

Self-Organization Preserved Graph Structure Learning with Principle of Relevant Information

Qingyun Sun^{1, 2}, Jianxin Li^{1, 2}, Beining Yang^{1, 2}, Xingcheng Fu^{1, 2}, Hao Peng¹, Philip S. Yu³

¹ Beijing Advanced Innovation Center for Big Data and Brain Computing, Beihang University, Beijing 100191, China

² School of Computer Science and Engineering, Beihang University, Beijing 100191, China

³ Department of Computer Science, University of Illinois at Chicago, Chicago, USA
{sunqy, lijx, yangbeining, fuxc, penghao}@buaa.edu.cn, psyu@uic.edu

Abstract

Most Graph Neural Networks follow the message-passing paradigm, assuming the observed structure depicts the ground-truth node relationships. However, this fundamental assumption cannot always be satisfied, as real-world graphs are always incomplete, noisy, or redundant. How to reveal the inherent graph structure in a unified way remains under-explored. We proposed **PRI-GSL**, a Graph Structure Learning framework guided by the **Principle of Relevant Information**, providing a simple and unified framework for identifying the self-organization and revealing the hidden structure. PRI-GSL learns a structure that contains the most relevant yet least redundant information quantified by von Neumann entropy and Quantum Jensen-Shannon divergence. PRI-GSL incorporates the evolution of quantum continuous walk with graph wavelets to encode node structural roles, showing in which way the nodes interplay and self-organize with the graph structure. Extensive experiments demonstrate the superior effectiveness and robustness of PRI-GSL.

Introduction

Graph Neural Networks (GNNs) (Wu et al. 2020) have gained popularity in recent years due to their remarkable success in representing graph data in diverse tasks and applications. Most of the existing GNNs follow the message-passing paradigm (Gilmer et al. 2017), i.e., exchanging information between neighbors along the graph structure. They take the raw graph structure as the path of information flow, assuming the observed structure perfectly depicts the ground-truth relations between nodes. However, these raw graphs are naturally collected from network-structure data (e.g., social networks), which are often noisy, incomplete, and independent of the downstream tasks. There is a gap between the raw structure and the optimal structure for specific tasks. The poor quality of graph structure leads to the poor quality of representations produced by GNNs, making GNNs prone to noise and adversarial attacks (Zügner, Akbarnejad, and Günnemann 2018; Sun et al. 2018, 2021).

Graph structure learning (Zhu et al. 2022) aims to learn a new structure of high quality simultaneously with the graph representations, which has received growing attention for its utility for improving representation quality and robustness.

Copyright © 2023, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

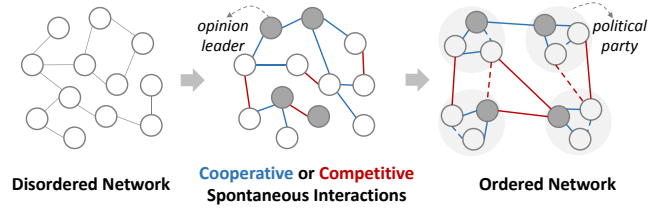


Figure 1: Self-organization in political network.

Most existing methods optimize the structure with heuristic assumptions (e.g., community (Wang et al. 2021)) or certain structure constraints (e.g., sparsity, low-rank, and smoothness (Jin et al. 2020; Sun et al. 2022b)). However, these assumptions and constraints cannot always be applicable to all graphs and tasks. *How to reveal the inherent graph structure in a unified way remains an under-explored question.*

Most of the graph data in the real-world shows “self-organization” property from molecules (Eigen and Schuster 1977) to social networks (Bonabeau et al. 1997), where the nodes organize their interactions spontaneously through the structure to create a global order amongst themselves. As the example in Fig. 1, the emergence of opinion leaders and the cooperation/competitive behaviors between people form political parties, making the political network more ordered. Since the structure navigates the information flow between nodes and decides the graph’s fundamental mechanism, it can be optimized to identify the organization and reduce the disorder of the noisy graph.

In this paper, we introduce the self-organized Principle of Relevant Information (PRI) to quantify the structure from an information-theoretic point of view. We propose a novel graph structure learning framework named PRI-GSL, which inherits the merits of PRI to identify the self-organization and reveal the inherent structure of graph. Rather than imposing statistical constraints on the graph data, PRI-GSL takes the structure learning as a trade-off between structure redundancy reduction and information preservation, and then use the von Neumann entropy and the Quantum Jensen-Shannon divergence to quantify them. To better capture the contribution of nodes in the self-organization evolution process, we use the quantum continuous walk evolution with multi-scale graph wavelets to characterize node

structural roles and incorporate them into structure learning. In this way, PRI-GSL enumerates the potential edges and preserves the most relevant yet least redundant ones, showing in which way the nodes interplay and self-organize with the graph structure.

- We propose PRI-GSL, an information-theoretic graph structure learning framework with the Principle of Relevant Information, providing a simple yet unified way to quantify the learned structure and unravel the graph self-organization.
- We use the quantum continuous walk with graph wavelets to encode node structural roles in a continuous and time-varying way, which is incorporated in structure learning to fully characterize the nodes in self-organization.
- Extensive experiment results demonstrate the superior effectiveness and robustness of PRI-GSL.

Related Work

Graph structure learning has gained more attention in recent years (Zhu et al. 2022) to improve the quality of graph representations by learning a better graph structure. Most existing works (Jin et al. 2020; Wang et al. 2021) optimize the structure with assumptions or certain constraints in a heuristic way. Substantial efforts have been made to give a theoretical quantification for the learned structure. SDRF (Topping et al. 2022) refines the structure based on the Ricci curvature of edges in a greedy pre-process strategy. SIB (Yu et al. 2020a) utilizes the information bottleneck principle to find the most predictive subgraph. VIB-GSL (Sun et al. 2022a) proposes a variational information bottleneck principle to learn a new structure for graph classification, which is not applicable to the node-level tasks. Graph-PRI (Yu et al. 2022) advances the Principle of Relevant Information for graph sparsification in an unsupervised way without considering node features and the specific downstream task.

Information theory provides a powerful methodology to describe general properties of arbitrarily complex systems. In information theory, there are two representative self-organizing principles: Information Bottleneck (IB) (Tishby, Pereira, and Bialek 2000) and Principle of Relevant Information (PRI) (Principe 2010). Both IB and PRI describe different forms of redundancy reduction and information preservation. The famed IB is formulated on the mutual information between independent and identically distributed (i.i.d.) data, which is difficult to model the complex node interactions imposed by the graph structure. PRI shares the spirit of the IB method but its formulation addresses the entropy and relative entropy of a single dataset (Principe 2010), which can be applied to graph data with well-defined information-theoretic tools.

