

BRYAN HOOI

Homepage: <http://bhooi.github.io>

+65 96503537 • bhooi@comp.nus.edu.sg

EDUCATION

Carnegie Mellon University (2014 - 2019) • Ph.D. in Machine Learning •

Stanford University (2010 - 2014) • B.S. (Honors) in Mathematics with Distinction, M.S. in Computer Science •

RESEARCH INTERESTS

- Machine Learning, Natural Language Processing, Graphs, Trustworthy AI

HONORS AND AWARDS

1. NUS School of Computing Faculty Teaching Excellence Award 2024
2. ECML-PKDD 2018 Runner-Up Best Student Data Mining Paper Award
3. KDD 2016 Best Paper Award (Research Track)
4. Undergraduate Research Award 2014, Stanford University

SERVICE

- Co-chair (Short Paper Track): CIKM '25
- Area Chair: ICLR '23, NeurIPS '23
- Session Chair: WSDM '22
- Senior Program Committee Member: IJCAI, AAAI, WSDM, KDD
- Workshop Organizer: KDD'21 Outlier Description and Detection, WebConf '24 Graph Foundation Models, ICLR'25 Reasoning and Planning in LLMs, ICLR'25 Machine Learning on Graphs, ICLR'25 AI for Nucleic Acids, AAAI'26 AI for Scientific Research

SELECTED CONFERENCE PUBLICATIONS

- Hui Chen, Miao Xiong, Yujie Lu, Wei Han, Ailin Deng, Yufei He, Jiaying Wu, Yibo Li, Yue Liu, Bryan Hooi. "MLR-Bench: Evaluating AI Agents on Open-Ended Machine Learning Research" NeurIPS 2025 (Datasets and Benchmarks Track)
- Yibo Li, Miao Xiong, Jiaying Wu, Bryan Hooi. "ConfTuner: Training Large Language Models to Express Their Confidence Verbally" NeurIPS 2025
- Haonan Yuan, Qingyun Sun, Junhua Shi, Xingcheng Fu, Bryan Hooi, Jianxin Li, Philip S. Yu. "GRAVER: Generative Graph Vocabularies for Robust Graph Foundation Models Fine-tuning" NeurIPS 2025
- Yue Liu, Shengfang Zhai, Mingzhe Du, Yulin Chen, Tri Cao, Hongcheng Gao, Cheng Wang, Xinfeng Li, Kun Wang, Junfeng Fang, Jiaheng Zhang, Bryan Hooi. "GuardReasoner-VL: Safeguarding VLMs via Reinforced Reasoning" NeurIPS 2025
- Wen Tao, Jing Tang, Alvin Chan, Bryan Hooi, Baolong Bi, Nanyun Peng, Yuansheng Liu, Yiwei Wang. "How to Make Large Language Models Generate 100% Valid Molecules?" EMNLP 2025
- Yulin Chen, Haoran Li, Yuexin Li, Yue Liu, Yangqiu Song, Bryan Hooi. "TopicAttack: An Indirect Prompt Injection Attack via Topic Transition" EMNLP 2025
- Lycheng Wu, Mengru Wang, Ziwen Xu, Tri Cao, Nay Oo, Bryan Hooi, Shumin Deng. "Automating Steering for Safe Multimodal Large Language Models" EMNLP 2025
- Shuyang Hao, Yiwei Wang, Bryan Hooi, Jun Liu, Muhao Chen, Zi Huang, Yujun Cai. "Making Every Step Effective: Jailbreaking Large Vision-Language Models Through Hierarchical KV Equalization" EMNLP 2025 (Findings)
- Zijie Lin, Bryan Hooi. "Enhancing Multi-Agent Debate System Performance via Confidence Expression" EMNLP 2025 (Findings)
- Yulin Chen, Haoran Li, Yuan Sui, Yangqiu Song, Bryan Hooi. "Backdoor-Powered Prompt Injection Attacks Nullify Defense Methods" EMNLP 2025 (Findings)
- Cheng Wang, Yue Liu, Baolong Bi, Duzhen Zhang, Zhong-Zhi Li, Yingwei Ma, Yufei He, Shengju Yu, Xinfeng Li, Junfeng Fang, Jiaheng Zhang, Bryan Hooi. "Safety in Large Reasoning Models: A Survey" EMNLP 2025 (Findings)
- Yuan Sui, Yufei He, Zifeng Ding, Bryan Hooi. "Can Knowledge Graphs Make Large Language Models More Trustworthy? An Empirical Study Over Open-ended Question Answering" ACL 2025
- Haoming Xu, Ningyuan Zhao, Liming Yang, Sendong Zhao, Shumin Deng, Mengru Wang, Bryan Hooi, Nay Oo, Huajun Chen, Ningyu Zhang. "ReLearn: Unlearning via Learning for Large Language Models" ACL 2025

