

Graph Representation Learning based on Deep Generative Gaussian Mixture Models

Ghazaleh Niknam^{a,*}, Soheila Molaei^{b,*}, Hadi Zare^{a,**}, David Clifton^{b,d}, Shirui Pan^c

^a*Department of Network Science and Technologies, University of Tehran*

^b*Department of Engineering Science, University of Oxford*

^c*School of Information and Communication Technology, Griffith University*

^d*Oxford-Suzhou Centre for Advanced Research (OSCAR), Suzhou, China*

Abstract

Graph representation learning is an effective tool for facilitating graph analysis with machine learning methods. Most GNNs, including Graph Convolutional Networks (GCN), Graph Recurrent Neural Networks (GRNN), and Graph Auto-Encoders (GAE), employ vectors to represent nodes in a deterministic way without exploiting the uncertainty in hidden variables. Deep generative models are combined with GAE in the Variational Graph Auto-Encoder (VGAE) framework to address this issue. While traditional VGAE-based methods can capture hidden and hierarchical dependencies in latent spaces, they are limited by the data's multimodality. Here, we propose the Gaussian Mixture Model (GMM) to model the prior distribution in VGAE. Furthermore, an adversarial regularization is incorporated into the proposed approach to ensure the fruitful impact of the latent representations on the results. We demonstrate the performance of the proposed method on clustering and link prediction tasks. Our experimental results on real datasets show remarkable performance compared to state-of-the-art methods.

Keywords: Graph representation learning, Node Embedding, Variational Graph Auto-Encoder, Adversarial mechanism, Gaussian Mixture Model

*Equal Contribution

**Corresponding author

1. Introduction

Graphs can model and represent the majority of complex systems across a wide range of domains [1]. These structures play a critical role in many applications, including biological protein-protein interaction networks [2], social media [3], transportation networks [4], and citation networks[5]. Recently, many studies have emerged on learning representations to encode structural information of the graph [6, 7, 8, 9, 10]. These graph representation learning algorithms convert graph data into a low-dimensional space. Downstream machine learning tasks such as node classification, link prediction, and clustering employ these representations [11].

Graph representation learning approaches can be divided into two categories: shallow embedding and deep embedding. Despite their popularity in recent years, shallow embedding approaches suffer from some limitations. Some of these limitations can be mentioned as follows: 1) Lack of parameter sharing, 2) Incapability to handle node features efficiently, and 3) Inability to apply in an inductive manner [12, 13]. On the other hand, the emergence of deep embedding approaches, aided by the introduction of the Graph Neural Network (GNN) paradigm, assists in overcoming these shortcomings. GNNs can generate representations of nodes based on the graph’s structure, and any feature information [14, 12].

GNNs ignore the data distribution, leading to poor representations and overfitting [6, 15]. Combining them with deep generative models to learn data distribution has significantly improved generated representations. Variational Graph Auto-Encoder (VGAE) framework and adversarial approach are widely studied in this scope. VGAE is an unsupervised framework based on deep generative and variational inference models [16, 17]. This framework considers the distribution of the encoder’s latent representation, which leads to complex data reasoning and, as a result, efficient embedding [18]. Adversarial Graph Auto-Encoder framework enforces the latent representation to match a prior distribution during the training process [12].

31 This paper considers some assumptions about observed data distributions
 32 for use in unsupervised clustering tasks. To accomplish this, we concentrate on
 33 modeling the multimodality of observed data in clustering. Probabilistic model-
 34 ing is mostly performed by considering simplified assumptions about data. Each
 35 sample is generated from one well-known distribution like Gaussian, named uni-
 36 modality assumption. This assumption has some major drawbacks and is too
 37 limited for certain phenomena. Researchers are usually faced with complex data
 38 where each sample may come from one of the multiple known (or unknown) dis-
 39 tributions called multimodal distributions. Intuitively, multimodal data contain
 40 several regions with a high amount of probability.

41 Technically, these models are called mixture models as a principled modeling
 42 approach to deal with such complex data. This paper uses the Gaussian Mixture
 43 Model (GMM) in VGAE to model the prior distribution. This combination
 44 improves the interpretability of the proposed model. Along with modelling with
 45 multimodality through GMM, the proposed method introduces a regularizer
 46 based on the adversarial mechanism concept to improve the results. The outline
 47 of our contributions is as follows,

- 48 • We leverage the VGAE framework by considering the Gaussian Mixture
 49 Model (GMM) to model the prior distribution of observed data.
- 50 • We introduce a regularization based on the adversarial mechanism concept
 51 to ensure that the latent embeddings boost our results.
- 52 • We demonstrate our model’s performance on clustering and link prediction
 53 tasks on real datasets.

54 **2. Related Work**

55 Deep learning methods have sparked widespread interest due to their high per-
 56 formance, efficiency, and ease of use. As a result, deep graph embedding ap-
 57 proaches based on the GNN paradigm have become the main focus of graph
 58 representation learning methods. GNNs can use parameter sharing to generate

node representations based on the graph’s structure and feature information. They can also appear inductively. GNNs are useful and widely used structures for the reasons stated above [14, 12, 13].

Graph Convolutional Network (GCN) [19] is one of the first GNN structures which generalizes the concept of convolutions to graph-structured data [20, 19]. Graph Auto-Encoder (GAE) is another approach introduced following the GCN, which is an unsupervised approach that generalizes the auto-encoder framework. GAE is a straightforward and robust framework comprised of an encoder and a decoder. A GCN model, which generates latent representations, is frequently used as the encoder. The decoder then reconstructs the input data, which is frequently an inner product of the latent variables [17].

Multi-Task Graph Auto-encoder (MTGAE) [21] is an auto-encoder-based method that can learn a joint representation of local graph structure and available node features. This method aims to learn link prediction and node classification simultaneously. There are some other studies based on GAE framework, such as [22, 23, 24, 25, 26]. Although GAE-based methods are effective, they disregard data distribution, leading to poor representations and overfitting [6, 15]. The GAE framework and deep generative models are combined for this purpose. By accounting for data distribution, deep generative models can represent complex dependencies and interactions between input and output data[27].

VGAE approach originates from deep generative concepts that generalize the Variational Auto-Encoder (VAE) [17, 27] method on graph-structured data. In VGAE, the encoder and the decoder are probabilistic. The main idea behind VGAE is that it embeds input data into a distribution rather than a point. A random sample Z is drawn from the distribution rather than generated directly by the encoder[28]. Using the VGAE framework, Xie et al. [29] provided a representation learning method for networked documents to model document contents and relations. Their proposed method reached from the Higher-order Graph Attention Network (HGAT), which examines and shuffles the document’s neighbourhood information in each order.

Multiresolution Graph Networks (MGN) and Multiresolution Graph Varia-

90 tional Autoencoders (MGVAE) are presented by [30] for using multiresolution
 91 and equivariant methods to learn and generate graphs. MGN performs higher-
 92 order message passing while designing graph encoders to cluster the graph and
 93 coarsen it to a lower resolution. They have encoded the hierarchy of coarsened
 94 graphs by MGVAE, which develops a hierarchical generative model based on
 95 MGN. NetVAE developed by Jin et al. [31] is a network embedding approach
 96 that addresses the orthogonality of network topology information and node at-
 97 tributes. Here, network structure compression and node attributes share the
 98 same encoder. On the other hand, a dual decoder, supported by GMM, is in-
 99 troduced for reconstructing network topologies and node attributes separately.

100 Adversarial-based approaches leverage alternative deep generative frame-
 101 works. Adversarially Regularized Variational Graph Autoencoder (ARVGAE)
 102 generalizes the adversarial autoencoder framework [32] to graph data [6]. In the
 103 ARVGAE, the latent representation is enforced to match a prior distribution
 104 by exploiting a min-max game strategy [6]. Random Walk Regularization for
 105 Graph Auto Encoders (RWR-VGAE) is a VGAE with a regularizer based on
 106 Random Walk with Restarts (RWR). This regularizer is used to capture the in-
 107 formation of the immediate neighbour of each node [33]. Lu et al. [34] proposed
 108 MRGAE, a network embedding approach to model network consistency across
 109 different views. MRGAE generates a second view of the input network based
 110 on node content capturing the relationship between them. By incorporating a
 111 multiview adversarial regularization module, they ensure consistency between
 112 the two views of the network.

113 Most of the earlier works assumed a Gaussian distribution to model the
 114 prior uncertainty of the observed data. This assumption suffers from the ineffi-
 115 ciency of modeling complex data with properties like multimodality. Some works
 116 employed other distributions rather than the traditional Gaussian assumption.
 117 Davidson et al. [35] introduced a method called S-VGAE, which considered the
 118 von Mises-Fisher (vMF) unimodal density to model the prior distribution and
 119 result in a better representation.

120 Zheng et al. [36] proposed DBGAN, a method for estimating the prior dis-

121 tribution of latent representation in a structure-aware manner. Rather than the
122 widely used Gaussian distribution assumption implicitly bridges the graph and
123 feature space through prototype learning. As a result, discriminative and robust
124 representations are generated for all nodes. The CGCN proposed by Hui et al.
125 [37] consists of two main modules, an attributed graph clustering module and a
126 semi-supervised node classification module. This approach investigates a mix-
127 ture of Gaussian in the clustering module to assign pseudo-labels to unlabeled
128 samples. CGCN ignores the benefits of the adversarial mechanism.

129 Our proposed method is based on the VGAE framework and uses a vari-
130 ant of the adversarial mechanism. Along with our primary aim of clustering
131 tasks, we consider the multimodal assumption of the observed data through a
132 theoretically and practically well-known GMM model. More specifically, the
133 correspondence of Gaussian density for each presumed cluster results in a more
134 efficient approach. Furthermore, a regularization based on a variant of the ad-
135 versarial mechanism is proposed to improve the result.

136 Focusing on positive and negative encoder samplings, we present a novel
137 perspective on the adversarial technique. Unlike the conventional adversarial-
138 based approaches [6, 33, 34], which treated encoder-generated representations as
139 negative samples, we address them as positive samples. Following our shift from
140 negative samples to positive ones and changing the nature of the regularizer, we
141 appraise the overall loss function. The inverse of KL-divergence is added to the
142 loss function, whereas traditional methods use the min-max technique.