Preliminary

Notions Given a graph $G = \{V, E\}$ where V is the set of N nodes and E is the edge set. $\mathbf{A} \in \mathbb{R}^{N \times N}$ is the adjacency matrix and $\mathbf{D}$ is the degree matrix. The Laplacian matrix of the graph G can be defined as $\mathbf{L} = \mathbf{D} - \mathbf{A} = \mathbf{U}\mathbf{\Lambda}\mathbf{U}^T$, where $\mathbf{U}$ is the eigenvector matrix, $\mathbf{\Lambda} = \text{Diag}(\lambda_1, \dots, \lambda_N)$ and $\lambda_1 < \lambda_2 \leq \dots \leq \lambda_N$ are the eigenvalues of $\mathbf{L}$.

Graph Structure Learning Given a graph G , graph structure learning (Zhu et al. 2022) aims to learn a new structure $\tilde{G}$ simultaneously with the graph representations with the objective function:

$$\mathcal{L} = \mathcal{L}_{task}(\tilde{G}, Y) + \alpha \mathcal{L}_{reg}(\tilde{G}, G), \quad (1)$$

where $\mathcal{L}_{task}$ is the task-specific objective with respect to the learned graph $\tilde{G}$ and the ground truth Y , $\mathcal{L}_{reg}$ imposes constraints on the learned graph and α is a hyper-parameter.

Principle of Relevant Information PRI (Principe 2010) is a self-organized information-theoretic principle that aims to perform mode decomposition of a random variable to obtain a reduced statistical representation. PRI formulates the redundancy reduction and information preservation as a trade-off between the entropy of reduced representation and its relative entropy given the original data.

Definition 1 (Principle of Relevant Information). *Given a random variable X , the Principle of Relevant Information aims to obtain a reduced representation T with:*

$$\mathcal{L}_{PRI} = \arg \min_T H(T) + \beta D(\mathbb{P}(T) || \mathbb{P}(X)), \quad (2)$$

where $H(T)$ is the entropy of T and $D(\mathbb{P}(T) || \mathbb{P}(X))$ is the divergence of distributions $\mathbb{P}(T)$ and $\mathbb{P}(X)$.

The first term $H(T)$ measures the redundancy of representation T and the second term $D(\mathbb{P}(T) || \mathbb{P}(X))$ measures the allowable distortion of the original data. The hyper-parameter β controls the level of distortion in T . PRI was commonly used in scalar random variables (Wei et al. 2021; Hoyos-Osorio et al. 2021) and defined by Rényi's formulation of entropy and divergence (Rényi et al. 1961).

Graph Structure Learning with Principle of Relevant Information

In this work, we propose a graph structure learning framework named PRI-GSL, which merits the Principle of Relevant Information as a guideline for controlling the structure quality. The overall architecture of PRI-GSL is shown in Figure 2. In this section, we first formulate the PRI loss for structure learning, then introduce the Role-aware graph learner and learning process of PRI-GSL.

PRI for Graph Structure Learning

In the PRI-GSL framework, PRI performs as a self-supervised regularizer for the quality of the learned graph structure. Motivated by the objective of PRI, the Graph Structure Learning Principle of Relevant Information is:

Definition 2 (PRI for Graph Structure Learning). *Given a graph G , the Principle of Relevant Information for graph structure learning aims to learn a refined graph $\tilde{G}$ with:*

$$\mathcal{L}_{PRI} = H(\tilde{G}) + \beta D(\tilde{G} || G), \quad (3)$$

The first term $H(\tilde{G})$ is the redundancy term, which measures the disorder of the learned graph $\tilde{G}$. The larger the $H(\tilde{G})$, the more disordered $\tilde{G}$ is. The second term $D(\tilde{G} || G)$

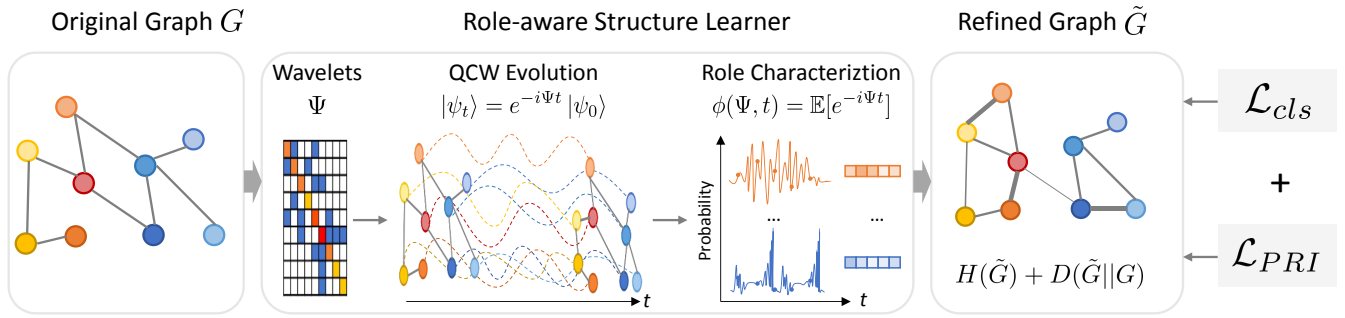


Figure 2: Overall Architecture of PRI-GSL.

is the distortion term, which measures the discrepancy between two graphs. The smaller the $D(\tilde{G}||G)$, the more similar are the distributions of $\tilde{G}$ and G . β denotes the trade-off between the redundancy reduction of $\tilde{G}$ and its discriminative description power of G . As β becomes larger, the emphasis is laid more on the distortion term, and more information from G is preserved in $\tilde{G}$.

Formulate PRI by von Neumann Entropy and Quantum Jensen-Shannon Divergence

The choice of entropy and divergence in PRI is application-specific. In this paper, we formulate the PRI loss $\mathcal{L}_{PRI}$ by the von Neumann entropy (VNE) (Nielsen and Chuang 2002) and quantum Jensen-Shannon (QJS) divergence (Lamberti et al. 2008) for graph data with complex interactions as in (Yu et al. 2022).

