, Schedule, Audience, and Materials

This half-day Lecture-style tutorial systematically guides participants from the theoretical foundations to the current research frontier of robust graph learning for web intelligence. It targets students, academic researchers, and industry practitioners in artificial intelligence, machine learning, and web-scale data analysis who aim to develop dependable, secure Graph Neural Network (GNN) models for web applications. All instructional materials including lecture slides, an annotated bibliography, and selected code with datasets will be made publicly available through a dedicated GitHub repository.

**Table 1: Classification of graph robustness methods.**

|                               | Category                           | Representative Methods                                                  |
|-------------------------------|------------------------------------|-------------------------------------------------------------------------|
| <b>Data-level Robustness</b>  | Heuristic Pre-processing           | Jaccard [29], SVD-GCN [4], MD-GNN [32], CARE-GNN [3]                    |
|                               | Graph Structure Learning           | GLNN [7], Pro-GNN [14], Gaug [39], SUBLIME [21], STABLE [18], RGSL [33] |
|                               | Graph Anomaly Detection            | DOMINANT [2], TAM [25], SI-HGAD [40], GGAD [27], NGS [28], AO-GNN [12]  |
| <b>Model-level Robustness</b> | Graph Adversarial Training         | RobustTrain [35], SAT [16], RobustDiff [8], GOOD-AT [17]                |
|                               | Robust Message-passing Mechanism   | GNNGuard [38], Advanced-GCN [18], Mid-GCN [11], F2GNN [10], DMP [36]    |
|                               | Out-of-Distribution Generalization | DisenGCN [23], PATTERN [37], EERM [30], DGNN [5], FLOOD [19]            |

### 3.1 The Landscape of Graph, Web, and the Fragility of Trust (30 mins)

This part will motivate the critical need for robustness on the web. We will begin by showcasing why GNNs are an exceptionally powerful tool for modeling the intricate webs of relationships on the web. We will then immediately pivot to the core problem: the inherent fragility of standard GNNs in high-stakes, adversarial environments [34]. This narrative will set the stage for the tutorial’s roadmap, introducing threat taxonomy and the two pillars of robustness: Data-Level Robustness and Model-Level Robustness.

### 3.2 A Taxonomy of Threats (30 mins)

- **Intentional Adversarial Attacks.** We focus on gray-box settings in which adversaries have partial knowledge of model inputs and craft subtle perturbations to graph structure or node/edge features to manipulate predictions or evade detection [43]. We will introduce representative methods like Nettack [42], TDGIA [41], and UGBA [1]. Examples include a fraudulent firm rewiring ownership or trading links to embed itself within reputable communities, inserting benign-looking intermediary counterparties to dilute suspicious connections, or split-transacting across accounts to slip rule-based thresholds while preserving topological stealth.
- **Inherent Data Imperfections.** A central shift in perspective is that a model’s adversary is not limited to human attackers: risk also arises from the intrinsic properties of web data. Common pathologies include noise [15], incompleteness [31], inconsistency (schema mismatches, stale attributes, temporal misalignment) [22], and class imbalance [20] endemic to fraud detection and rare-event modeling.

### 3.3 The Robustness Toolbox: A Hierarchical Defense Strategy (100 mins)

This section forms the technical core of the tutorial, presenting a hierarchical defense strategy with state-of-the-art methods for engineering resilient GNNs, as shown in Table 1.

At the **data** level, we will explore methods that sanitize the input graph itself, a process often termed graph purification or denoising. This ranges from heuristic pre-processing techniques leveraging intrinsic graph properties (homophily, low-rankness, feature smoothness) [4, 29, 32], to more advanced graph structure learning frameworks like Pro-GNN [14] and STABLE [18]. The

cutting edge here involves anomaly detection, which identifies and removes outliers nodes and edges in the graph that do not conform to the expected patterns or structure of the data [2, 5, 12].

At the **model** level, we will discuss fortifying the GNN’s architecture and training procedure to withstand worst-case perturbations and distribution shifts. We first introduce adversarial training, which hardens the model by optimizing on adversarially perturbed examples constructed to approximate worst-case attacks [6, 9, 24]. We also consider intrinsically robust architectures [10, 11, 18], especially defenses involves robust message-passing mechanisms like GNNGuard [38], which prune suspicious information flow, and recent generative approaches such as ASD-VAE [13]. Finally, because reliability in dynamic web environments demands resilience to unknown threats and concept drift, we ground these techniques in the broader literature on OOD generalization [19, 23, 30, 37], explicitly targeting robustness under distribution shift through both training objectives and architectural choices.