143 3. Problem Statement

144 3.1. Notation

145 A graph G is represented as $G = (V, E)$, where $V = \{v_1, \dots, v_N\}$ denotes a
146 set of N nodes and E represents a set of ε edges in this graph. $\mathbf{X} \in R^{N \times d}$
147 denotes a feature matrix where each row shows a d -dimensional feature vector
148 for each node v_i in the graph. $\mathbf{A} \in R^{N \times N}$ represents the adjacency matrix
149 where $A_{i,j} = 1$ if $e_{i,j} \in E$, otherwise $A_{i,j} = 0$. The goal is to map the nodes

150 $v_i \in V$ to low-dimensional vectors $\mathbf{z}_i \in R^F$. The mapping function is as follows:
 151 $f : (\mathbf{A}, \chi) \rightarrow \mathbf{Z}$, where $\mathbf{Z} \in R^{N \times F}$. Each node embedding $\mathbf{z}_i \in \mathbf{Z}$ is an F -
 152 dimensional vector in latent space, where $F \ll d$. Table 1 summarizes the notations used in this paper.

Table 1: The summary of notations.

Symbols	Meaning
$\mathbf{A}$	The adjacency matrix of graph G
N	Number of nodes in the graph
χ	The features matrix of G
F	Number of features in χ
$D_{KL}(q p)$	KL-divergence between q and p
$\mathbf{z}$	The latent variable in Variational Inference
$\mathbf{X}, \mathbf{W}, \mathbf{C}$	The latent variables in the proposed method
$\mathbf{X}', \mathbf{W}', \mathbf{C}'$	The latent variables in the proposed method of fake data for regularization part
ϕ	The set of parameters of the neural network in the encoder part
θ	The set of parameters of the neural network in the decoder part
β	The set of parameters of the GNN related to each GMM component in the proposed method
ϕ_C	The set of parameters of the GNN related to C in variational factors in the proposed method
ϕ_W	The set of parameters of the GNN related to W in variational factors in the proposed method
K	Number of components in GMM
π	Mixing probability of GMM

153

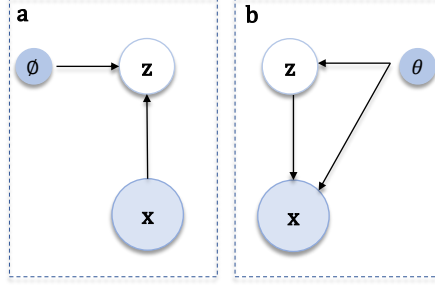


Figure 1: Graphical models for Variational Auto-Encoders. **a.** The recognition or inference model maps input data to the latent variable $\mathbf{z}$. **b.** The generative model reconstructs the input data by the latent variable.

4. Preliminaries

4.1. Variational Graph Auto-Encoders

VGAE was introduced by [28]. This framework is a generalization of the VAE framework [27, 17] to graph structure data. There are two values z (embeddings) and x (input data), in the VAE framework. Based on the Bayes theorem, the posterior distribution is calculated as shown in Equation (1).

$$p_{\theta}(z|x) = \frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(x)} = \frac{p_{\theta}(x|z)p_{\theta}(z)}{\int p_{\theta}(z)p_{\theta}(x|z)dz} \quad (1)$$

Since the true posterior density, shown in Equation (1), is mostly intractable, Rezende et al. [27] introduced a recognition model $q_{\phi}(z|x)$ as an approximation for the posterior distribution. The recognition model commonly chooses Gaussian distribution with a diagonal covariance structure. Consequently, the recognition model (or inference model), $q_{\phi}(z|x)$ is the probabilistic encoder of VAE, and the generative model, $p_{\theta}(x|z)$ is the probabilistic decoder of VAE. A Graphical model for VAE is shown in Figure 1.

Inference model of VGAE. In VGAE, the inputs are the adjacency matrix A , and the feature matrix χ . Each node in these inputs maps to a latent vector

as shown in Equation (2).

$$q_\phi(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi}) = \prod_{i=1}^N q_\phi(\mathbf{z}_i|\mathbf{A}, \boldsymbol{\chi}), \quad (2)$$

$$q_\phi(\mathbf{z}_i|\mathbf{A}, \boldsymbol{\chi}) = \mathcal{N}(\mathbf{z}_i|\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$$

Here $\boldsymbol{\mu}$ represents a matrix of mean vectors and $\boldsymbol{\Sigma}$ denotes covariance matrix. $\boldsymbol{\mu}_i$ and $\boldsymbol{\Sigma}_i$ denote i -th vector of these matrices. The distribution of the recognition model is parameterized by two GNNs as presented in Equation (3).

$$\begin{aligned} \boldsymbol{\mu} &= GNN_\mu(\mathbf{A}, \boldsymbol{\chi}) \\ \boldsymbol{\Sigma} &= GNN_\Sigma(\mathbf{A}, \boldsymbol{\chi}) \end{aligned} \quad (3)$$

Here, GNN_μ and GNN_Σ are two GNNs that produce $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$. These GNNs can be any of the different types of GNNs such as GCN [19], GCN with Chebyshev filters [38], GAT (Graph Attention Network)[39], and GraphSAGE [40]. The two-layer GCN shown in Equation (4) is commonly used for this purpose [28].

$$GCN(\mathbf{A}, \boldsymbol{\chi}) = \widehat{\mathbf{A}} ReLU(\widehat{\mathbf{A}}\boldsymbol{\chi}\mathbf{W}_0)\mathbf{W}_1 \quad (4)$$

Here, $\mathbf{W}_0$ and $\mathbf{W}_1$ are the weight matrices of each layer, and $\mathbf{W}_0$ is shared in the first layer of two GCNs. $\widehat{\mathbf{A}}$ is the normalized adjacency matrix introduced as follows,

$$\begin{aligned} \tilde{\mathbf{A}} &= \mathbf{A} + \mathbf{I} \\ \tilde{D}_{ii} &= \sum_j \tilde{A}_{ij} \\ \widehat{\mathbf{A}} &= \tilde{\mathbf{D}}^{-\frac{1}{2}} \tilde{\mathbf{A}} \tilde{\mathbf{D}}^{-\frac{1}{2}} \end{aligned} \quad (5)$$

In Equation (5), $\mathbf{I}$ represents the identity matrix of size $\mathbf{A}$, and $\mathbf{D}$ represents degree matrix of the graph.

Generative model of VGAE. After inferring the latent vectors, the adjacency matrix can be reconstructed with sampling from $p_\theta(\mathbf{A}|\mathbf{Z})$. The generative model can be any kind of GNN, but most of the time, an inner product between latent

185 variables, as shown in Equation (6), is used for this purpose.

$$p(\mathbf{A}|\mathbf{Z}) = \prod_{i=1}^N \prod_{j=1}^N p(A_{ij}|\mathbf{z}_i, \mathbf{z}_j) \quad (6)$$

$$p(A_{ij}|\mathbf{z}_i, \mathbf{z}_j) = \text{Sigmoid}(\mathbf{z}_i^T \mathbf{z}_j)$$

186 Here, $\mathbf{z}_i$ and $\mathbf{z}_j$ are the i th and j th vectors of $\mathbf{Z}$ and *Sigmoid* denotes the logistic
187 sigmoid function.

188 **Learning of VGAE.** To learn the model, the approximate posterior, which is
189 parameterized by variational parameters ϕ , should be close to the true posterior.
190 The Kullback–Leibler divergence (KL-divergence) is used for this purpose. KL-
191 divergence is a metric that measures the distance between two distributions.
192 The goal is to minimize this KL-divergence as shown in Equation (7) [28].

$$\phi = \operatorname{argmin}_{\phi} D_{KL}[q_{\phi}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi}) || p_{\theta}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi})] \quad (7)$$

193 In this equation, D_{KL} indicates KL-divergence, ϕ represents the parameters
194 of posterior estimation ($q_{\phi}(\mathbf{Z}|\mathbf{A})$), and θ denotes parameters of real posterior
195 ($p_{\theta}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi})$). Calculation of KL-divergence depends on log marginal likelihood,
196 which is intractable. Instead, maximizing the Evidence Lower Bound (ELBO),
197 shown in Equation (8), is used for the learning process [28].

$$L_{ELBO} = \mathbb{E}_{q_{\phi}}(\log \frac{p_{\theta}(\mathbf{A}|\mathbf{Z})}{q_{\phi}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi})}) \quad (8)$$

198 Which can be rewritten as,

$$L_{ELBO} = \mathbb{E}_{q_{\phi}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi})}[\log p_{\theta}(\mathbf{A}|\mathbf{Z})] - D_{KL}(q_{\phi}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi}) || p_{\theta}(\mathbf{Z})) \quad (9)$$

199 L_{ELBO} denotes the ELBO loss function, which has two terms. The first term is
200 associated with data reconstruction in a generative model. The second term (the
201 KL-divergence term) is a loss function regularizer that measures the distance
202 between the posterior estimation $q_{\phi}(\mathbf{Z}|\mathbf{A}, \boldsymbol{\chi})$ and the prior distribution $p_{\theta}(\mathbf{Z})$.
203 This equation should be optimized concerning θ and ϕ by some optimization
204 method such as SGD (Stochastic Gradient Descent), Adagrad, Adam, or mini-
205 batch stochastic gradient descent [41, 42].