- Zhiyuan Hu, Yuliang Liu, Jinman Zhao, Suyuchen Wang, Yan Wang, Wei Shen, Qing Gu, Anh Tuan Luu, See-Kiong Ng, Zhiwei Jiang, Bryan Hooi. “LongRecipe: Recipe for Efficient Long Context Generalization in Large Language Models” ACL 2025
- Yulin Chen, Haoran Li, Yuan Sui, Yufei He, Yue Liu, Yangqiu Song, Bryan Hooi. “Can Indirect Prompt Injection Attacks Be Detected and Removed?” ACL 2025
- Yulin Chen, Haoran Li, Zihao Zheng, Dekai Wu, Yangqiu Song, Bryan Hooi. “Defense Against Prompt Injection Attack by Leveraging Attack Techniques” ACL 2025
- Zhecheng Li, Yiwei Wang, Bryan Hooi, Yujun Cai, Zhen Xiong, Nanyun Peng, Kai-Wei Chang. “Vulnerability of LLMs to Vertically Aligned Text Manipulations” ACL 2025
- Zhecheng Li, Yiwei Wang, Bryan Hooi, Yujun Cai, Nanyun Peng, Kai-Wei Chang. “DRS: Deep Question Reformulation With Structured Output” ACL (Findings) 2025
- James Xu Zhao, Jimmy Z.J. Liu, Bryan Hooi, See-Kiong Ng. “How Does Generation Length Affect Long-Form Factuality” ACL (Findings) 2025
- Yuan Sui, Yufei He, Nian Liu, Xiaoxin He, Kun Wang, Bryan Hooi. “FiDeLiS: Faithful Reasoning in Large Language Models for Knowledge Graph Question Answering” ACL (Findings) 2025
- Yue Liu, Xiaoxin He, Miao Xiong, Jinlan Fu, Shumin Deng, Yingwei Ma, Jiaheng Zhang, Bryan Hooi. “FlipAttack: Jailbreak LLMs via Flipping” ICML 2025
- Zhi Zheng, Zhuoliang Xie, Zhenkun Wang, Bryan Hooi. “Monte Carlo Tree Search for Comprehensive Exploration in LLM-Based Automatic Heuristic Design” ICML 2025
- Yuan Li, Jun Hu, Zemin Liu, Bryan Hooi, Jia Chen, Bingsheng He. “Adapting Precomputed Features for Efficient Graph Condensation” ICML 2025
- Haonan Yuan, Qingyun Sun, Junhua Shi, Xingcheng Fu, Bryan Hooi, Jianxin Li, Philip S. Yu. “How Much Can Transfer? BRIDGE: Bounded Multi-Domain Graph Pre-training and Prompt Learning with Generalization Error” ICML 2025
- Ailin Deng, Tri Cao, Zhirui Chen, Bryan Hooi. “Words or Vision: Do Vision-Language Models Have Blind Faith in Text?” CVPR 2025
- Shuyang Hao, Bryan Hooi, Jun Liu, Kai-Wei Chang, Zi Huang, Yujun Cai. “Exploring Visual Vulnerabilities via Multi-Loss Adversarial Search for Jailbreaking Vision-Language Models” CVPR 2025
- Cheng Wang, Yiwei Wang, Yujun Cai, Bryan Hooi. “Tricking Retrievers with Influential Tokens: An Efficient Black-Box Corpus Poisoning Attack” NAACL 2025