### 3.4 Future Research Directions (20 mins)

This section synthesizes the tutorial’s key insights and outlines a concrete research agenda for trustworthy AI, with emphasis on graph-centric modeling, robustness, and governance at production scale. We will discuss prioritized directions and actionable research questions that can guide near-term work and longer-term community efforts. The discussion will probe critical questions at the intersection of research and practice, such as:

- What is the role of Large Language Models (LLMs) in augmenting and interpreting robust graph learning methods?
- Beyond accuracy, how do we ensure fairness and mitigate bias in robust graph learning systems?
- What are the emerging regulatory and ethical hurdles for deploying advanced AI in the web?

## 4 Comparable Tutorials

Several related tutorials have been presented at previous AI conferences, covering topics such as graph anomaly detection and graph distribution shift. In contrast, this tutorial introduces a new series of events dedicated specifically to robust graph learning on the web, establishing a clear distinction from earlier tutorials by focusing on the unique challenges, methodologies, and deployment considerations on the web.

Qiao et al. [26] organized a tutorial "Deep learning for graph anomaly detection" in IJCAI 2025, which aims to identify rare observations in graphs. The tutorial presented a comprehensive introduction of DLGAD from three key technical perspectives, including GNN backbone, proxy task design, and graph anomaly measure. Besides, our proposed tutorial also covers other areas related to robustness like adversarial attacks and defenses.

Wang et al. organized a tutorial "Beyond Graph Distribution Shifts: LLMs, Adaptation, and Generalization" in IJCAI 2025, which presents a comprehensive overview of three emerging and synergistic directions for tackling distribution shifts in graph learning. Graph distribution shifts exist in various scenarios due to attribute noises, structural shifts, and robust learning is one of the possible solutions for tackling graph distribution shift.

NaderiAlizadeh et al. organized a tutorial "Graph Neural Networks: Architectures, Fundamental Properties and Applications" in AAAI 2025. They explore four fundamental properties of GNNs: equivariance to permutations, stability to deformations, transferability across scales, and generalization to unseen data. In contrast, our tutorial focuses on a specific graph learning setting, which is more applicable in the web systems.

## Acknowledgments

The tutorial is supported by the National Natural Science Foundation of China under Grant No. 62576333, 62406307, the Strategic Priority Research Program of the Chinese Academy of Sciences under Grant No. XDB0680201, the Beijing Natural Science Foundation under Grant No. JQ25015.