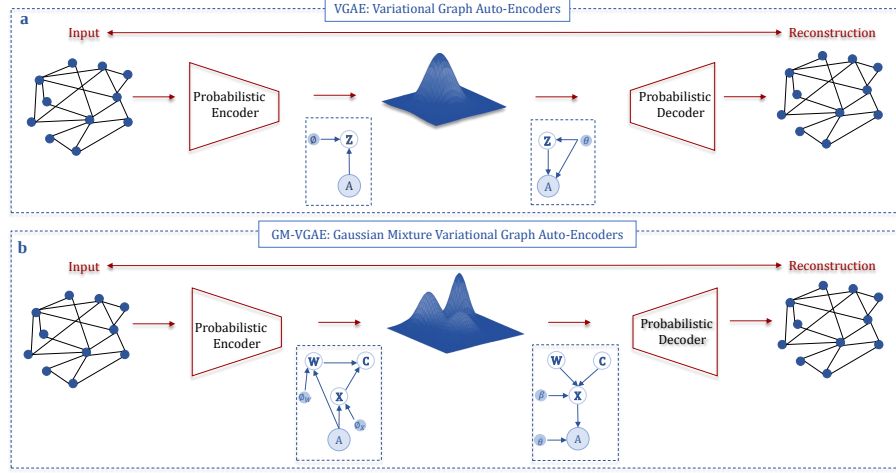


Figure 2: Comparison of the general framework of Variational graph auto-encoder (VGAE) and Gaussian mixture Variational graph auto-encoder (GM-VGAE). a) In VGAE, the input data as an adjacency matrix $\mathbf{A}$ maps to a unimodal latent space. The inference model (encoder), with parameter ϕ , maps input data to the latent variable $\mathbf{Z}$. The generative model (decoder), with parameter ϕ , reconstructs the input data by the latent variable. b) In GM-VGAE, input data ($\mathbf{A}$) maps to a multimodal latent space. This latent space comprises three latent variables $\mathbf{X}$, $\mathbf{W}$, and $\mathbf{C}$, which follow a GMM. The inference model maps input data to the latent variables. ϕ_C and ϕ_W are GNN parameters related to $\mathbf{C}$ and $\mathbf{W}$ in variational factors of the encoder, respectively. The decoder with parameters θ and β reconstructs the input data.

4.2. Gaussian Mixture Models

GMM belongs to the parametric probability density function family. It is a weighted sum of K components. Each of these components follows a multivariate Gaussian distribution. The weight of mixing, also known as mixture probability or mixture coefficient, is denoted by π_i for i -th component, where $\sum_{i=1}^K \pi_i = 1$. The components of GMM can capture the multimodal nature of the data [43].

213 5. The proposed model

214 5.1. Gaussian Mixture Variational Graph Auto-Encoder

215 In regular VGAE, the prior distribution over the latent variables is assumed to
 216 be a Gaussian distribution [28]. While VGAE is applied to many domains, it
 217 suffers from facing multimodality observed data. Figure 2 shows the general
 218 framework of the regular VGAE and our proposed method, called Gaussian
 219 Mixture Variational Graph Auto-Encoder (GM-VGAE). There is a set of latent
 220 variables $\mathbf{X}$, $\mathbf{W}$, and $\mathbf{C}$ in the GM-VGAE instead of just one latent variable in
 221 VGAE.

222 **Generative model of GM-VGAE.** The definition of the generative model,
 223 based on the latent variables and the observed sample, is shown in Equation (10).

$$\begin{aligned} \mathbf{W} &\sim \mathcal{N}(0, \mathbf{I}) \\ \mathbf{C} &\sim \text{Mult}(\pi) \\ \mathbf{X}|\mathbf{W}, \mathbf{C} &\sim \prod_{k=1}^K \mathcal{N}(\boldsymbol{\mu}_{\mathbf{c}_k}(\mathbf{W}; \beta), \boldsymbol{\Sigma}_{\mathbf{c}_k}(\mathbf{W}; \beta))^{\mathbf{c}_k} \end{aligned} \quad (10)$$

224 Here, K represents the number of components in the mixture model. This pa-
 225 rameter is defined as a hyperparameter of the model. $\mathbf{W}$ follows a Gaussian
 226 distribution with mean zero and covariance matrix $\mathbf{I}$. Eventually, $\mathbf{C}$ is a one-hot
 227 vector, which represents the mixing coefficient of the Gaussian mixture compo-
 228 nents. This vector is sampled from the mixing probability, which is denoted by
 229 π . Here, π is equal to $\frac{1}{K}$ so that it is uniformly distributed.

230 In this model, there is a GNN parameterized by β , which $\mathbf{W}$ feeds to it as
 231 input. This GNN produces a set of K numbers of $(\boldsymbol{\mu}_{\mathbf{c}_k})$ and $(\boldsymbol{\Sigma}_{\mathbf{c}_k})$. Here, each of
 232 $\boldsymbol{\mu}_{\mathbf{c}_k}$ and $\boldsymbol{\Sigma}_{\mathbf{c}_k}$ computed with a GCN with parameter sharing. Finally, according
 233 to the definitions provided, $\mathbf{X}|\mathbf{W}$ is a Gaussian mixture. The adjacency matrix
 234 $\mathbf{A}$ is reconstructed from another GNN parameterized by θ . The generative
 235 model is shown in Equation (11).

$$\mathbf{A}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}(\mathbf{X}; \theta), \boldsymbol{\Sigma}(\mathbf{X}; \theta)) \quad (11)$$

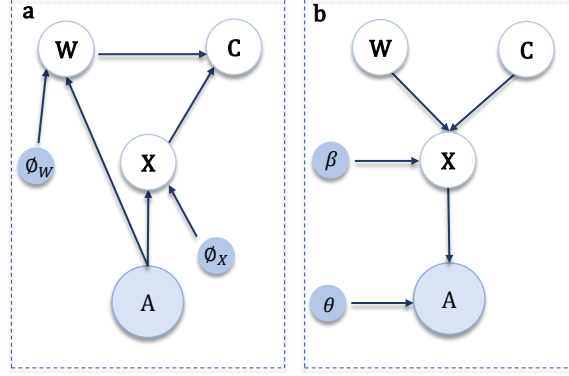


Figure 3: Graphical models for Gaussian mixture variational graph autoencoder. **a.** The recognition or inference model, which maps input data to the latent variables $\mathbf{X}$, $\mathbf{W}$, and $\mathbf{C}$. **b.** The generative model reconstructs the input data by the latent variables.

In this study, the choice of decoder is followed by [28] to avoid complexity. So we use an inner product between latent variables to reconstruct the adjacency matrix instead of a GNN. This decoder is shown in Equation (12).

$$p(\mathbf{A}|\mathbf{X}) = \prod_{i=1}^N \prod_{j=1}^N p(A_{ij}|\mathbf{x}_i, \mathbf{x}_j) \quad (12)$$

$$p(A_{ij}|\mathbf{x}_i, \mathbf{x}_j) = \text{Sigmoid}(\mathbf{x}_i^T \mathbf{x}_j)$$

Here A_{ij} represents the elements of the adjacency matrix. Note that in this case, the decoder, which generates the reconstruction of observed data, lacks parameters.

Inference model of GM-VGAE. Based on the mean-field variational family,

$$q(\mathbf{X}, \mathbf{W}, \mathbf{C}|\mathbf{A}, \chi) = \prod_{i=1}^N q_{\phi_{\mathbf{X}}}(\mathbf{x}_i|\mathbf{A}_i, \chi_i) q_{\phi_{\mathbf{W}}}(\mathbf{w}_i|\mathbf{A}_i, \chi_i) p_{\beta}(\mathbf{c}_i|\mathbf{x}_i, \mathbf{w}_i) \quad (13)$$

In this equation, β represents the set of GNN parameters related to each GMM component, ϕ_X demonstrates the set of GNN parameters related to X in variational factors, and ϕ_W shows the set of GNN parameters related to W in

246 variational factors. The z-posterior is as follows,

$$\begin{aligned} p_\beta(\mathbf{c}_j = 1|\mathbf{X}, \mathbf{W}) &= \frac{p(\mathbf{c}_j = 1)p(\mathbf{X}|\mathbf{c}_j = 1, \mathbf{W})}{\sum_{k=1}^K p(\mathbf{c}_k = 1)p(\mathbf{X}|\mathbf{c}_k = 1, \mathbf{W})} \\ &= \frac{\pi_j \mathcal{N}(\mathbf{X}|\mu_j(\mathbf{W}; \beta), \sigma_j(\mathbf{W}; \beta))}{\sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{X}|\mu_k(\mathbf{W}; \beta), \sigma_k(\mathbf{W}; \beta))} \end{aligned} \quad (14)$$

247 The graphical illustration of generative and inference models of our method is
248 demonstrated in Figure 3.

249 **Learning of GM-VGAE.** The ELBO of the proposed model is shown in
250 Equation (15).

$$L_{ELBO} = \mathbb{E}_q\left(\frac{p(\mathbf{A}, \mathbf{X}, \mathbf{W}, \mathbf{C})}{q(\mathbf{X}, \mathbf{W}, \mathbf{C}|\mathbf{A}, \boldsymbol{\chi})}\right) \quad (15)$$

251 in which,

$$p(\mathbf{A}, \mathbf{X}, \mathbf{W}, \mathbf{C}) = p(\mathbf{W})p(\mathbf{C})p(\mathbf{X}|\mathbf{W}, \mathbf{C})p(\mathbf{A}|\mathbf{X}) \quad (16)$$

252 Based on Equation (13) and Equation (16) the ELBO can be written as,

$$\begin{aligned} L_{ELBO} = & \mathbb{E}_{q(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})}[\log p(\mathbf{A}|\mathbf{X})] - \mathbb{E}_{q(\mathbf{W}|\mathbf{A}, \boldsymbol{\chi})p(\mathbf{C}|\mathbf{X}, \mathbf{W})}[KL(q_{\phi_{\mathbf{X}}}(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})||p_\beta(\mathbf{X}|\mathbf{W}, \mathbf{C}))] \\ & - KL(q_{\phi_{\mathbf{W}}}(\mathbf{W}|\mathbf{A}, \boldsymbol{\chi})||p(\mathbf{W})) - \mathbb{E}_{q(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})q(\mathbf{W}|\mathbf{A}, \boldsymbol{\chi})}[KL(p_\beta(\mathbf{C}|\mathbf{X}, \mathbf{W})||p(\mathbf{C}))] \end{aligned} \quad (17)$$