For the first redundancy term $H(\tilde{G})$, we propose to measure the structure redundancy by von Neumann entropy (VNE) (Nielsen and Chuang 2002), which has been used in a variety of graph learning studies (Dasoulas et al. 2020; Yu et al. 2022). von Neumann entropy quantifies the spectral complexity (or disorder) of graph structure by taking the graph as a quantum system through a mapping between discrete Laplacians and quantum states. A density matrix ρ is a Hermitian and positive semi-definite matrix that is used to encode the probability distributions and describe the state of a quantum mechanical system. For the graph $\tilde{G}$, the von Neumann entropy $H_{vN}(\tilde{G})$ is defined as:

$$H_{vN}(\tilde{G}) = H_{vN}(\tilde{\rho}) = -\text{tr}(\tilde{\rho} \log \tilde{\rho}) = -\sum_{i=1}^N (\lambda_i \log \lambda_i), \quad (4)$$

where $\tilde{\rho}$ is the graph density matrix of $\tilde{G}$, $\text{tr}(\cdot)$ denotes trace and $\{\lambda_i\}$ are the eigenvalues of $\tilde{\rho}$. Typically, both the Laplacian matrix and the normalized Laplacian matrix can be used for the mapping from graphs to states (Minello, Rossi, and Torsello 2019). We define the density matrix $\tilde{\rho} = \frac{\tilde{\mathbf{L}}}{\text{tr}(\tilde{\mathbf{L}})} = \frac{\tilde{\mathbf{L}}}{2|E|}$ based on the Laplacian matrix $\tilde{\mathbf{L}}$ of $\tilde{G}$, which models the continuous information diffusion process (De Domenico and Biamonte 2016).

For the second distortion term $D(\tilde{G}||G)$, we use the quantum Jensen-Shannon (QJS) divergence (Lamberti et al.

2008) between the graph density matrices ($\tilde{\rho}$ and ρ) of G and $\tilde{G}$. The quantum Jensen-Shannon divergence has been widely used as a generalization of the classical Jensen-Shannon divergence to quantum states of graph data (De Domenico et al. 2015; Bai et al. 2015), which is symmetric, negative definite and bounded ($0 \leq D_{QJS} \leq 1$).

$$D_{QJS}(\tilde{G}||G) = H_{vN}\left(\frac{\tilde{\rho} + \rho}{2}\right) - \frac{1}{2}H_{vN}(\tilde{\rho}) - \frac{1}{2}H_{vN}(\rho). \quad (5)$$

Combining the redundancy term in Eq. (4) and the distortion term in Eq. (5), we can obtain the following objective:

$$\begin{aligned} \mathcal{L}_{PRI} &= H_{vN}(\tilde{G}) + \beta D_{QJS}(\tilde{G}||G) \\ &= \beta H_{vN}\left(\frac{\tilde{\rho} + \rho}{2}\right) + \frac{2-\beta}{2}H_{vN}(\tilde{\rho}) - \frac{\beta}{2}H_{vN}(\rho) \\ &\equiv \beta H_{vN}\left(\frac{\tilde{\rho} + \rho}{2}\right) + \frac{2-\beta}{2}H_{vN}(\tilde{\rho}). \end{aligned} \quad (6)$$

We neglect $H_{vN}(\rho)$ in the last line because it's a constant value during optimization. The above formalism provides a unified way of quantification and comparison for the learned structure. Then we use a graph structure learner to obtain $\tilde{G}$.

Role-aware Graph Structure Learner

In this section, we introduce the Role-aware Graph Structure Learner in PRI-GSL, which aims to learn a better graph $\tilde{G}$ that preserves the graph self-organization. Considering the evolution process in self-organization, we characterize the nodes' roles in a continuous and time-varying way and then incorporate both the merits of features as well as structural roles to refine the graph.

Structural Role Encoding The structural role of the node represents its contribution to the overall information flow of the graph, which can provide key insights into the identification of graph organization. We propose to model graph state by quantum continuous walk and use the time-evolution operator with graph wavelets to generate role encodings.

(1) Model graph state by QCW. Recall that we apply the von Neumann Entropy and the QJS divergence for PRI formulation, which takes the whole graph as a quantum system. To investigate the nodes' roles in this quantum system, we use the quantum continuous walk (QCW) (Childs 2010; Bai

et al. 2015) to build maps of how information flows through the graph in the perspective of graph state evolution. QCW is the quantum mechanical counterpart of the continuous-time random walk in a graph, which describes the propagation of a quantum particle evolving continuously in time on the nodes. The QCW on a graph G is defined as a dynamical process over the nodes that takes place on a N -dimensional Hilbert space $\mathbb{H} = \text{span}(\{|v\rangle, v \in V\})$. The evolution of the walker is governed by the Schrödinger equation

$$\frac{d}{dt} |\psi_t\rangle = -i\mathcal{H} |\psi_t\rangle. \quad (7)$$

$|\psi_t\rangle$ represents the state of the walk at time t , which is a time-dependent amplitude vector on nodes. $\mathcal{H}$ is the Hamiltonian operator, which accounts for the total energy of the graph and governs the time evolution of the quantum continuous walk. Given an initial state $|\psi_0\rangle \in \mathbb{H}$, the state of walker $|\psi_t\rangle$ evolves in time according to

$$|\psi_t\rangle = U(t) |\psi_0\rangle, \quad U(t) := e^{-i\mathcal{H}t}, \quad (8)$$

where $U(t)$ is the unitary time-evolution operator.

(2) Diffusion by graph wavelets. To give a full characterization of the structural properties, we use the spectral graph wavelets Ψ (Hammond, Vandergheynst, and Gribonval 2011) as the Hamiltonian $\mathcal{H}$ in QCW. Then the structural information residing in the graph is encoded in $U(t)$. $U(t)$ is a polynomial in Ψ for all t , thus any matrix that commutes with Ψ also commutes with $U(t)$ (Coutinho and Godsil 2021). We adopt the heat kernel $g_s(\Lambda) = e^{-\Lambda s}$ to obtain spectral graph wavelets, where the scaling parameter s controls the spread radii of the diffusion process and larger s allows farther diffusion. The spectral wavelet basis is

$$\Psi_s = \mathbf{U}\Lambda_s\mathbf{U}^T = (\Psi_s(1) |\Psi_s(2) | \cdots |\Psi_s(N)), \quad (9)$$

where $\Lambda_s = g_s(\Lambda)$. In this way, the spectral graph wavelet $\Psi_s(a)$ centered at node v_a associated with filter g_s will be given by an N -dimensional vector:

$$\Psi_s(a) = \mathbf{U}\Lambda_s\mathbf{U}^T\delta_a, \quad (10)$$

where δ_a is the one-hot vector of node v_a . $\Psi_s(a)$ is a N -dimensional vector where the b -th wavelet coefficient of $\Psi_s(a)$ represents the information that v_a received from v_b . Nodes playing similar roles have similar wavelet coefficient.