## References

- [1] Enyan Dai, Minhua Lin, Xiang Zhang, and Suhang Wang. 2023. Unnoticeable backdoor attacks on graph neural networks. In *WWW*. 2263–2273.
- [2] Kaize Ding, Jundong Li, Rohit Bhanushali, and Huan Liu. 2019. Deep anomaly detection on attributed networks. In *SDM*. SIAM, 594–602.
- [3] Yingdong Dou, Zhiwei Liu, Li Sun, Yutong Deng, Hao Peng, and Philip S Yu. 2020. Enhancing graph neural network-based fraud detectors against camouflaged fraudsters. In *CIKM*. 315–324.
- [4] Negin Entezari, Saba A Al-Sayouri, Amirali Darvishzadeh, and Evangelos E Papalexakis. 2020. All you need is low (rank) defending against adversarial attacks on graphs. In *Proceedings of the 13th international conference on web search and data mining*. 169–177.
- [5] Shaohua Fan, Xiao Wang, Chuan Shi, Kun Kuang, Nian Liu, and Bai Wang. 2022. Debiased graph neural networks with agnostic label selection bias. *IEEE transactions on neural networks and learning systems* 35, 4 (2022), 4411–4422.
- [6] Fuli Feng, Xiangnan He, Jie Tang, and Tat-Seng Chua. 2019. Graph adversarial training: Dynamically regularizing based on graph structure. *IEEE Transactions on Knowledge and Data Engineering* 33, 6 (2019), 2493–2504.
- [7] Xiang Gao, Wei Hu, and Zongming Guo. 2019. Exploring structure-adaptive graph learning for robust semi-supervised classification. *arXiv:1904.10146* (2019).
- [8] Lukas Gosch, Simon Geisler, Daniel Sturm, Bertrand Charpentier, Daniel Zügner, and Stephan Günnemann. 2023. Adversarial training for graph neural networks: Pitfalls, solutions, and new directions. *NeurIPS* 36 (2023), 58088–58112.
- [9] Lukas Gosch, Simon Geisler, Daniel Sturm, Bertrand Charpentier, Daniel Zügner, and Stephan Günnemann. 2023. Adversarial training for graph neural networks: Pitfalls, solutions, and new directions. *NeurIPS* 36 (2023), 58088–58112.
- [10] Guanghui Hu, Yang Liu, Qing He, and Xiang Ao. 2024. F2gnn: An adaptive filter with feature segmentation for graph-based fraud detection. In *ICASSP*.
- [11] Jincheng Huang, Lun Du, Xu Chen, Qiang Fu, Shi Han, and Dongmei Zhang. 2023. Robust mid-pass filtering graph convolutional networks. In *Proceedings of the ACM Web Conference 2023*. 328–338.
- [12] Mengda Huang, Yang Liu, Xiang Ao, Kuan Li, Jianfeng Chi, Jinghua Feng, Hao Yang, and Qing He. 2022. AUC-oriented graph neural network for fraud detection. In *Proceedings of the ACM web conference 2022*. 1311–1321.
- [13] Xinke Jiang, Zidi Qin, Jiarong Xu, and Xiang Ao. 2024. Incomplete graph learning via attribute-structure decoupled variational auto-encoder. In *WSDM*. 304–312.
- [14] Wei Jin, Yao Ma, Xiaorui Liu, Xianfeng Tang, Suhang Wang, and Jiliang Tang. 2020. Graph structure learning for robust graph neural networks. In *ACM SIGKDD*.
- [15] Zhao Kang, Haiqi Pan, Steven CH Hoi, and Zenglin Xu. 2019. Robust graph learning from noisy data. *IEEE transactions on cybernetics* 50, 5 (2019).
- [16] Jintang Li, Jiaying Peng, Liang Chen, Zibin Zheng, Tingting Liang, and Qing Ling. 2022. Spectral adversarial training for robust graph neural network. *IEEE TKDE* 35, 9 (2022), 9240–9253.
- [17] Kuan Li, YiWen Chen, Yang Liu, Jin Wang, Qing He, Minhao Cheng, and Xiang Ao. 2023. Boosting the Adversarial Robustness of Graph Neural Networks: An OOD Perspective. In *ICLR*.
- [18] Kuan Li, Yang Liu, Xiang Ao, Jianfeng Chi, Jinghua Feng, Hao Yang, and Qing He. 2022. Reliable representations make a stronger defender: Unsupervised structure refinement for robust gnn. In *ACM SIGKDD*. 925–935.
- [19] Yang Liu, Xiang Ao, Fuli Feng, Yunshan Ma, Kuan Li, Tat-Seng Chua, and Qing He. 2023. FLOOD: A flexible invariant learning framework for out-of-distribution generalization on graphs. In *ACM SIGKDD*. 1548–1558.
- [20] Yang Liu, Xiang Ao, Zidi Qin, Jianfeng Chi, Jinghua Feng, Hao Yang, and Qing He. 2021. Pick and choose: a GNN-based imbalanced learning approach for fraud detection. In *Proceedings of the web conference 2021*. 3168–3177.
- [21] Yixin Liu, Yu Zheng, Daokun Zhang, Hongxu Chen, Hao Peng, and Shirui Pan. 2022. Towards unsupervised deep graph structure learning. In *Proceedings of the ACM Web Conference 2022*. 1392–1403.