253 Here i indices have been dropped to simplify the notations and presume one
254 node at each time. In this equation, the first term is the reconstruction term,
255 and the following are X -conditional prior term, W -prior term, and C -prior term
256 [44] in order. The reconstruction term is used to calculate the difference between
257 the input and the reconstruction. This term has the same form as the VGAE's
258 ELBO reconstruction term. It can be estimated using Monte-Carlo samples
259 from $q(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})$. The conditional prior term has a negative sign and a positive
260 value. This term should be small to maximize the objective function. As a
261 result, the conditional prior term serves as a regularizer. To approximate this
262 term, Monte Carlo can be used as shown in Equation (18).

$$\begin{aligned} & \mathbb{E}_{q(\mathbf{W}|\mathbf{A}, \boldsymbol{\chi})p(\mathbf{C}|\mathbf{X}, \mathbf{W})}[KL(q_{\phi_{\mathbf{X}}}(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})||p_\beta(\mathbf{X}|\mathbf{W}, \mathbf{C}))] \approx \\ & \frac{1}{M} \sum_{j=1}^M \sum_{k=1}^K p_\beta(\mathbf{c}_k = 1|\mathbf{x}^j, \mathbf{w}^j) KL(q_{\phi_{\mathbf{X}}}(\mathbf{X}|\mathbf{A}, \boldsymbol{\chi})||p_\beta(\mathbf{X}|\mathbf{w}^j, \mathbf{c}_k = 1)) \end{aligned} \quad (18)$$

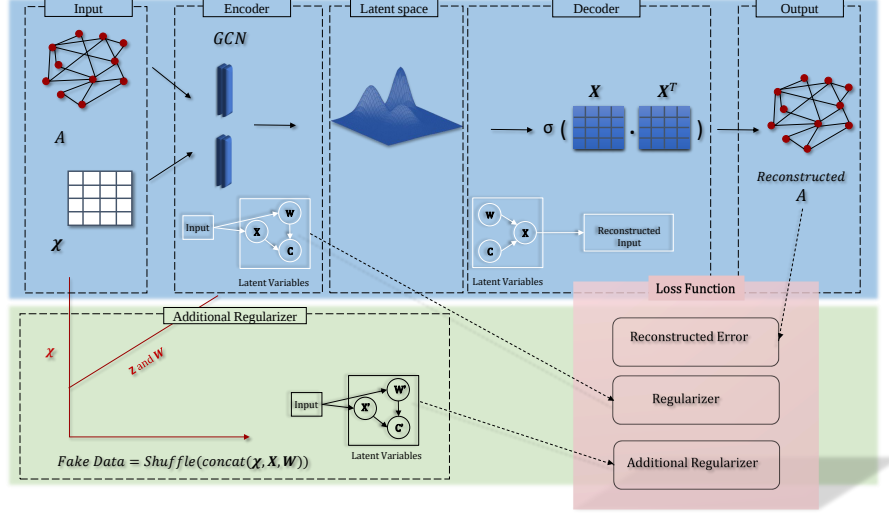


Figure 4: A high-level overview of our method. In the upper half of the figure, the adjacency matrix and the feature matrix are fed into the GCN-based probabilistic encoder. $\mathbf{X}$, $\mathbf{W}$, and $\mathbf{C}$ are the latent variables inferred from the encoder. The decoder reconstructs the adjacency matrix using these latent variables by the inner product of latent vectors. The lower half of the figure is related to the additional regularizer that needs to generate fake data. The cost function consists of three parts, the first part measures the reconstruction error, and the second part is the KL-divergence-based regularizer. The third part is related to the discriminator, in which the fake data is generated from the shuffled combination of the adjacency matrix, $\mathbf{X}$ and $\mathbf{W}$.

263 The $\mathbf{W}$ -prior term in the ELBO equation is also a regularizer and can be com-
 264 puted analytically. $\mathbf{C}$ is a discrete latent variable in our model. $\mathbf{C}$ -posterior
 265 measures how far $\mathbf{X}$ is from each cluster position created by $\mathbf{W}$ to determine
 266 cluster assignment probability. ELBO's $\mathbf{C}$ -prior term tries to reduce the KL di-
 267 vergence between the $\mathbf{C}$ -posterior and the uniform prior. It would aim to bring
 268 the clusters together by maximizing overlap and bringing the means closer to-
 269 gether. This term has a positive value and appears with a negative sign in the
 270 ELBO equation, indicating that it is a regularizer. Figure 4 shows a high-level
 271 overview of our proposed method.

Table 2: Summary of the citation network datasets.

Dataset	Nodes	Edges	Features	Classes
Cora	2,708	5,429	1,433	7
Citeseer	3,327	4,732	3,703	6
PubMed	19,717	44,338	500	3
Cornell	183	295	1703	5
Wisconsin	251	499	1703	5
Texas	183	309	1703	5

272 5.2. Regularizer based on adversarial mechanism concept

273 This paper introduces an additional regularizer inspired by the adversarial
274 mechanism concept. In the adversarial mechanism, a generator provides fake
275 data and tries to deceive a discriminator. For this purpose, the cost function
276 is changed in such a way to involve the competition between the generator and
277 the discriminator. Here, to generate fake data, the feature matrix of input
278 data (χ) and the hidden variables inferred from the input data ($\mathbf{X}$ and $\mathbf{W}$)
279 concatenate together as $concat(\mathbf{X}, \mathbf{W}, \chi)$. Then a shuffling process performs on
280 them as $Shuffle(concat(\mathbf{X}, \mathbf{W}, \chi))$. After that, these fake data are fed to the
281 inference model, and hidden variables of fake data are inferred ($\mathbf{X}', \mathbf{W}', \mathbf{C}'$). In
282 the following, the regularization part of the ELBO is calculated by these fake
283 latent variables, and the negative inverse of that is added to the ELBO. Thus,
284 if KL-divergence cannot detect the difference between the fake and real, the
285 denominator is low, so the fraction is high, and an additional negative burden
286 is added to the ELBO. This part acts as the discriminator. The process is
287 shown in Figure 4. Using the loss function defined for our model and the newly
288 introduced regularizer, the ELBO can be rewritten as L_{RELBO} as shown in
289 Equation (19). To simplify the equation, subscripts have been removed.

$$\begin{aligned}
L_{RELBO} = L_{ELBO} - & \left[\mathbb{E}[KL(q(\mathbf{X}'|\mathbf{A}, \chi')||q(\mathbf{X}|\mathbf{A}, \chi))] \right. \\
& \left. + KL(q(\mathbf{W}'|\mathbf{A}, \chi')||q(\mathbf{W}|\mathbf{A}, \chi)) + \mathbb{E}[KL(p(\mathbf{C}'|\mathbf{X}', \mathbf{W}')||p(\mathbf{C}|\mathbf{X}, \mathbf{W}))] \right]^{(-1)}
\end{aligned}
\tag{19}$$

290 A summary of the main steps of GM-VGAE is as follows:

- 291 • Encoder
 - 292 – Inferring latent variables $(\mathbf{X}, \mathbf{W}, \mathbf{C})$ from real data
- 293 • Regularizer
 - 294 – Concatenating $(\mathbf{X}, \mathbf{W}, \boldsymbol{\chi})$ named *Temp*
 - 295 – Generating Fake Data by Shuffling *Temp*
 - 296 – Inferring latent variables from Fake Data $(\mathbf{X}', \mathbf{W}', \mathbf{C}')$
- 297 • Decoder
 - 298 – Reconstructing the input data based on latent variables inferred from
 - 299 real data
 - 300 – Evaluating the objective function: 1) Reconstruction term, 2) Reg-
 - 301 ularization terms based on $(\mathbf{X}, \mathbf{W}, \mathbf{C})$, and 3) Additional regulariza-
 - 302 tion terms based on $(\mathbf{X}', \mathbf{W}', \mathbf{C}')$

303 6. Experiments

304 In this section, the results of our experiments are given. These experiments
305 are performed to compare the proposed method versus the state-of-the-art al-
306 gorithms. We demonstrate the efficiency of our approach in two downstream
307 machine-learning tasks, clustering and link prediction.

308 6.1. Datasets

309 Experiments on six graph datasets were performed. The first three of them,
310 Cora, Citeseer, and Pubmed, are three well-known citation networks [45, 46].
311 Furthermore, we utilized Cornell, Wisconsin, and Texas (WebKB dataset) datasets
312 following [47]. Table 2 reports a summary of these datasets.

313 **Citation networks.** Cora, Citeseer, and Pubmed fall into this category. The
314 datasets include sparse bag-of-words feature vectors for each document and a

list of document-to-document citation links. In these datasets, nodes correspond to publications and (undirected) edges to citations, i.e., there would be an undirected connection between them whenever a document cites another document. Using the citation links as edges, we create a binary, symmetric adjacency matrix $\mathbf{A}$. A class label was assigned to each document.

WebKB. Cornell, Wisconsin, and Texas are webpage datasets. Carnegie Mellon University collected these datasets from computer science departments of different universities. Nodes denote web pages, and edges represent hyperlinks between these web pages. The datasets include bag-of-words feature vectors for each web page. The web pages are divided into five classes, project, student, staff, course, and faculty.

6.2. Baselines and state-of-the-art methods

We compared the performance of our method with state-of-the-art algorithms: **Spectral Clustering [48]:** This approach seeks to learn social embedding, which generates representations from eigenvectors of the normalized graph Laplacian of G .

DeepWalk [49]: This approach learns latent representations, which these representations encode social relations in a continuous vector space. To do so, Deepwalk learns structural regularities present within short random walks.

SEAL[50]: This approach is a link prediction framework that creates a local subgraph around each target link and uses GNN to learn a function to convert the patterns of the subgraph into link existence.

GAE[28]: This approach is an unsupervised framework for graph structure data, consisting of an encoder and a decoder. This approach leverages both topological and content information.

VGAE [28]: This approach is based on the GAE framework, which consists of a probabilistic encoder and a probabilistic decoder.

ARGA [6]: This approach is based on the GAE. This method enforces the latent representation to match a prior distribution via an adversarial training scheme.