- Yifei Sun, Zemin Liu, Bryan Hooi, Yang Yang, Rizal Fathony, Jia Chen, Bingsheng He. “Multi-Label Node Classification with Label Influence Propagation” ICLR 2025
- Jinlan Fu, Shenzhen Huangfu, Hao Fei, Xiaoyu Shen, Bryan Hooi, Xipeng Qiu, See-Kiong Ng. “CHiP: Cross-modal Hierarchical Direct Preference Optimization for Multimodal LLMs” ICLR 2025
- Yuexin Li, Hiok Kuek Tan, Qiaoran Meng, Mei Lin Lock, Tri Cao, Shumin Deng, Nay Oo, Hoon Wei Lim, Bryan Hooi. “PhishIntel: Toward Practical Deployment of Reference-based Phishing Detection” WebConf 2025 (Demo Track)
- Yufei He, Yuan Sui, Xiaoxin He, Yue Liu, Yifei Sun, Bryan Hooi. “UniGraph2: Learning a Unified Embedding Space to Bind Multimodal Graphs” WebConf 2025
- Tri Cao*, Chengyu Huang*, Yuexin Li*, Huilin Wang, Amy He, Nay Oo, Bryan Hooi. “PhishAgent: A Robust Multimodal Agent for Phishing Webpage Detection” AAAI 2025 (AI for Social Impact Track)
- Jun Hu, Bryan Hooi, Bingsheng He, Yinwei Wei. “Modality-Independent Graph Neural Networks with Global Transformers for Multimodal Recommendation” AAAI 2025
- Cheng Wang, Yiwei Wang, Bryan Hooi, Yujun Cai, Nanyun Peng, Kai-Wei Chang. “Con-ReCall: Detecting Pre-training Data in LLMs via Contrastive Decoding” COLING 2025
- Yufei He, Bryan Hooi. “UniGraph: Learning a Cross-Domain Graph Foundation Model From Natural Language” KDD 2025
- Nan Chen, Zemin Liu, Bryan Hooi, Bingsheng He, Jun Hu, Jia Chen. “NodeImport: Imbalanced Node Classification with Node Importance Assessment” KDD 2025
- Jianqing Xu, Shen Li, Jiaying Wu, Miao Xiong, Ailin Deng, Jiazhen Ji, Yuge Huang, Guodong Mu, Wenjie Feng, Shouhong Ding, Bryan Hooi. “ID3: Identity-Preserving-yet-Diversified Diffusion Models for Synthetic Face Recognition” NeurIPS 2024
- Xiaoxin He, Yijun Tian, Yifei Sun, Nitesh V. Chawla, Thomas Laurent, Yann LeCun, Xavier Bresson, Bryan Hooi. “G-Retriever: Retrieval-Augmented Generation for Textual Graph Understanding and Question Answering” NeurIPS 2024
- Zhiyuan Hu, Chumin Liu, Xidong Feng, Yilun Zhao, See-Kiong Ng, Anh Tuan Luu, Junxian He, Pang Wei Koh, Bryan Hooi. “Uncertainty of Thoughts: Uncertainty-Aware Planning Enhances Information Seeking in Large Language Models” NeurIPS 2024
- Rishabh Anand, Chaitanya K. Joshi, Alex Morehead, Arian Jamasb, Charles Harris, Simon V. Mathis, Kieran Didi, Bryan Hooi, Pietro Liò. “RNA-FrameFlow: Flow Matching for de novo 3D RNA Backbone Design” MLCB 2024