(3) Characterize by the time-evolution operator. Since the time-evolution operator $U(t)$ reflects the graph state evolution, we treat the wavelets Ψ as probability distributions over graph and use $\phi(\Psi, t) = \mathbb{E}[U(t)] = \mathbb{E}[e^{-i\Psi t}]$ as the characteristic function to uncover nodes' roles in information diffusion. The empirical characteristic function of v_a is:

$$\phi_s(v_a, t) = \frac{1}{N} \sum_{n=1}^N e^{-i\Psi_s(a)t}. \quad (11)$$

$\phi_s(v_a, t)$ can capture all the moments (including higher-order moments) of the given distribution $\Psi_s(a)$. We sample at T different time points on the time-evolution operator and then concatenate the values:

$$\mathbf{h}_s(v_a) = [\text{Re}(\phi_s(v_a, t)), \text{Im}(\phi_s(v_a, t))]_{t=t_1, t_2, \dots, T}. \quad (12)$$

In general, the nodes play different roles across different scales. Hence we utilize a multi-scale wavelet diffusion strategy to capture both the local and global structural roles of nodes. We integrate information across different radii of neighborhoods by jointly considering a set of different values of s . We can obtain multi-scale structure role encodings $\mathbf{h}_a \in \mathbb{R}^{2TM}$ by concatenating the structural encodings at several different scales $S : \{s_1, s_2, \dots, s_M\}$:

$$\mathbf{h}_a = (\mathbf{h}_{s_1}(v_a) | \mathbf{h}_{s_2}(v_a) | \cdots | \mathbf{h}_{s_M}(v_a)). \quad (13)$$

Iterative Structure Learning After obtaining the structural role encodings, we use a metric function that accounts for both feature information and the role-based similarities to measure the possibility of edge existence. PRI-GSL is agnostic to various metric functions and we choose the multi-head cosine similarity function here:

$$a_{ij}^{R(e-1)} = \frac{1}{m} \sum_{h=1}^m \cos(\mathbf{W}_h \cdot (\mathbf{z}_i^{(e-1)} | \mathbf{h}_i^{(e-1)}), \mathbf{W}_h \cdot (\mathbf{z}_j^{(e-1)} | \mathbf{h}_j^{(e-1)})), \quad (14)$$

where m is the number of heads, $\mathbf{W}_h$ is the weight matrix of the h -th head, $\mathbf{z}_i^{(e-1)}$ and $\mathbf{h}_i^{R(e-1)}$ denote the representation vector and the structural role encoding vector of node v_i in the $(e-1)$ -th epoch, and $|$ denotes the concatenation operation. With the above structure learning strategy, we can obtain a role-aware adjacency matrix in the e -th epoch:

$$\mathbf{A}_R^{(e)} = \{a_{ij}^{R(e-1)}\}, i, j \in \{1, 2, \dots, N\}. \quad (15)$$

Learning Process of PRI-GSL

Dynamic Structure Fusion The input graph structure determines the learning performance to a certain extent. To avoid the non-convergence or unstable training brought by the poor quality of learned structure at the beginning of training, we hence incorporate the original graph structure $\mathbf{A}$ as supplementary to formulate an optimized graph structure $\tilde{\mathbf{A}}$:

$$\tilde{\mathbf{A}}^{(e)} = \gamma \mathbf{D}^{-\frac{1}{2}} \mathbf{A} \mathbf{D}^{-\frac{1}{2}} + (1 - \gamma) \cdot \text{RowNorm}(\mathbf{A}_R^{(e)}), \quad (16)$$

where $\text{RowNorm}(\cdot)$ denotes the row-wise normalization function, γ is a constant that control the contribution of original structure. Here we use a dynamic decay mechanism for γ to enable the role-aware structure $\mathbf{A}_R^{(e)}$ to play a more and more important role during training. Then the refined structure is inputted into a GNN encoder for node representation vectors $\mathbf{Z} \in \mathbb{R}^{N \times d}$ and classification:

$$\mathbf{Z}^{(e)} = \text{GNN-Encoder}(\tilde{\mathbf{A}}^{(e)}, \mathbf{X}). \quad (17)$$

Objective of PRI-GSL The overall loss $\mathcal{L}$ of PRI-GSL is composed of two terms, the classification loss $\mathcal{L}_{cls}$ and the structure PRI loss $\mathcal{L}_{\text{PRI}}$ in Eq. (6), given by:

$$\begin{aligned} \mathcal{L} &= \mathcal{L}_{cls} + \alpha \mathcal{L}_{\text{PRI}} \\ &= H_{\text{CE}}(\tilde{G}, Y) + \alpha (H_{\text{vN}}(\tilde{G}) + \beta D_{\text{QJS}}(\tilde{G} || G)), \end{aligned} \quad (18)$$

where $H_{\text{CE}}(\cdot)$ is the cross-entropy loss for classification and α is a hyper-parameter to balance the two loss terms. The overall process of PRI-GSL is shown in Algorithm 1.

Algorithm 1: The overall process of PRI-GSL for node classification

Input: Graph G with node labels Y ; Number of training epochs $Epochs$; Wavelet scale set S ; Timepoints T ; Hyper-parameters α , β and γ .

Output: Refined graph $\tilde{G}$; Predicted labels $\hat{Y}$.

Parameter initialization;

for $e = 1, 2, \dots, Epochs$ **do**

// Structural Role Encoding

for $s \in S$ **do**

$\Psi_s^{(e-1)} \leftarrow \text{Eq. (9)}, \phi_s(v_a, t) \leftarrow \text{Eq. (11)};$

$\mathbf{h}_s^{(e-1)}(v_a) \leftarrow \text{Eq. (12)};$

end

$\mathbf{h}_a^{(e-1)} \leftarrow \text{Eq. (13)};$

// Graph Structure Learning

$a_{ij}^{R(e)} \leftarrow \text{Eq. (14)}, \mathbf{A}_R^{(e)} = \{a_{ij}^{R(e)}\};$

$\tilde{\mathbf{A}}^{(e)} \leftarrow \text{Eq. (16)}, \tilde{G}^{(e)} \leftarrow (\mathbf{X}, \tilde{\mathbf{A}}^{(e)});$