Table 3: Performance comparison of different models for the Link Prediction task

Model	Cora		Citeseer		Pubmed	
	AUC	AP	AUC	AP	AUC	AP
SC	84.6	88.5	80.5	85.0	84.2	87.8
DW	83.1	85.0	80.5	83.6	84.4	84.1
SEAL	91.8	92.9	87.5	88.7	95.4	95.5
GAE	91.0	92.0	89.5	89.9	96.4	96.5
VGAE	91.4	92.6	90.8	92.0	94.4	94.7
ARGA	92.4	93.2	91.9	93.0	96.8	97.1
ARVGA	92.4	92.6	92.4	93.0	96.5	96.8
RWR-GAE	92.9	92.7	92.1	91.5	96.2	96.3
RWR-VGAE	92.6	92.7	92.3	92.4	95.3	95.2
s-VGAE	92.7	93.2	90.3	91.5	97.1	97.1
DGI	89.8	89.7	95.5	95.7	91.2	92.2
GIC	93.5	93.3	97	96.8	93.7	93.5
CGCN	96.38	96.25	96.01	96.25	96.9	95.56
HGCAE-P	95.6	95.5	96.7	97	96.2	96
GM-VGAE	99.0	98.7	99.2	99.1	98.7	98.5

345 **ARVGA** [6]: This approach is VGAE based version of ARGA.

346 **RWR-GAE** [33]: This approach is based on the GAE, which uses a random
347 walk-based method to regularize the representations learned by the encoder.

348 **RWR-VGAE** [33]: This approach is VGAE based version of RWR-GAE.

349 **S-VGAE** [35]: This approach is based on VGAE, which consider the von
350 Mises-Fisher (vMF) distribution as the prior distribution for link prediction
351 purpose.

352 **DGI** [51]: The robust convolutional graph architecture obtains some graph
353 patch representations in this method. The maximization of local mutual infor-
354 mation in these graph patch representations can be used to obtain node embed-
355 dings that account for the graph’s global structural properties. As the number

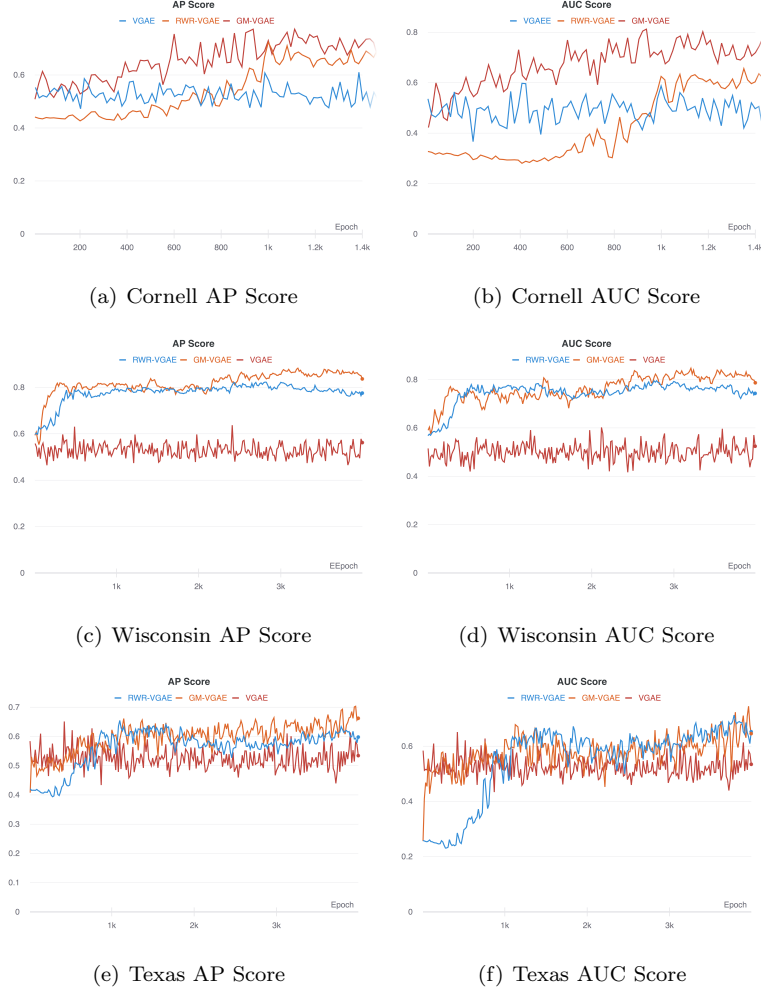


Figure 5: AP and AUC Scores on Cornell, Wisconsin and Texas Datasets

of hidden units has a significant impact on the performance of this method, we employ the technique in two settings: once with 128 hidden units (DGI-128) and once with 512 hidden units (DGI-512)

GIC [52]: This approach is an unsupervised learning technique for representation learning in graphs. GIC recognizes nodes with similar representations, clusters them, and maximizes mutual information based on the DGI technique.

HGCAE-P [53]: This approach is a GAE-based representation learning method

Table 4: Performance comparison of different models for the Clustering task on Cora.

Model	Acc	NMI	F1	Precision	ARI
SC	36.7	12.7	31.8	19.3	3.1
DW	48.4	32.7	39.2	36.1	24.3
GAE	59.6	37.4	59.5	59.6	27.4
VGAE	62.5	37.1	62.5	62.5	31.9
ARGA	66.8	48.9	66.8	66.8	42.2
ARVGA	54.4	43.3	54.4	54.4	31.0
RWR-GAE	59.3	43.1	57.7	57.7	34.1
RWR-VGAE	57.7	43.1	57.7	57.7	37.2
CGCN	71.5	54.4	67.7	71.5	48.2
DGI-128	59.0	38.6	60.6	59.0	33.6
DGI-512	71.3	56.4	67.2	71.3	51.1
GIC	72.5	53.7	69.2	72.5	50.8
HGCAE-P	74.6	59.9	71.73	74.3	52.0
GM-VGAE	74.8	59.7	73.9	74.8	52.2

that derives its representation from a geometry-aware message passing auto-encoder. All operations in auto-encoding are performed in hyperbolic space.

CGCN [37]: CGCN consists of an attributed graph clustering module and a semi-supervised node classification module that enhances the learning ability of semi-supervised node classification using samples with clustering assignments.

6.3. Tasks

In this study, we perform two tasks, clustering and link prediction. In the link prediction task, the edges and non-edges among the test set nodes are predicted after reconstructing the input graph using a decoder. On the other hand, the K-means clustering algorithm is performed on learned node embeddings to specify the clusters.

Table 5: Performance comparison of different models for the Clustering task on Citeseer.

Model	Acc	NMI	F1	Precision	ARI
SC	23.9	5.6	29.9	17.9	3.1
DW	33.7	8.8	27.0	24.8	24.3
GAE	41.2	14.2	41.2	41.2	9.7
VGAE	39.1	13.3	39.1	39.1	7.0
ARGA	50.8	26.9	50.8	50.8	21.4
ARVGA	59.7	33.9	59.7	59.7	33.0
RWR-GAE	44.0	25.2	44.0	44.0	18.0
RWR-VGAE	51.9	28.9	51.9	51.9	25.7
CGCN	67.4	42.3	63.2	66.4	43.5
DGI-128	57.9	30.9	53.4	57.2	27.9
DGI-512	68.8	44.4	64.3	67.8	45
GIC	69.6	45.3	65	69.6	46.5
HGCAE-P	71.5	45.3	67.2	69.5	46.9
GM-VGAE	69.5	43.2	69.1	59.1	45.5

6.4. Metrics

In this paper, in order to evaluate the efficiency of clustering, the following criteria were used by [54]: Average rand index (ARI), normalized mutual information (NMI), precision, F-score (F1), and accuracy (ACC). On the other hand, the average precision (AP) and the area under a receiver operating characteristic curve (AUC) were used following [28] to evaluate the link prediction task.

6.5. Settings

Our model is initialized using Glorot initialization [55] and trained with an initial learning rate of 0.01 for a maximum of 100 epochs using Adam optimizer [56]. We terminate the training if the validation accuracy does not improve for 10 consecutive steps; as a result, most runs finish in less than 200 steps. It

Table 6: Performance comparison of different models for the Clustering task on Pubmed.

Model	Acc	NMI	F1	Precision	ARI
GAE	67.2	22.4	67.2	67.2	24.5
VGAE	67.3	22.5	67.3	67.3	24.5
ARGA	61.8	21.4	61.8	61.8	19.7
ARVGA	41.8	4.5	41.8	41.8	1.9
RWR-GAE	66.4	26.8	65.1	68.1	26.7
RWR-VGAE	67.2	26.4	67.2	67.3	27.3
CGCN	71	30.2	69.7	71	32.4
DGI-128	49	15.1	45	48.2	14.5
DGI-512	53.3	18.1	47.4	53.1	16.6
GIC	67.3	31.8	65.2	67.3	29.1
HGCAE-P	74.8	37.7	73.2	73.2	35.9
GM-VGAE	74.8	37.5	74	74.2	36.9

applied a fixed dropout rate [57] of 0.5 to the input and hidden layers. Also, we considered $\mathcal{L}_2$ regularization of 0.0005 on the weights. We use a two-layer GCN model as an encoder with the parametric ReLU (PReLU) [58] non-linearity. For training the VAE, we use hyper-parameters provided by [28] and apply full-batch gradient descent while using the reparameterization trick [17].

We employ the embedding dimension of 16 for our proposed method. The settings for other methods are as follows. The [59] implementation for SC with an embedding dimension of 128 is used. For DW, we used the implementation provided by [49] with the same settings as in their paper, namely an embedding dimension of 128, 10 random walks of length 80 per node, and a context size of 10, trained for a single epoch. Pytorch implementation of [28] with an embedding dimension of 16 is used for GAE and VGAE. For ARGA and ARVGA, we use the [6] implementation with an embedding dimension of 16.