- Jun Hu, Bryan Hooi, Bingsheng He. “Efficient Heterogeneous Graph Learning via Random Projection” TKDE 2024
- Jihai Zhang, Xiang Lan, Xiaoye Qu, Yu Cheng, Mengling Feng, Bryan Hooi. “Learning the Unlearned: Mitigating Feature Suppression in Contrastive Learning” ECCV 2024
- Adam Goodge, Bryan Hooi, Wee Siong Ng. “When Text and Images Don’t Mix: Bias-Correcting Language-Image Similarity Scores for Anomaly Detection” BMVC 2024
- Yuexin Li, Chengyu Huang, Shumin Deng, Mei Lin Lock, Tri Cao, Nay Oo, Bryan Hooi, Hoon Wei Lim. “KnowPhish: Large Language Models Meet Multimodal Knowledge Graphs for Enhancing Reference-Based Phishing Detection” USENIX Security 2024
- Jiaying Wu, Jiafeng Guo, Bryan Hooi. “Fake News in Sheeps Clothing: Robust Fake News Detection Against LLM-Empowered Style Attacks” KDD 2024
- Jintian Zhang, Xin Xu, Ningyu Zhang, Ruibo Liu, Bryan Hooi, Shumin Deng. “Exploring Collaboration Mechanisms for LLM Agents: A Social Psychology View” ACL 2024
- Adrien Benamira, Thomas Peyrin, Trevor Yap, Tristan Guerand, Bryan Hooi. “Truth Table Net: Scalable, Compact and Verifiable Neural Networks with a Dual Convolutional Small Boolean Circuit Networks Form” IJCAI 2024
- Wenjie Feng, Li Wang, Bryan Hooi, See Kiong Ng, Shenghua Liu. “Interrelated Dense Pattern Detection in Multilayer Networks” TKDE 2024
- Ailin Deng, Zhirui Chen, Bryan Hooi. “Seeing is Believing: Mitigating Hallucination in Large Vision-Language Models via CLIP-Guided Decoding” Workshop on Reliable and Responsible Foundation Models, ICLR 2024
- Xu Liu, Junfeng Hu, Yuan Li, Shizhe Diao, Yuxuan Liang, Bryan Hooi, Roger Zimmermann. “UniTime: A Language-Empowered Unified Model for Cross-Domain Time Series Forecasting” WebConf 2024
- Miao Xiong, Zhiyuan Hu, Xinyang Lu, Yifei Li, Jie Fu, Junxian He, Bryan Hooi. “Can LLMs Express Their Uncertainty? An Empirical Evaluation of Confidence Elicitation in LLMs” ICLR 2024
- Juncheng Liu, Bryan Hooi, Kenji Kawaguchi, Yiwei Wang, Chaosheng Dong, Xiaokui Xiao. “Scalable and Effective Implicit Graph Neural Networks on Large Graphs” ICLR 2024
- Xiaoxin He, Xavier Bresson, Thomas Laurent, Adam Perold, Yann LeCun, Bryan Hooi. “Harnessing Explanations: LLM-to-LM Interpreter for Enhanced Text-Attributed Graph Representation Learning” ICLR 2024
- Wei Zhuo, Zemin Liu, Bryan Hooi, Bingsheng He, Guang Tan, Rizal Fathony, Jia Chen. “Partitioning Message Passing for Graph Fraud Detection” ICLR 2024
- Nan Chen, Zemin Liu, Bryan Hooi, Bingsheng He, Rizal Fathony, Jun Hu, Jia Chen. “Consistency Training with Learnable Data Augmentation for Graph Anomaly Detection with Limited Supervision” ICLR 2024
- Zhiyuan Hu, Chumin Liu, Yue Feng, Anh Tuan Luu, Bryan Hooi. “PoetryDiffusion: Towards Joint Semantic and Metrical Manipulation in Poetry Generation” AAAI 2024
- Jun Hu, Bryan Hooi, Shengsheng Qian, Quan Fang, Changsheng Xu. “MGDCF: Distance Learning via Markov Graph Diffusion for Neural Collaborative Filtering” TKDE 2024
- Yifan Zhang, Bryan Hooi. “HiPA: Enabling One-Step Text-to-Image Diffusion Models via High-Frequency-Promoting Adaptation” arXiv 2023
- Shumin Deng, Ningyu Zhang, Nay Oo, Bryan Hooi. “Towards A Unified View of Answer Calibration for Multi-Step Reasoning” arXiv 2023
- Miao Xiong, Ailin Deng, Pang Wei Koh, Jiaying Wu, Shen Li, Jianqing Xu, Bryan Hooi. “Proximity-Informed Calibration for Deep Neural Networks” NeurIPS 2023
- Yifan Zhang, Daquan Zhou, Bryan Hooi, Kai Wang, Jiashi Feng. “Expanding Small-Scale Datasets with Guided Imagination” NeurIPS 2023
- Yiwei Wang, Yujun Cai, Muhao Chen, Yuxuan Liang, Bryan Hooi. “Primacy Effect of ChatGPT” EMNLP 2023
- Yiwei Wang, Bryan Hooi, Fei Wang, Yujun Cai, Yuxuan Liang, Wenxuan Zhou, Jing Tang, Manjuan Duan, Muhao Chen. “How Fragile is Relation Extraction under Entity Replacements?” CoNLL 2023
- Xu Liu, Yutong Xia, Yuxuan Liang, Junfeng Hu, Yiwei Wang, Lei Bai, Chao Huang, Zhengguang Liu, Bryan Hooi, Roger Zimmermann. “LargeST: A Benchmark Dataset for Large-Scale Traffic Forecasting” NeurIPS 2023 (Datasets and Benchmarks Track)
- Baixiang Huang, Bryan Hooi and Kai Shu. “TAP: A Comprehensive Data Repository for Traffic Accident Prediction in Road Networks” SIGSPATIAL 2023
- Jiaying Wu, Shen Li, Ailin Deng, Miao Xiong, Bryan Hooi. “Prompt-and-Align: Prompt-Based Social Alignment for Few-Shot Fake News Detection” CIKM 2023
- Zhiyuan Hu, Yue Feng, Anh Tuan Luu, Bryan Hooi, Aldo Lipani. “Unlocking the Potential of User Feedback: Leveraging Large Language Model as User Simulators to Enhance Dialogue System” CIKM 2023 (Short Paper)
- Jiaying Wu, Bryan Hooi. “DECOR: Degree-Corrected Social Graph Refinement for Fake News Detection” KDD 2023