// Learn Node Representations

$\mathbf{Z}^{(e)} = \text{GNN-Encoder}(\tilde{\mathbf{A}}^{(e)}, \mathbf{X});$

// Optimize

$\mathcal{L}_{\text{PRI}}^{(e)} = H_{\text{vN}}(\tilde{G}^{(e)}) + \beta D_{\text{QJS}}(\tilde{G}^{(e)} || G);$

$\mathcal{L}^{(e)} = \mathcal{L}_{\text{cls}}^{(e)}(\tilde{G}^{(e)}, Y) + \alpha \mathcal{L}_{\text{PRI}}^{(e)};$

Update model parameters to minimize $\mathcal{L}^{(e)}$.

end

Approximation Recall that the computation of von Neumann entropy and spectral wavelets requires the full eigenvalue decomposition of the Laplacian matrix, which takes $O(N^3)$ time. The von Neumann entropy can be approximated with linear complexity $O(|V| + |E|)$ (Chen et al. 2019). In the experiments, we still use the basic von Neumann entropy. As for the graph wavelets, we use the Chebyshev polynomial approximation (Shuman, Vandergheynst, and Frossard 2011) to compute Ψ , reducing the computational complexity to $O(K|E|)$, where K is the order of Chebyshev polynomials.

Properties of $\tilde{G}$ Learned by PRI-GSL

PRI-GSL provides a unified way to control the quality of learned graph structure in terms of *sparsity*, *centrality*, and *nuisance invariance* property.

Sparsity and Centrality The von Neumann entropy has close connections with the structure sparsity and centrality. As indicated in (Passerini and Severini 2008), given a graph G , let $G' = G + \{x, y\}$ with $V(G') = V(G)$ and $E(G') = E(G) \cup \{x, y\}$, then $H_{\text{vN}}(\rho_{G'}) \geq \frac{d_{G'}-2}{d_{G'}} H_{\text{vN}}(\rho_G)$. The von Neumann entropy tends to grow with the increasing number of edges. The graph centrality (i.e., the extent to which a graph is organized around some central nodes) can be measured as a quantum relative entropy between the relative degree distribution and the uniform distribution (Simmons, Coon, and Datta 2018), which is given by $D(\rho_G || I_N) := \text{tr}(\rho_G(\log \rho_G - \log I/N)) = \log N - H_{\text{vN}}(G)$. The von

Neumann entropy tends to grow with the increasing regularity of the graph. The above conclusions suggest that minimizing $H_{\text{vN}}(\tilde{G})$ leads to a sparse and centralized structure.

Nuisance Invariance $\tilde{G}$ only preserves the most relevant yet least redundant information in the observed graph G and is invariant to nuisances in data. Suppose $G_n \in G$ the task-irrelevant nuisance in G , the relevance of $\tilde{G}$ and G_n can be formulated as the divergence between their conditional distributions predictions of the desired labels Y (Yu et al. 2020b): $\mathbb{E} [D_{\text{KL}}(p(Y|\tilde{G})) || p(Y|G_n)]$. Since G_n is irrelevant with Y , we have $p(Y|G_n) = p(Y)$. Minimizing the cross-entropy loss (i.e., the mutual information $I(Y; \tilde{G})$ between $\tilde{G}$ and Y) is equivalent to minimizing the relevance between $\tilde{G}$ and G_n :

$$\begin{aligned} & \mathbb{E} [D_{\text{KL}}(p(Y|\tilde{G})) || p(Y|G_n)] \\ &= \mathbb{E} [D_{\text{KL}}(p(Y|\tilde{G})) || p(Y)] \\ &= \iint \left(p(Y|\tilde{G}) \log \frac{p(Y|\tilde{G})}{p(Y)} \right) p(\tilde{G}) \quad (19) \\ &= \iint p(Y|\tilde{G}) \log \frac{p(Y, \tilde{G})}{p(Y)p(\tilde{G})} \\ &= I(Y; \tilde{G}). \end{aligned}$$

Experiments

We evaluate PRI-GSL on node classification and graph denoising tasks to verify its capability of improving the effectiveness and robustness of graph representation learning. Then we provide the analyses of the PRI loss, the structural role encodings, and the learned structure.

Experimental Settings

Datasets We select datasets with different homophily ratios h (Pei et al. 2020) to analyze methods' generalization on graphs with different properties. The evaluation datasets are Squirrel, Chameleon (Rozemberczki, Allen, and Sarkar 2021), Actor (Pei et al. 2020), CiteSeer, PubMed, Cora (Sen et al. 2008) and Photo (Shchur et al. 2018).

Baselines We consider three types of baselines: (1) *Graph neural networks*: GCN (Kipf and Welling 2016), GAT (Veličković et al. 2017) and GraphSAGE (Hamilton, Ying, and Leskovec 2017); (2) *Graph sparsification methods*: DropEdge (Rong et al. 2019), NeuralSparse (Zheng et al. 2020), and Graph-PRI (Yu et al. 2022); (3) *Graph structure learning methods*: IDGL (Chen, Wu, and Zaki 2020), Pro-GNN (Jin et al. 2020), SDRF (Topping et al. 2022), and SLAPS (Fatemi, El Asri, and Kazemi 2021).

Parameter Settings We re-implement the NeuralSparse (Zheng et al. 2020) and SDRF (Topping et al. 2022). The parameters of baseline methods are set to the suggested value in their papers or carefully tuned for

Method	Squirrel $h=0.22$	Actor $h=0.24$	Chameleon $h=0.25$	CiteSeer $h=0.72$	PubMed $h=0.79$	Cora $h=0.83$	Photo $h=0.83$
GCN	22.22±1.24	22.85±1.64	32.16±2.76	66.31±1.12	74.11±3.65	79.10±0.77	85.61±2.20
GAT	22.64±1.25	23.53±1.33	32.05±2.56	63.85±2.45	73.02±2.52	77.86±1.48	87.97±2.73
GraphSAGE	28.79±1.74	24.22±1.44	37.05±2.35	64.80±1.83	72.61±2.95	75.23±1.31	86.23±2.53
DropEdge	22.35±1.12	23.84±1.20	32.62±2.75	66.68±1.38	75.97±0.82	79.30±0.84	86.05±1.78
NeuralSparse	29.02±1.10	24.50±1.42	47.30±2.22	67.82±1.18	74.87±2.77	81.47±1.43	89.40±1.85
Graph-PRI	28.44±2.10	23.81±2.31	42.39±1.99	69.24±1.25	76.25±1.44	79.07±1.12	88.30±2.11
IDGL	29.13±2.94	27.44±5.80	49.80±4.80	67.94±0.28	75.32±1.45	83.23±0.62	88.89±2.55
Pro-GNN	27.18±1.28	24.82±2.81	48.54±4.87	66.68±2.02	75.44±3.54	82.14±0.58	87.28±1.85
SDRF	>1 day	>1 day	41.05±1.17	69.97±0.28	>1 day	81.94±0.59	>1 day
SLAPS	25.29±1.06	23.10±3.39	40.24±1.80	68.58±1.46	75.64±0.77	79.27±1.54	88.48±2.47
PRI-GSL	33.87±2.08	28.55±2.04	51.83±2.44	69.34±2.64	76.77±3.20	83.67±2.09	92.40±1.15

Table 1: Accuracy $\pm$ standard deviation (%) of node classification. (Bold: best result; Underlined: runner up.)