We use an implementation of [33] with hyper-parameters provided by [28]

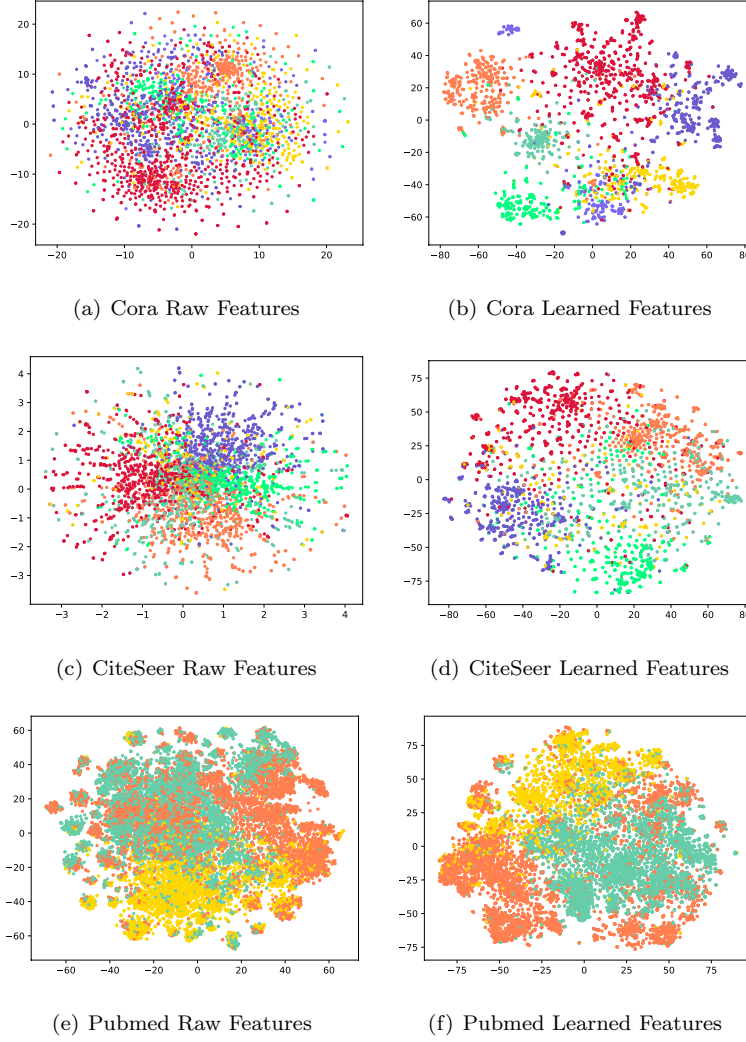


Figure 6: t-SNE embeddings of the nodes in the Cora, CiteSeer, and Pubmed datasets from the raw features (**left**) and features from a learned model (**right**). The clusters of the learned GM-VGAE model’s representations are clearly defined.

401 for the autoencoder in RWR-GAE and RWR-VGAE. For the RandomWalk
 402 Regularization network’s hyper-parameters, the number of walks is set to 50,
 403 and the window size and walk length are set to 30. CGCN embedding dimension
 404 is set to 16 in [37] implementation. DGI implementation is used in two modes,

with embedding sizes of 128 and 512. For GIC, we use the implementation of [60] with an embedding size of 16. Eventually, we use the implementation of [53] for HGCAE with embedding size 16.

6.6. Complexity

In terms of time complexity, when N is the number of nodes, F is the representation size, and d is the dimension of each node’s features, the encoder has a time complexity of $O(N^2dF)$. An encoder with a complexity of $O(NdF)$ is possible if we deem the adjacency matrix sparse. The time complexity of expectation-maximization steps is $O(N)$. We may change this to $O(NK)$ for large sample sizes. Using an inner product decoder, it has an $O(N^2)$ time complexity. Overall, $O(NdF) + O(N) + O(N^2)$ represents the temporal complexity for one epoch.

6.7. Results

For three citation datasets, Cora, Citseer, and Pubmed, Table 3 shows the mean AP and AUC for the link prediction task over 10 runs. Our proposed method achieves best-in-class results across all three datasets for two assessment criteria. The most promising outcomes are bolded. The comparison methods for the link prediction task can be considered as representative of four categories.

Spectral Clustering and DeepWalk belong to the shallow embedding approaches. SEAL, DGI, and GIC are GCN-based deep embedding approaches. GAE, ARGAE, and RWR-GAE are auto-encoder-based methods. VGAE, ARVGA, RWR-VGAE, S-VGAE, HGCAE-p, and CGCN are deep generative-based methods. In a general analysis, it can be seen that deep embedding approaches have improved the results of shallow embedding approaches. Auto-encoder-based approaches have improved the outcomes of deep embedding approaches. DGI and GIC have an implicit auto-encoder structure, which has led to good results. Finally, deep generative-based techniques have improved the results of auto-encoder-based approaches. The proposed method has improved the results of all groups.

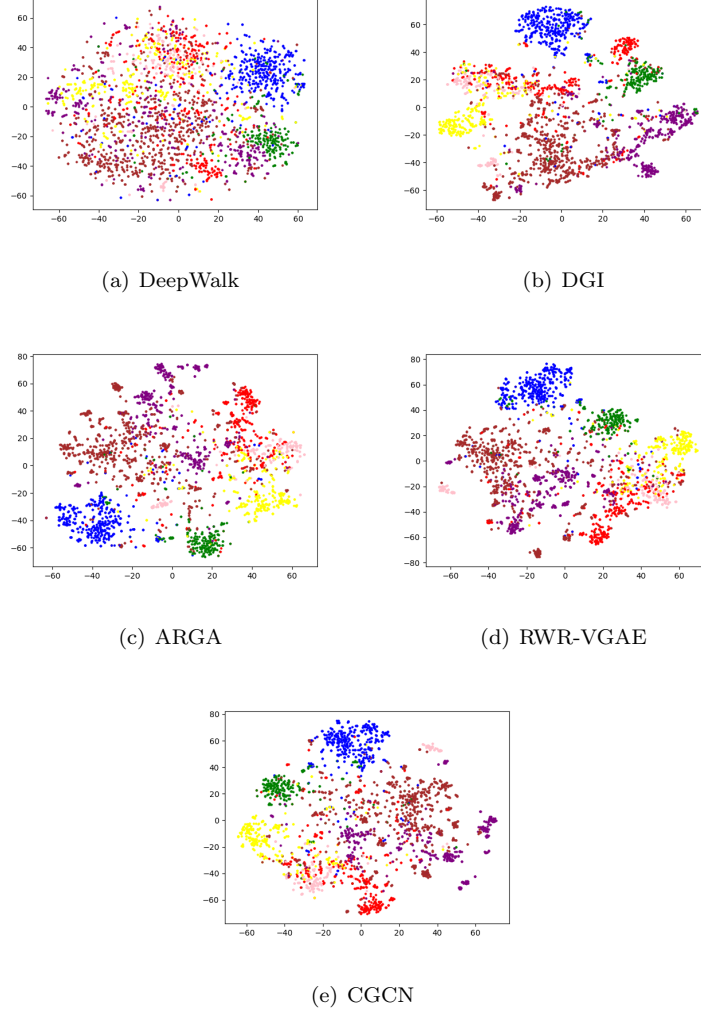


Figure 7: t-SNE embeddings of the nodes in the Cora from the learned features by some of the baseline methods

434 According to Table 3, the comparison between the proposed method and the
 435 best results for each group is as follows: GM-VGAE increases AUC by 14.4%
 436 and AP by 10.2% compared with Spectral Clustering in Cora, AUC by 18.7%,
 437 and AP by 14.1% in Citeseer, AUC by 14.5%, and AP by 10.7% in Pubmed. The
 438 analysis of GM-VGAE and GIC demonstrates 5.5% and 5.4%, 2.2% and 2.3%

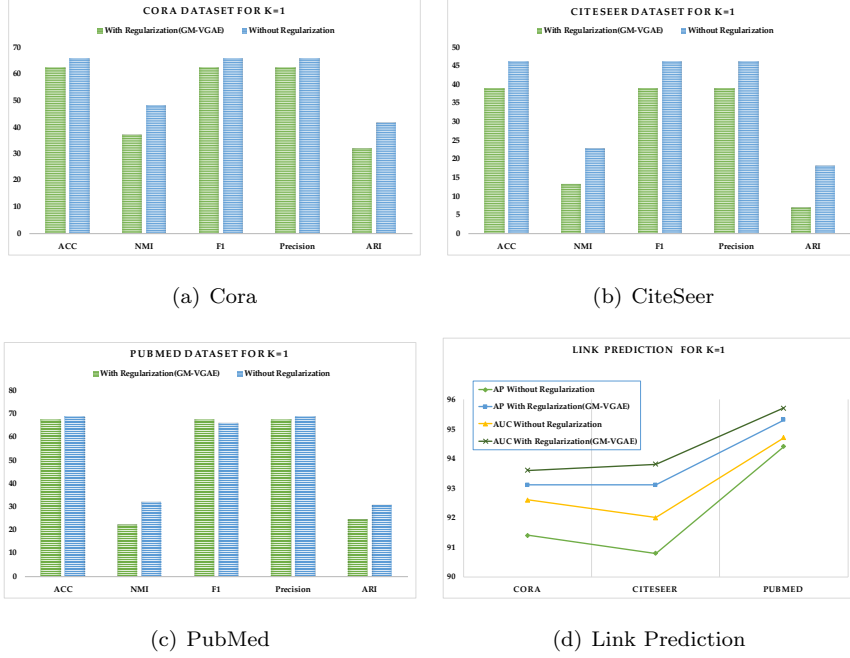


Figure 8: Impact of the regularization on Datasets for $k = 1$

improvement in AUC and AP in the Cora and CiteSeer datasets, respectively, and 5% enhancement in both AUC and AP in Pubmed dataset.

GM-VGAE outperforms AUC by 6.1%, 7.1%, 2.51%, and AP by 6%, 7.6%, 2.2% in Cora, CiteSeer, and Pubmed datasets, respectively, compared with RWR-GAE. The superior result of previous methods belongs to CGCN. Our proposed method evinces 2.62% and 2.4%, 3.19% and 2.85%, and 1.8% and 2.94% progress in AUC and AP in Cora, CiteSeer, and Pubmed, respectively, in contrast to CGCN.