- Siddharth Bhatia, Mohit Wadhwa, Kenji Kawaguchi, Neil Shah, Philip S. Yu, Bryan Hooi. “Sketch-Based Anomaly Detection in Streaming Graphs” KDD 2023
- Shumin Deng, Shengyu Mao, Ningyu Zhang, Bryan Hooi. “SPEECH: Structured Prediction with Energy-Based Event-Centric Hyperspheres” ACL 2023
- Ailin Deng, Miao Xiong, Bryan Hooi. “Great Models Think Alike: Improving Model Reliability via Inter-Model Latent Agreement” ICML 2023
- Xiaoxin He, Bryan Hooi, Thomas Laurent, Adam Perold, Yann LeCun, Xavier Bresson. “A Generalization of ViT/MLP-Mixer to Graphs” ICML 2023
- Kaixin Wang, Kuangqi Zhou, Jiashi Feng, Bryan Hooi, Xinchao Wang. “Reachability-Aware Laplacian Representation in Reinforcement Learning” ICML 2023
- Yuwen Li, Miao Xiong, Bryan Hooi. “GraphCleaner: Detecting Mislabelled Samples in Popular Graph Learning Benchmarks” ICML 2023
- Mingqi Yang, Wenjie Feng, Yanming Shen, Bryan Hooi. “Towards Better Graph Representation Learning with Parameterized Decomposition & Filtering” ICML 2023
- Yifan Zhang, Bingyi Kang, Bryan Hooi, Shuicheng Yan, Jiashi Feng. “Deep Long-Tailed Learning: A Survey” TPAMI 2023
- Jianqing Xu, Shen Li, Ailin Deng, Miao Xiong, Jiaying Wu, Jiaxiang Wu, Shouhong Ding, Bryan Hooi. “Probabilistic Knowledge Distillation for Face Ensembles” CVPR 2023
- Shumin Deng, Chengming Wang, Zhoubo Li, Ningyu Zhang, Zelin Dai, Hehong Chen, Feiyu Xiong, Ming Yan, Qiang Chen, Mosha Chen, Jiaoyan Chen, Jeff Z. Pan, Bryan Hooi, Huajun Chen. “Construction and Applications of Billion-Scale Pre-trained Multimodal Business Knowledge Graph” ICDE 2023
- Jiaxin Jiang, Yuan Li, Bingsheng He, Bryan Hooi, Jia Chen, Johan Kok Zhi Kang. “Spade: A real-time fraud detection framework on evolving graphs” VLDB 2023
- Yiwei Wang, Bryan Hooi, Yozen Liu, Neil Shah. “Graph Explicit Neural Networks: Explicitly Encoding Graphs for Efficient and Accurate Inference” WSDM 2023
- Yiwei Wang, Bryan Hooi, Yozen Liu, Tong Zhao, Zhichun Guo, Neil Shah. “Flashlight: Scalable Link Prediction with Effective Decoders” LoG 2022
- Kuangqi Zhou, Kaixin Wang, Jian Tang, Jiashi Feng, Bryan Hooi, Peilin Zhao, Tingyang Xu, Xinchao Wang. “Jointly Modelling Uncertainty and Diversity for Active Molecular Property Prediction” LoG 2022
- Yifan Zhang, Bryan Hooi, Lanqing Hong, Jiashi Feng. “Self-Supervised Aggregation of Diverse Experts for Test-Agnostic Long-Tailed Recognition” NeurIPS 2022
- Juncheng Liu, Bryan Hooi, Kenji Kawaguchi, Xiaokui Xiao. “MGNNI: Multiscale Graph Neural Networks with Implicit Layers” NeurIPS 2022
- Miao Xiong, Shen Li, Wenjie Feng, Ailin Deng, Jihai Zhang, Bryan Hooi. “Birds of a Feather Trust Together: Knowing When to Trust a Classifier via Adaptive Neighborhood Aggregation” TMLR 2022
- Aashish Kolluri, Teodora Baluta, Bryan Hooi, Prateek Saxena. “LPGNet: Link Private Graph Networks for Node Classification” ACM CCS 2022
- Xu Liu, Yuxuan Liang, Chao Huang, Yu Zheng, Bryan Hooi and Roger Zimmermann. “When Do Contrastive Learning Signals Help Spatio-Temporal Graph Forecasting?” SIGSPATIAL 2022
- Ailin Deng, Shen Li, Miao Xiong, Zhirui Chen, Bryan Hooi. “Trust, but Verify: Using Self-Supervised Probing to Improve Trustworthiness” ECCV 2022
- Yiwei Wang, Yujun Cai, Yuxuan Liang, Henghui Ding, Changhu Wang, Bryan Hooi. “Time-Aware Neighbor Sampling on Temporal Graphs” IJCNN 2022
- Jiaying Wu, Bryan Hooi. “Probing Spurious Correlations in Popular Event-Based Rumor Detection Benchmarks” ECML-PKDD 2022
- Juncheng Liu, Yiwei Wang, Bryan Hooi, Renchi Yang, Xiaokui Xiao. “LSCALE: Latent Space Clustering-Based Active Learning for Node Classification” ECML-PKDD 2022
- Adam Goodge, Bryan Hooi, See Kiong Ng, Wee Siong Ng. “ARES: Locally Adaptive Reconstruction-based Anomaly Scoring” ECML-PKDD 2022
- Kaixin Wang, Navdeep Kumar, Kuangqi Zhou, Bryan Hooi, Jiashi Feng, Shie Mannor. “The Geometry of Robust Value Functions” ICML 2022
- Shen Li, Bryan Hooi. “Neural PCA for Flow-Based Representation Learning” IJCAI 2022
- Ailin Deng, Adam Goodge, Lang Yi Ang, Bryan Hooi. “CADET: Calibrated Anomaly Detection for Mitigating Hardness Bias” IJCAI 2022
- Yiwei Wang, Muhao Chen, Wenxuan Zhou, Yujun Cai, Yuxuan Liang, Dayiheng Liu, Baosong Yang, Juncheng Liu, Bryan Hooi. “Should We Rely on Entity Mentions for Relation Extraction? Debiasing Relation Extraction with Counterfactual Analysis” NAACL 2022
- Nicholas Lim, Bryan Hooi, See-Kiong Ng, Yong Liang Goh, Renrong Weng, Rui Tan. “Hierarchical Multi-Task Graph Recurrent Network for Next POI Recommendation” SIGIR 2022