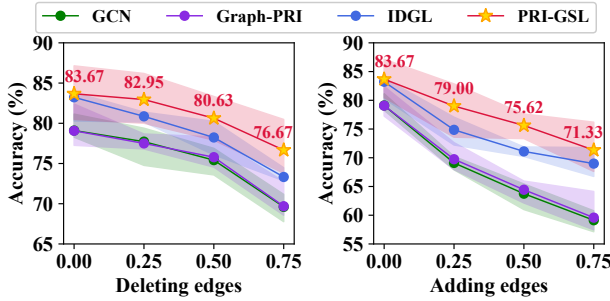


Figure 3: PRI-GSL on noisy graphs.

fairness. For the GNN encoders, we use a 2-layer GCN for node classification. We set the representation dimension $d=32$, the Chebyshev polynomial order $K=10$, the number of time points $T=4$, the number of scales $M=2$, and the number of heads $m=4$. The other hyper-parameters (α , β , and γ) are tuned for each dataset.

Evaluation Results and Analysis

Node Classification We set the number of nodes in each class to be 20/30 for training/validation and take the remaining nodes for the test. The accuracy and standard deviation on 10 randomly split are shown in Table 1. The best results are shown in bold and the runner-ups are underlined. Our PRI-GSL achieves the best performance on both homophilic and heterophilic datasets, showing the effectiveness of utilizing the self-organization property when mining the latent inherent structure. Generally, the graph structure learning methods show better performance than GNNs and graph sparsification methods. Pro-GNN and SLAPS perform well on the homophilic graphs but show unsatisfactory performance on the heterophilic graphs (Actor, Chameleon, and Squirrel), demonstrating the limitation of using heuristic assumptions for structural constraints. Although Graph-PRI also uses PRI for graph sparsification, it achieves fewer improvements since it does not take the node feature and the downstream task into consideration.

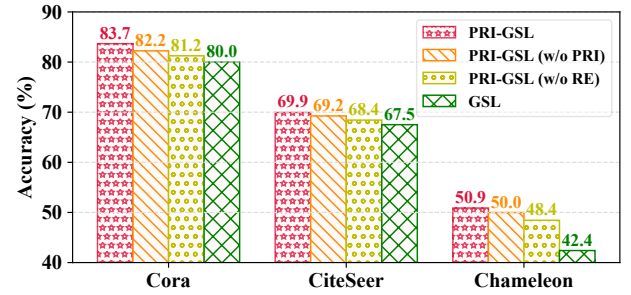


Figure 4: Ablation study results of PRI-GSL.

Graph Denoising To evaluate PRI-GSL’s ability to move noisy information, we generate synthetic noisy datasets by adding/deleting edges on Cora following (Chen, Wu, and Zaki 2020; Sun et al. 2022a). Specifically, we randomly add/delete 25%, 50%, 75% edges for 5 times and show the mean accuracy (solid line) and standard deviation (shaded region) in Fig. 3. The performance of GCN and Graph-PRI decreases dramatically with the increasing noise level. The structure learning methods, IDGL and PRI-GSL, show more robustness compared to vanilla GCN and Graph-PRI. Adding edge hurts more than deleting edges, indicating the importance of removing redundant information in the structure. PRI-GSL consistently shows better performance under different levels of external noise than the other baselines. Even though Graph-PRI shares the same spirit with PRI-GSL, it fails to distinguish whether the noise is task-relevant and shows the same poor robustness as GCN. The performance of PRI-GSL on noisy graphs also demonstrates the nuisance invariance property.

Ablation Study To illustrate the advantages of the guidance of PRI and structural role information, we compare PRI-GSL with three variants: (1) PRI-GSL (w/o PRI) that removes the PRI loss, (2) PRI-GSL (w/o RE) that removes the role encodings, and (3) GSL that removes both the PRI loss and the role encoding. The results of variants on 5 random split datasets are shown in Fig. 4. As we can observe, both the PRI loss and the structural role encoding benefit

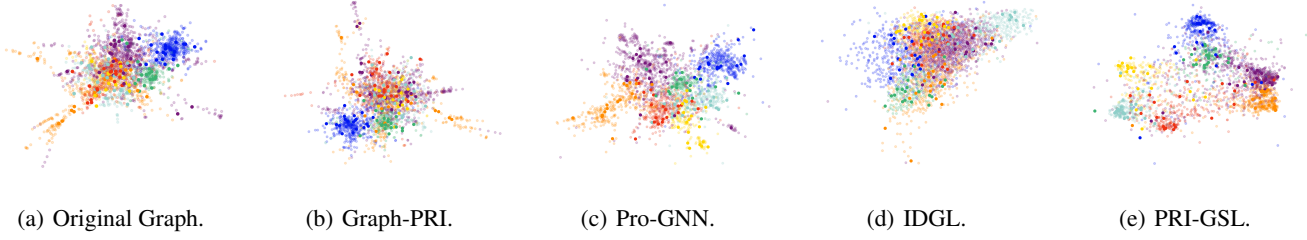


Figure 5: Visualization of the original graph of Cora and learned graphs by Graph-PRI, Pro-GNN, IDGL, and PRI-GSL.

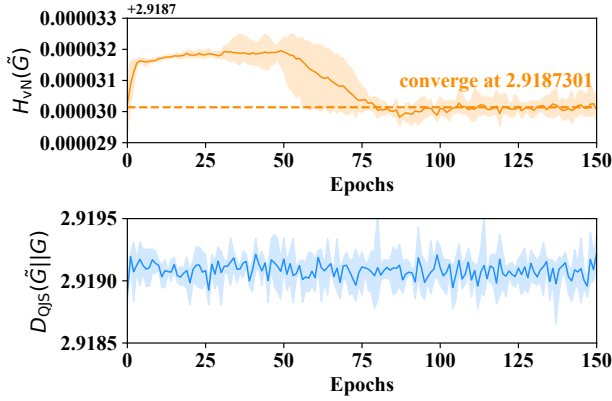


Figure 6: The variations of $H_{vN}(\tilde{G})$ and $D_{QJS}(\tilde{G}||G)$.

the classification, where the structural role encoding brings more improvement. This suggests that it's important to capture the node's contribution to the graph information flow when identifying the graph organization.