In general, it can be said in the link prediction task that our method showed notable performance on all three benchmarks compared to all state-of-the-art techniques. We repeated the link prediction task on three webpage datasets, Cornell, Wisconsin, and Texas, to further investigate. The results of this experiment are presented in Figure 5. These results show that our method on these datasets also achieved good results. It can also be seen that, by increasing the

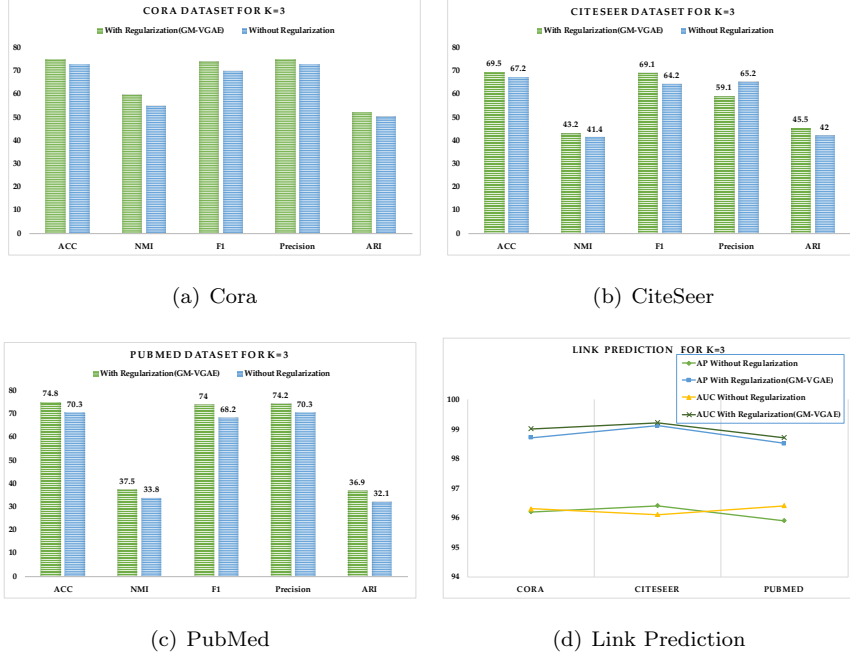


Figure 9: Impact of the regularization on Datasets for $k = 3$

number of runs, the model showed more improvement and better performance compared to other methods.

Tables 4 to 6 show how our model works on three different citation datasets when it comes to node clustering. After 10 runs of each experiment, we announce the mean of the clustering metrics.

Our proposed method shows notable performance on the clustering problem regarding the clustering results on the Cora dataset. The GM-VGAE approach increases accuracy by 38.1%, NMI by 47%, F1 by 42.1%, precision by 55.5 %, and ARI by 49.4% contrary to Spectral Clustering. The maximum results of clustering among the comparative methods on the Cora dataset are provided by HGCAE-P. Our proposed method augment accuracy by 2.2%, F1 by 2.17%, precision by 0.5%, and ARI by 0.2% against HGCAE-P.

In general, the results of our proposed method on Cora and PubMed showed a significant superiority over other methods in the clustering task. Moreover,

our method also represented remarkable performance on CiteSeer in this task over all of the methods but HGCAE-P; however, our results slightly vary with this model.

6.8. Qualitative Analysis

t-SNE is a variant of Stochastic Neighbour Embedding that preserves the data’s local structure while revealing some important global structures. It visualizes high-dimensional data in a two- or three-dimensional space. Figure 6 shows a variety of regular t-SNE plots [61] of the representations learned by the GM-VGAE algorithm on the Cora, CiteSeer, and Pubmed datasets to help understand the efficacy of GM-VGAE. Colours indicate the document’s type. Compared to the raw features, the trained representations in 2D space show noticeable clustering proportional to the number of topic classes in each dataset. For better comparison, Figure 7 shows the plots of some of the baseline methods in the Cora dataset. We plotted the representations learned by DW from shallow embedding approaches, DGI from GCN-based approaches, ARGAs and RWR-GAE from auto-encoder-based approaches, and CGCN from deep generative-based approaches.

6.9. Ablation Study

Tables 7 and 8 show the results of comparing the values 1, 2, and 3 for the hyperparameter k in clustering and link prediction tasks. As can be noticed, the best value of k for all datasets is 3. However, the best value of k can be varied depending on the new diverse datasets. The first column of the tables shows the affected results of the GMM method. As demonstrated, at $k = 2$ with applying the GMM, remarkable progress in the results is achieved which regarding clustering results on the Cora dataset (Table 7), GM-VGAE with $k = 2$ increases accuracy by 7.78%, NMI by 7.03%, F1 by 8.18%, precision by 7.785 %, and ARI by 10.68% against GM-VGAE with $k = 1$. This boost indicates the validity of our claim that using GMM positively affects results.

Table 7: Effect of parameter k on the result of clustering.

Dataset	Metric	K = 1	K = 2	K = 3
Cora	ACC	65.72	73.5	74.8
	NMI	48.27	55.3	59.7
	F1	65.72	73.9	73.9
	Precision	65.72	73.5	74.8
	ARI	41.82	52.2	52.2
Citeseer	ACC	46.20	59.5	69.5
	NMI	22.93	33.2	43.2
	F1	46.20	59.1	69.1
	Precision	46.20	59.1	59.1
	ARI	18.26	32.5	45.5
Pubmed	ACC	68.83	72.6	74.8
	NMI	32.0	34.5	37.5
	F1	65.72	72.2	74
	Precision	68.83	73.2	74.2
	ARI	30.91	35.9	36.9

495 To evaluate the effectiveness of adversarial regularization, the proposed method
 496 is investigated in two modes: with and without regularization. Figure 8 shows
 497 the effectiveness of the adversarial regularization without considering GMM to
 498 further emphasize the impact of adversarial regularization. The analysis illus-
 499 trates the definite effect of regularization on the results of both the clustering
 500 and link prediction tasks on all datasets. Figure 9 highlights the potency of
 501 the adversarial regularization for $k = 3$ to verify the effect of regularization on
 502 improving the results of both clustering and link prediction tasks.

Table 8: Effect of parameter k on the result of Link prediction.

Dataset	Metric	K = 1	K = 2	K = 3
Cora	AUC	93.61	98.5	99
	AP	93.17	98.22	98.7
Citeseer	AUC	93.82	98.75	99.2
	AP	93.10	98.39	99.1
Pubmed	AUC	95.38	98.2	98.7
	AP	95.72	98.0	98.5

7. Conclusion

In this paper, we proposed a variant of the VGAE. In this proposed framework, we assume a GMM to model the prior distribution instead of adopting a Gaussian distribution in a regular VGAE. This assumption intends to perform a specific task. Here, this particular task is unsupervised clustering. Moreover, inspired by the adversarial mechanism concept, a regularization was introduced. Eventually, we examined the performance of the proposed method on six open graph datasets, Cora, Citseer, and PubMed, which are well-known citation networks, and Cornell, Wisconsin, and Texas, which are webpage datasets. To examine the performance, we performed clustering and link prediction tasks on citation networks and link prediction tasks on webpage datasets. Our proposed method demonstrated remarkable performance on these tasks compared to state-of-the-art algorithms.

516 8. Code Availability

517 A reference to GM-VGAE implementation can be found at:
518 <https://github.com/SoheilaMolaei/GM-VGAE>.

519 9. Acknowledgements

520 This research was supported by the Wellcome Trust under grant 217650/Z/19/Z.
521 DAC was supported by the NIHR Oxford Biomedical Research Centre, the In-
522 noHK Hong Kong Centre for Cerebro-cardiovascular Health Engineering (COCHE),
523 and the Pandemic Sciences Institute at the University of Oxford.

524 References

- 525 [1] R. Angles, C. Gutierrez, Survey of graph database models, ACM Comput-
526 ing Surveys (CSUR) 40 (1) (2008) 1–39.
- 527 [2] A. Fout, J. Byrd, B. Shariat, A. Ben-Hur, Protein interface prediction using
528 graph convolutional networks, in: Proceedings of the 31st International
529 Conference on Neural Information Processing Systems, 2017, pp. 6533–
530 6542.
- 531 [3] A. Sanchez-Gonzalez, N. Heess, J. T. Springenberg, J. Merel, M. Riedmiller,
532 R. Hadsell, P. Battaglia, Graph networks as learnable physics engines for
533 inference and control, in: International Conference on Machine Learning,
534 PMLR, 2018, pp. 4470–4479.
- 535 [4] H. Hamedmoghadam, M. Jalili, H. L. Vu, L. Stone, Percolation of hetero-
536 geneous flows uncovers the bottlenecks of infrastructure networks, Nature
537 communications 12 (1) (2021) 1–10.
- 538 [5] S. Molaei, H. Zare, H. Veisi, Deep learning approach on information dif-
539 fusion in heterogeneous networks, Knowledge-Based Systems 189 (2020)
540 105–153.

- [6] S. Pan, R. Hu, G. Long, J. Jiang, L. Yao, C. Zhang, Adversarially regularized graph autoencoder for graph embedding, in: Proceedings of the 27th International Joint Conference on Artificial Intelligence, 2018, pp. 2609–2615.
- [7] F. Hu, Y. Zhu, S. Wu, W. Huang, L. Wang, T. Tan, Graphair: Graph representation learning with neighborhood aggregation and interaction, Pattern Recognition 112 (2021) 107745.
- [8] Z. Hou, Y. Cen, Y. Dong, J. Zhang, J. Tang, Automated unsupervised graph representation learning, IEEE Transactions on Knowledge & Data Engineering (01) (2021) 1–13.
- [9] S. Molaei, N. G. Bousejin, H. Zare, M. Jalili, Deep node clustering based on mutual information maximization, Neurocomputing 455 (2021) 274–282.
- [10] S. Molaei, N. G. Bousejin, H. Zare, M. Jalili, S. Pan, Learning graph representations with maximal cliques, IEEE Transactions on Neural Networks and Learning Systems (2021) 1–8.
- [11] J. Zhou, G. Cui, S. Hu, Z. Zhang, C. Yang, Z. Liu, L. Wang, C. Li, M. Sun, Graph neural networks: A review of methods and applications, AI Open 1 (2020) 57–81.
- [12] W. L. Hamilton, Graph representation learning, Synthesis Lectures on Artificial Intelligence and Machine Learning 14 (3) (2020) 1–159.
- [13] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, S. Y. Philip, A comprehensive survey on graph neural networks, IEEE transactions on neural networks and learning systems 32 (1) (2020) 4–24.
- [14] W. L. Hamilton, R. Ying, J. Leskovec, Representation learning on graphs: Methods and applications, IEEE Data(base) (2017) 1–24.
- [15] D. Charte, F. Charte, M. J. del Jesus, F. Herrera, An analysis on the use of autoencoders for representation learning: Fundamentals, learning task case studies, explainability and challenges, Neurocomputing 404 (2020) 93–107.