- Siddharth Bhatia, Arjit Jain, Shivin Srivastava, Kenji Kawaguchi, Bryan Hooi. “MemStream: Memory-Based Streaming Anomaly Detection” WebConf 2022
- Adam Goodge, Bryan Hooi, See Kiong Ng, Wee Siong Ng. “LUNAR: Unifying Local Outlier Detection Methods via Graph Neural Networks” AAAI 2022
- Shenghua Liu, Bin Zhou, Quan Ding, Bryan Hooi, Zhengbo Zhang, Huawei Shen, Xueqi Cheng. “Time Series Anomaly Detection with Adversarial Reconstruction Networks” IEEE TKDE 2022
- Shimiao Li, Amritanshu Pandey, Bryan Hooi, Christos Faloutsos, Larry Pileggi. “Dynamic Graph-Based Anomaly Detection in the Electrical Grid” IEEE Transactions on Power Systems 2022
- Xiaobing Sun, Wenjie Feng, Shenghua Liu, Yuyang Xie, Siddharth Bhatia, Bryan Hooi, Wenhan Wang, Xueqi Cheng. “MonLAD: Money Laundering Agents Detection in Transaction Streams” WSDM 2022
- Siddharth Bhatia, Rui Liu, Bryan Hooi, Minji Yoon, Kijung Shin, Christos Faloutsos. “Real-Time Anomaly Detection in Edge Streams” TKDD 2022