Empirical Behavior of $H_{vN}(\tilde{G})$ and $D_{QJS}(\tilde{G}||G)$ We analyze the learning dynamics of PRI-GSL by measuring the variations of $H_{vN}(\tilde{G})$ and $D_{QJS}(\tilde{G}||G)$ on Cora with $\alpha=0.1$ and $\beta=1$ in Fig. 6. The shadowed area is enclosed by the min and max value of four training runs. The solid line in the middle is the mean value of each epoch. $H_{vN}(\tilde{G})$ first increases for about 50 epochs with the structure exploration and then decreases to converge after about the 80-th epoch, indicating that the learned structure is with high certainty. $D_{QJS}(\tilde{G}||G)$ bumps during the training process. This may be because the model continues to seek a balance of structure redundancy and distortion during the training process.

Hyper-parameter Analysis We analyze the impact of hyper-parameters including α controlling the importance of the PRI loss in Eq. (18) and β trading off redundancy and distortion in Eq. (6). The results are shown in Fig. 7. PRI-GSL achieves the best performance with $\alpha=1$ on Cora and $\alpha=0.4$ on CiteSeer, which indicates that PRI-GSL benefits from the PRI loss. As for the β in the PRI loss, when the distortion term has a weight more than 2 compared to the redundancy term, PRI-GSL could reach satisfactory performance on both datasets. This suggests that the distortion term dominates the PRI loss in PRI-GSL. That is to say, the learned

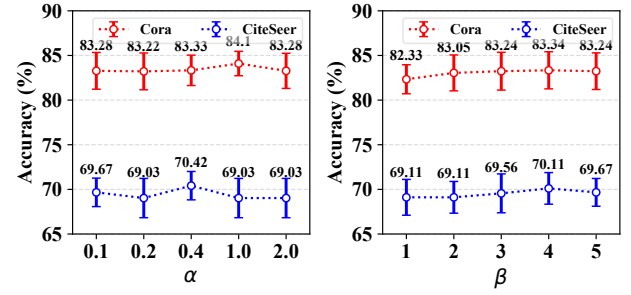


Figure 7: Parameter sensitivity of α and β .

structure should preserve enough information from the original graph to perform well on the downstream task.

Visualization

In Figure 5, we visualize the original graph structure of Cora and the graphs learned by PRI-Graph, Pro-GNN, IDGL, and PRI-GSL using *networkx*. The nodes' colors indicate their classes, the labeled nodes are solid and the unlabeled nodes are hollow. The edges are not shown for clarity and the layout of nodes represents their connectivities. Graph-PRI has little effect on the overall property of graph structure. Even though Pro-GNN and IDGL can make nodes within different classes more separate, there are still some overlapping and entangled areas. Benefiting from the structural role encoding, PRI-GSL can obtain the structure with separate clusters with similar shapes and clearer class boundaries, showing how the nodes within a class are organized.

Conclusion

In this work, we propose an information-theoretic framework for graph structure learning named PRI-GSL. We formulate the Principle of Relevant Information for graph data to quantify the structure redundancy and distortion, which acts as a potentially unifying guidance of structure learning. We propose a role-aware structure learner based on the quantum continuous walk evolution to unravel the self-organization of the graph. The learned structure enjoys the property of sparse, centralized, and nuisance invariance. Extensive experiments demonstrate the superior effectiveness and robustness of PRI-GSL.

Acknowledgments

The corresponding author is Jianxin Li. The authors are supported by the NSFC through grant No.U20B2053, and in part by NSF under grants III-1763325, III-1909323, III-2106758, and SaTC-1930941.