- [16] T. Kipf, E. Fetaya, K.-C. Wang, M. Welling, R. Zemel, Neural relational inference for interacting systems, in: International Conference on Machine Learning, PMLR, 2018, pp. 2688–2697.
- [17] D. P. Kingma, M. Welling, Auto-encoding variational bayes, ICLR (2014) 14–16.
- [18] I. Higgins, L. Matthey, X. Glorot, A. Pal, B. Uria, C. Blundell, S. Mohamed, A. Lerchner, Early visual concept learning with unsupervised deep learning., CoRR abs/1606.05579 (2016) 1–12.
- [19] T. N. Kipf, M. Welling, Semi-supervised classification with graph convolutional networks, ICLR (2017) 1–14.
- [20] J. Bruna, W. Zaremba, A. Szlam, Y. LeCun, Spectral networks and deep locally connected networks on graphs, in: 2nd International Conference on Learning Representations, ICLR 2014, 2014, pp. 1–14.
- [21] P. V. Tran, Multi-task graph autoencoders, NIPS (2018) 1–5.
- [22] G. Salha, R. Hennequin, J.-B. Remy, M. Moussallam, M. Vazirgiannis, Fastgae: Scalable graph autoencoders with stochastic subgraph decoding, Neural Networks 142 (2021) 1–19.
- [23] J. Park, M. Lee, H. J. Chang, K. Lee, J. Y. Choi, Symmetric graph convolutional autoencoder for unsupervised graph representation learning, in: 2019 IEEE/CVF International Conference on Computer Vision (ICCV), IEEE Computer Society, 2019, pp. 6518–6527.
- [24] G. Salha, S. Limnios, R. Hennequin, V.-A. Tran, M. Vazirgiannis, Gravity-inspired graph autoencoders for directed link prediction, in: Proceedings of the 28th ACM International Conference on Information and Knowledge Management, 2019, pp. 589–598.
- [25] A. Salehi, H. Davulcu, Graph attention auto-encoders, in: 2020 IEEE 32nd International Conference on Tools with Artificial Intelligence (IC-TAI), IEEE, 2020, pp. 989–996.

- [26] S. Fan, X. Wang, C. Shi, E. Lu, K. Lin, B. Wang, One2multi graph autoencoder for multi-view graph clustering, in: Proceedings of The Web Conference 2020, 2020, pp. 3070–3076.
- [27] D. J. Rezende, S. Mohamed, D. Wierstra, Stochastic backpropagation and approximate inference in deep generative models, in: International conference on machine learning, PMLR, 2014, pp. 1278–1286.
- [28] T. N. Kipf, M. Welling, Variational graph auto-encoders, NIPS Workshop on Bayesian Deep Learning (2016) 1–12.
- [29] Q. Xie, J. Huang, P. Du, M. Peng, Graph relational topic model with higher-order graph attention auto-encoders, in: Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021, 2021, pp. 2604–2613.
- [30] T. S. Hy, R. Kondor, Multiresolution graph variational autoencoder, arXiv preprint arXiv:2106.00967 (2021) 1–15.
- [31] D. Jin, B. Li, P. Jiao, D. He, W. Zhang, Network-specific variational auto-encoder for embedding in attribute networks, in: Proceedings of the 28th International Joint Conference on Artificial Intelligence, 2019, pp. 2663–2669.
- [32] A. Makhzani, J. Shlens, N. Jaitly, I. Goodfellow, B. Frey, Adversarial autoencoders, ICLR (2016) 1–16.
- [33] P.-Y. Huang, R. Frederking, et al., Rwr-gae: Random walk regularization for graph auto encoders, arXiv preprint arXiv:1908.04003 (2019) 1–7.
- [34] Q. Lu, N. de Silva, D. Dou, T. H. Nguyen, P. Sen, B. Reinwald, Y. Li, Exploiting node content for multiview graph convolutional network and adversarial regularization, in: Proceedings of the 28th International Conference on Computational Linguistics, 2020, pp. 545–555.
- [35] T. R. Davidson, L. Falorsi, N. De Cao, T. Kipf, J. M. Tomczak, Hyper-spherical variational auto-encoders, in: 34th Conference on Uncertainty

- in Artificial Intelligence 2018, UAI 2018, Association For Uncertainty in
Artificial Intelligence (AUAI), 2018, pp. 856–865.
- [36] S. Zheng, Z. Zhu, X. Zhang, Z. Liu, J. Cheng, Y. Zhao, Distribution-
induced bidirectional generative adversarial network for graph representa-
tion learning, in: Proceedings of the IEEE/CVF Conference on Computer
Vision and Pattern Recognition, 2020, pp. 7224–7233.
- [37] B. Hui, P. Zhu, Q. Hu, Collaborative graph convolutional networks: Un-
supervised learning meets semi-supervised learning, in: Proceedings of the
AAAI Conference on Artificial Intelligence, Vol. 34, 2020, pp. 4215–4222.
- [38] M. Defferrard, X. Bresson, P. Vandergheynst, Convolutional neural net-
works on graphs with fast localized spectral filtering, in: Proceedings of
the 30th International Conference on Neural Information Processing Sys-
tems, 2016, pp. 3844–3852.
- [39] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, Y. Bengio,
Graph attention networks, in: International Conference on Learning Rep-
resentations, 2018, pp. 1–18.
- [40] W. L. Hamilton, R. Ying, J. Leskovec, Inductive representation learning
on large graphs, in: Proceedings of the 31st International Conference on
Neural Information Processing Systems, 2017, pp. 1025–1035.
- [41] J. Duchi, E. Hazan, Y. Singer, Adaptive subgradient methods for online
learning and stochastic optimization., Journal of machine learning research
12 (7) (2011) 2121–2159.
- [42] J. Duchi, Y. Singer, Efficient online and batch learning using forward back-
ward splitting, The Journal of Machine Learning Research 10 (2009) 2899–
2934.
- [43] D. A. Reynolds, Gaussian mixture models, Encyclopedia of biometrics 741
(2009) 659–663.

- [44] N. Dilokthanakul, P. A. Mediano, M. Garnelo, M. C. Lee, H. Salimbeni, K. Arulkumaran, M. Shanahan, Deep unsupervised clustering with gaussian mixture variational autoencoders, ICLR (2017) 1–12.
- [45] P. Sen, G. Namata, M. Bilgic, L. Getoor, B. Galligher, T. Eliassi-Rad, Collective classification in network data, AI magazine 29 (3) (2008) 93–93.
- [46] G. Namata, B. London, L. Getoor, B. Huang, U. EDU, Query-driven active surveying for collective classification, in: 10th International Workshop on Mining and Learning with Graphs, Vol. 8, 2012, pp. 1–8.
- [47] H. Pei, B. Wei, K. C.-C. Chang, Y. Lei, B. Yang, Geom-gcn: Geometric graph convolutional networks, ICLR (2020) 1–10.
- [48] L. Tang, H. Liu, Leveraging social media networks for classification, Data Mining and Knowledge Discovery 23 (3) (2011) 447–478.
- [49] B. Perozzi, R. Al-Rfou, S. Skiena, Deepwalk: Online learning of social representations, in: Proceedings of the 20th ACM SIGKDD international conference on Knowledge discovery and data mining, 2014, pp. 701–710.
- [50] M. Zhang, Y. Chen, Link prediction based on graph neural networks, Advances in Neural Information Processing Systems 31 (2018) 5165–5175.
- [51] P. Velickovic, W. Fedus, W. L. Hamilton, P. Liò, Y. Bengio, R. D. Hjelm, Deep graph infomax, ICLR (Poster) 2 (3) (2019) 1–17.
- [52] C. Mavromatis, G. Karypis, Graph infoclust: Maximizing coarse-grain mutual information in graphs, in: Pacific-Asia Conference on Knowledge Discovery and Data Mining, Springer, 2021, pp. 541–553.
- [53] J. Park, J. Cho, H. J. Chang, J. Y. Choi, Unsupervised hyperbolic representation learning via message passing auto-encoders, in: 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), IEEE, 2021, pp. 5512–5522.

- [54] R. Xia, Y. Pan, L. Du, J. Yin, Robust multi-view spectral clustering via low-rank and sparse decomposition, in: Proceedings of the Twenty-Eighth AAAI Conference on Artificial Intelligence, 2014, pp. 2149–2155.
- [55] X. Glorot, Y. Bengio, Understanding the difficulty of training deep feed-forward neural networks, in: Proceedings of the thirteenth international conference on artificial intelligence and statistics, JMLR Workshop and Conference Proceedings, 2010, pp. 249–256.
- [56] D. P. Kingma, J. Ba, Adam: A method for stochastic optimization, in: ICLR (Poster), 2015, pp. 1–15.
- [57] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, R. Salakhutdinov, Dropout: a simple way to prevent neural networks from overfitting, The Journal of Machine Learning Research 15 (1) (2014) 1929–1958.
- [58] K. He, X. Zhang, S. Ren, J. Sun, Delving deep into rectifiers: Surpassing human-level performance on imagenet classification, in: Proceedings of the IEEE international conference on computer vision, 2015, pp. 1026–1034.
- [59] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, et al., Scikit-learn: Machine learning in python, the Journal of machine Learning research 12 (2011) 2825–2830.
- [60] C. Mavromatis, G. Karypis, Graph infoclust: Maximizing coarse-grain mutual information in graphs, in: PAKDD, 2021.
- [61] L. v. d. Maaten, G. Hinton, Visualizing data using t-sne, Journal of machine learning research 9 (Nov) (2008) 2579–2605.