- [22] Zhiwei Liu, Yingdong Dou, Philip S Yu, Yutong Deng, and Hao Peng. 2020. Alleviating the inconsistency problem of applying graph neural network to fraud detection. In *ACM SIGIR*. 1569–1572.
- [23] Jianxin Ma, Peng Cui, Kun Kuang, Xin Wang, and Wenwu Zhu. 2019. Disentangled graph convolutional networks. In *ICML*. PMLR, 4212–4221.
- [24] Shirui Pan, Ruiqi Hu, Sai-fu Fung, Guodong Long, Jing Jiang, and Chengqi Zhang. 2019. Learning graph embedding with adversarial training methods. *IEEE transactions on cybernetics* 50, 6 (2019), 2475–2487.
- [25] Hezhe Qiao and Guansong Pang. 2023. Truncated affinity maximization: One-class homophily modeling for graph anomaly detection. *NeurIPS* (2023).
- [26] Hezhe Qiao, Hanghang Tong, Bo An, Irwin King, Charu Aggarwal, and Guansong Pang. 2025. Deep graph anomaly detection: A survey and new perspectives. *IEEE Transactions on Knowledge and Data Engineering* (2025).
- [27] Hezhe Qiao, Qingsong Wen, Xiaoli Li, Ee-Peng Lim, and Guansong Pang. 2024. Generative semi-supervised graph anomaly detection. *NeurIPS* (2024).
- [28] Zidi Qin, Yang Liu, Qing He, and Xiang Ao. 2022. Explainable graph-based fraud detection via neural meta-graph search. In *CIKM*. 4414–4418.
- [29] Huijun Wu, Chen Wang, Yuriy Tyshetskiy, Andrew Docherty, Kai Lu, and Liming Zhu. 2019. Adversarial examples on graph data: Deep insights into attack and defense. *arXiv preprint arXiv:1903.01610* (2019).
- [30] Qitian Wu, Hengrui Zhang, Junchi Yan, and David Wipf. 2022. Handling distribution shifts on graphs: An invariance perspective. *ICLR* (2022).
- [31] Riting Xia, Huibo Liu, Anchen Li, Xueyan Liu, Yan Zhang, Chunxu Zhang, and Bo Yang. 2025. Incomplete graph learning: A comprehensive survey. *Neural Networks* (2025), 107682.
- [32] Yang Xiao, Jie Li, and Wengui Su. 2021. A lightweight metric defence strategy for graph neural networks against poisoning attacks. In *International Conference on Information and Communications Security*. Springer, 55–72.
- [33] Xuanting Xie, Wenyu Chen, and Zhao Kang. 2025. Robust graph structure learning under heterophily. *Neural Networks* 185 (2025), 107206.
- [34] Jiarong Xu, Junru Chen, Siqi You, Zhiqing Xiao, Yang Yang, and Jiangang Lu. 2021. Robustness of deep learning models on graphs: A survey. *AI Open* 2 (2021).
- [35] Kaidi Xu, Hongge Chen, Sijia Liu, Pin-Yu Chen, Tsui-Wei Weng, Mingyi Hong, and Xue Lin. 2019. Topology attack and defense for graph neural networks: An optimization perspective. *arXiv preprint arXiv:1906.04214* (2019).
- [36] Liang Yang, Mengzhe Li, Liyang Liu, Chuan Wang, Xiaochun Cao, et al. 2021. Diverse message passing for attribute with heterophily. *NeurIPS* (2021).
- [37] Gilad Yehudai, Ethan Fetaya, Eli Meir, Gal Chechik, and Hagai Maron. 2021. From local structures to size generalization in graph neural networks. In *ICML*.
- [38] Xiang Zhang and Marinka Zitnik. 2020. GnnGuard: Defending graph neural networks against adversarial attacks. *NeurIPS* 33 (2020), 9263–9275.
- [39] Tong Zhao, Yozen Liu, Leonardo Neves, Oliver Woodford, Meng Jiang, and Neil Shah. 2021. Data augmentation for graph neural networks. In *AAAI*, Vol. 35.
- [40] Dongcheng Zou, Hao Peng, and Chunyang Liu. 2024. A structural information guided hierarchical reconstruction for graph anomaly detection. In *CIKM*.
- [41] Xu Zou, Qinkai Zheng, Yuxiao Dong, Xinyu Guan, Evgeny Kharlamov, Jialiang Lu, and Jie Tang. 2021. Tdgia: Effective injection attacks on graph neural networks. In *ACM SIGKDD*. 2461–2471.
- [42] Daniel Zügner, Amir Akbarnejad, and Stephan Günnemann. 2018. Adversarial attacks on neural networks for graph data. In *ACM SIGKDD*. 2847–2856.
- [43] Daniel Zügner, Oliver Borchert, Amir Akbarnejad, and Stephan Günnemann. 2020. Adversarial attacks on graph neural networks: Perturbations and their patterns. *ACM TKDD* 14, 5 (2020), 1–31.
- [44] D. Zügner and S. Günnemann. 2019. Adversarial attacks on graph neural networks via meta learning. In *ACM SIGKDD*. 246–256.