References

- Bai, L.; Rossi, L.; Torsello, A.; and Hancock, E. R. 2015. A quantum Jensen–Shannon graph kernel for unattributed graphs. *Pattern Recognition*, 48(2): 344–355.
- Bonabeau, E.; Theraulaz, G.; Deneubourg, J.-L.; Aron, S.; and Camazine, S. 1997. Self-organization in social insects. *Trends in ecology & evolution*, 12(5): 188–193.
- Chen, P.-Y.; Wu, L.; Liu, S.; and Rajapakse, I. 2019. Fast incremental von neumann graph entropy computation: Theory, algorithm, and applications. In *ICML*, 1091–1101. PMLR.
- Chen, Y.; Wu, L.; and Zaki, M. 2020. Iterative deep graph learning for graph neural networks: Better and robust node embeddings. In *NeurIPS*.
- Childs, A. M. 2010. On the relationship between continuous-and discrete-time quantum walk. *Communications in Mathematical Physics*, 294(2): 581–603.
- Coutinho, G.; and Godsil, C. 2021. Graph spectra and continuous quantum walks. *Unpublished notes*.
- Dasoulas, G.; Nikolentzos, G.; Scaman, K.; Virmaux, A.; and Vazirgiannis, M. 2020. Ego-based entropy measures for structural representations. *arXiv preprint arXiv:2003.00553*.
- De Domenico, M.; and Biamonte, J. 2016. Spectral entropies as information-theoretic tools for complex network comparison. *Physical Review X*, 6(4): 041062.
- De Domenico, M.; Nicosia, V.; Arenas, A.; and Latora, V. 2015. Structural reducibility of multilayer networks. *Nature communications*, 6(1): 1–9.
- Eigen, M.; and Schuster, P. 1977. A principle of natural self-organization. *Naturwissenschaften*, 64(11): 541–565.
- Fatemi, B.; El Asri, L.; and Kazemi, S. M. 2021. SLAPS: Self-supervision improves structure learning for graph neural networks. In *NIPS*, volume 34, 22667–22681.
- Gilmer, J.; Schoenholz, S. S.; Riley, P. F.; Vinyals, O.; and Dahl, G. E. 2017. Neural message passing for quantum chemistry. In *ICML*, 1263–1272. PMLR.
- Hamilton, W.; Ying, Z.; and Leskovec, J. 2017. Inductive representation learning on large graphs. In *NeurIPS*, 1024–1034.
- Hammond, D. K.; Vandergheynst, P.; and Gribonval, R. 2011. Wavelets on graphs via spectral graph theory. *Applied and Computational Harmonic Analysis*, 30(2): 129–150.
- Hoyos-Osorio, J.; Alvarez-Meza, A.; Daza-Santacoloma, G.; Orozco-Gutierrez, A.; and Castellanos-Dominguez, G. 2021. Relevant information undersampling to support imbalanced data classification. *Neurocomputing*, 436: 136–146.
- Jin, W.; Ma, Y.; Liu, X.; Tang, X.; Wang, S.; and Tang, J. 2020. Graph structure learning for robust graph neural networks. In *ACM SIGKDD*, 66–74.
- Kipf, T. N.; and Welling, M. 2016. Semi-Supervised Classification with Graph Convolutional Networks. In *ICLR*.
- Lamberti, P. W.; Majtey, A. P.; Borrás, A.; Casas, M.; and Plastino, A. 2008. Metric character of the quantum Jensen-Shannon divergence. *Physical Review A*, 77(5): 052311.
- Minello, G.; Rossi, L.; and Torsello, A. 2019. On the von Neumann entropy of graphs. *Journal of Complex Networks*, 7(4): 491–514.
- Nielsen, M. A.; and Chuang, I. 2002. Quantum computation and quantum information. *American Association of Physics Teachers*.
- Passerini, F.; and Severini, S. 2008. The von Neumann entropy of networks. *arXiv preprint arXiv:0812.2597*.
- Pei, H.; Wei, B.; Chang, K. C.-C.; Lei, Y.; and Yang, B. 2020. Geom-gcn: Geometric graph convolutional networks. In *ICLR*.
- Principe, J. C. 2010. *Information theoretic learning: Renyi’s entropy and kernel perspectives*. Springer Science & Business Media.
- Rényi, A.; et al. 1961. On measures of entropy and information. In *Proceedings of the fourth Berkeley symposium on mathematical statistics and probability*, volume 1. Berkeley, California, USA.
- Rong, Y.; Huang, W.; Xu, T.; and Huang, J. 2019. Dropedge: Towards deep graph convolutional networks on node classification. In *ICLR*.
- Rozemberczki, B.; Allen, C.; and Sarkar, R. 2021. Multi-scale attributed node embedding. *Journal of Complex Networks*, 9(2): cnab014.
- Sen, P.; Namata, G.; Bilgic, M.; Getoor, L.; Galligher, B.; and Eliassi-Rad, T. 2008. Collective classification in network data. *AI magazine*, 29(3): 93–93.
- Shchur, O.; Mumme, M.; Bojchevski, A.; and Günnemann, S. 2018. Pitfalls of graph neural network evaluation. *arXiv preprint arXiv:1811.05868*.
- Shuman, D. I.; Vandergheynst, P.; and Frossard, P. 2011. Chebyshev polynomial approximation for distributed signal processing. In *DCOSS*, 1–8. IEEE.
- Simmons, D. E.; Coon, J. P.; and Datta, A. 2018. The von Neumann Theil index: characterizing graph centralization using the von Neumann index. *Journal of Complex Networks*, 6(6): 859–876.
- Sun, L.; Dou, Y.; Yang, C.; Wang, J.; Yu, P. S.; He, L.; and Li, B. 2018. Adversarial attack and defense on graph data: A survey. *arXiv preprint arXiv:1812.10528*.
- Sun, Q.; Li, J.; Peng, H.; Wu, J.; Fu, X.; Ji, C.; and Philip, S. Y. 2022a. Graph Structure Learning with Variational Information Bottleneck. In *AAAI*, volume 36, 4165–4174.
- Sun, Q.; Li, J.; Peng, H.; Wu, J.; Ning, Y.; Yu, P. S.; and He, L. 2021. SUGAR: Subgraph neural network with reinforcement pooling and self-supervised mutual information mechanism. In *Web Conference*, 2081–2091.
- Sun, Q.; Li, J.; Yuan, H.; Fu, X.; Peng, H.; Ji, C.; Li, Q.; and Yu, P. S. 2022b. Position-aware structure learning for graph topology-imbalance by relieving under-reaching and over-squashing. In *CIKM*, 1848–1857.

Tishby, N.; Pereira, F. C.; and Bialek, W. 2000. The information bottleneck method. *arXiv preprint physics/0004057*.

Topping, J.; Di Giovanni, F.; Chamberlain, B. P.; Dong, X.; and Bronstein, M. M. 2022. Understanding over-squashing and bottlenecks on graphs via curvature. In *ICLR*.

Veličković, P.; Cucurull, G.; Casanova, A.; Romero, A.; Lio, P.; and Bengio, Y. 2017. Graph Attention Networks. In *ICLR*.

Wang, R.; Mou, S.; Wang, X.; Xiao, W.; Ju, Q.; Shi, C.; and Xie, X. 2021. Graph Structure Estimation Neural Networks. In *Web Conference*, 342–353.

Wei, Y.; Yu, S.; Giraldo, L. S.; and Príncipe, J. C. 2021. Multiscale principle of relevant information for hyperspectral image classification. *Machine Learning*, 1–26.

Wu, Z.; Pan, S.; Chen, F.; Long, G.; Zhang, C.; and Philip, S. Y. 2020. A comprehensive survey on graph neural networks. *IEEE transactions on neural networks and learning systems*, 32(1): 4–24.

Yu, J.; Xu, T.; Rong, Y.; Bian, Y.; Huang, J.; and He, R. 2020a. Graph Information Bottleneck for Subgraph Recognition. In *ICLR*.

Yu, S.; Alesiani, F.; Yin, W.; Jenssen, R.; and Principe, J. C. 2022. Principle of Relevant Information for Graph Sparsification. In *UAI*.

Yu, S.; Shaker, A.; Alesiani, F.; and Principe, J. C. 2020b. Measuring the discrepancy between conditional distributions: Methods, properties and applications. In *IJCAI*.

Zheng, C.; Zong, B.; Cheng, W.; Song, D.; Ni, J.; Yu, W.; Chen, H.; and Wang, W. 2020. Robust graph representation learning via neural sparsification. In *ICML*, 11458–11468.

Zhu, Y.; Xu, W.; Zhang, J.; Liu, Q.; Wu, S.; and Wang, L. 2022. Deep Graph Structure Learning for Robust Representations: A Survey. In *IJCAI*.

Zügner, D.; Akbarnejad, A.; and Günnemann, S. 2018. Adversarial attacks on neural networks for graph data. In *ACM SIGKDD*, 2847–2856.