



Contents lists available at ScienceDirect

Information Processing and Management

journal homepage: www.elsevier.com/locate/ipm

Rank-based self-training for graph convolutional networks

Daniel Carlos Guimarães Pedronette^{a,*}, Longin Jan Latecki^b^a Department of Statistics, Applied Mathematics and Computing (DEMAC), São Paulo State University (UNESP), Rio Claro, Brazil^b Department of Computer and Information Sciences, Temple University, Philadelphia, USA

ARTICLE INFO

Keywords:

Graph convolutional networks

Self-training

Rank model

Semi-supervised learning

ABSTRACT

Graph Convolutional Networks (GCNs) have been established as a fundamental approach for representation learning on graphs, based on convolution operations on non-Euclidean domain, defined by graph-structured data. GCNs and variants have achieved state-of-the-art results on classification tasks, especially in semi-supervised learning scenarios. A central challenge in semi-supervised classification consists in how to exploit the maximum of useful information encoded in the unlabeled data. In this paper, we address this issue through a novel self-training approach for improving the accuracy of GCNs on semi-supervised classification tasks. A margin score is used through a rank-based model to identify the most confident sample predictions. Such predictions are exploited as an expanded labeled set in a second-stage training step. Our model is suitable for different GCN models. Moreover, we also propose a rank aggregation of labeled sets obtained by different GCN models. The experimental evaluation considers four GCN variations and traditional benchmarks extensively used in the literature. Significant accuracy gains were achieved for all evaluated models, reaching results comparable or superior to the state-of-the-art. The best results were achieved for rank aggregation self-training on combinations of the four GCN models.

1. Introduction

Mainly grounded by deep-learning approaches, classification methods experimented a remarkable development in last decade. However, in spite of the huge advances achieved, performing classification tasks based on scarce labeled data still remains a challenging task, since deep models often require large amount of data for training. In this scenario, semi-supervised classification re-emerge as a promising approach, capable of also exploiting useful information encoded in the unlabeled data. Actually, semi-supervised learning (SSL) is halfway between supervised and unsupervised, being suitable to operate with large amounts of unlabeled data and a small quantity of labeled data (Chapelle, Schölkopf, & Zien, 2010; Triguero, García, & Herrera, 2015). A direct and central motivation for semi-supervised learning relies on the fact that labeled data is typically much harder to obtain compared to unlabeled data (Tian, Yu, Xue, & Sebe, 2004).

One of the earliest approaches of semi-supervised learning is self-training, also known as self-labeling or self-supervision (Chapelle et al., 2010; Scudder, 1965). The main idea consists in a wrapper approach that repeatedly exploits a supervised learning method. Firstly, the supervised method is trained based on labeled data only. Subsequently, the unlabeled data is labeled according to the trained supervised model. Then, a selected sub-set of predictions is exploited as additional labeled data to re-train the supervised method (Chapelle et al., 2010). Naturally, the possible combinations between supervised models and approaches to select the additional labeled data open a broad and promising range of self-training approaches (Triguero et al., 2015).

* Corresponding author.

E-mail address: daniel.pedronette@unesp.br (D.C.G. Pedronette).

On the other hand, graphs have been employed as a representation tool in a wide range of real-world scenarios, mainly due to their expressive power (Sun, Lin, & Zhu, 2020). Graph analytics approaches enable a better understanding of what is behind the data, being useful in diverse applications and domains. Tasks as node classification, node recommendation, and link prediction often can benefit from graph-based representations in several areas, from social networks to physical and biological systems (Cai, Zheng, & Chang, 2018). Additionally, the semi-supervised learning literature also exploits graphs as a way to encode the geometry of both labeled and unlabeled data in order to improve supervised methods (Chapelle et al., 2010).

For graph-based semi-supervised learning, a remarkable recent development has been achieved by Graph Neural Networks (GNNs) (Wu, Pan, Chen, Long, Zhang, & Yu, 2019) and, more specifically, by Graph Convolutional Networks (GCNs) (Kipf & Welling, 2017). The GCN models integrate local node features and graph topology in the convolutional layers (Li, Han, & Wu, 2018). In fact, graph convolution enables the extension of standard convolution from Euclidean to non-Euclidean domain, defined by graph-structured data (Sun et al., 2020). Recently, GCNs (Kipf & Welling, 2017) and subsequent variants (Bianchi, Grattarola, Livi, & Alippi, 2019; Klicpera, Bojchevski, & Günnemann, 2019; Pei, Wei, Chang, Lei, & Yang, 2020; Velickovic, et al., 2018; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) have achieved state-of-the-art results in diverse applications, specially involving semi-supervised classification tasks. Many variants propose changes in the network structure, mostly by including or removing components. For instance, in Wu, Souza, Zhang, Fifty, Yu, and Weinberger (2019), the weight matrices between consecutive layers are collapsed and non-linearities removed. In Bianchi et al. (2019), a convolutional layer based on auto-regressive moving average filter is inserted.

In a distinct and promising research direction, other recent works (Bai, Zhang, & Torr, 2019; Klicpera, Weißenberger, & Günnemann, 2019; Li et al., 2018; Sun et al., 2020) have exploited traditional machine learning strategies for further improving GCN capacities. Mainly motivated by the recent impressive GCN results, some relevant techniques in machine learning have been revisited in conjunction with GCN models (Li et al., 2018). In Bai et al. (2019), hypergraphs are exploited to learn deep embeddings on the high-order graph-structured data. A diffusion process is employed in Klicpera, Weißenberger, and Günnemann (2019) for aggregating information from a larger neighborhood. The neighborhood is constructed based on the sparsification of a generalized form of graph diffusion.

In this context, some few self-training approaches have been investigated for GCNs very recently (Li et al., 2018; Sun et al., 2020; Zhou, Zhang, & Huang, 2019). The GCN model is described as a special form of Laplacian smoothing in Li et al. (2018), discussing scenarios of success and fails of the model. While the smoothing is pointed as the key reason for successful scenarios, potential over-smoothing can occur for many convolutional layers. A self-training approach is proposed to overcome such limitations. In a more recent work (Sun et al., 2020), the method is extended to multiple self-training stages. An aligning mechanism is used on the output of a deep-clustering approach. The clusters are used to assign pseudo-labels for each unlabeled data point in the embedding space.

In this paper, we propose a novel Rank-Based Self-Training approach for Graph Convolutional Networks. Our work focuses on the selection of data points in the embedding space to be used as additional labeled data. How to accurately select effective predictions is a challenging task for self-training approaches, few exploited in recent related work (Li et al., 2018; Sun et al., 2020). Typically, the selection is performed by considering the top softmax scores per class (Li et al., 2018; Sun et al., 2020). We propose to employ a margin-based selection score inspired by active learning approaches (Settles, 2009). The nodes are represented according to the score in a rank-based model for the whole dataset, allowing to handle unbalanced additional labeled data per class. Based on the rank model, an expanded labeled set is defined for a second-stage training and classification. The main contributions, differences and novelties with respect to related works can be summarized as follows:

- The proposed self-training approach is not restricted to a specific GCN model. Different from recent related self-training methods (Li et al., 2018; Sun et al., 2020; Zhou et al., 2019), which are focused on a GCN model, the proposed method was validated on four GCN models (Bianchi et al., 2019; Kipf & Welling, 2017; Klicpera, Bojchevski, & Günnemann, 2019; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) with a log-soft-max as the last layer. Actually, the rank-based formulation is versatile and can be easily extended to other learning tasks;
- The proposed margin score addresses the challenging task of self-labeled data selection, which has a central problem in self-training methods. Most of approaches consider the *softmax* score of the predicted class to estimate the confidence of prediction. In contrast, the proposed margin confidence considers not only the class predicted, but the relationship between the first and second higher scores. The motivation is based on the conjecture that even a high *softmax* score can provide a low accurate estimation when first and second top prediction scores are similar. The proposed score provides a simple, yet effective label prediction estimation;
- Our approach employ a global ranking of nodes, being able to handle data with unbalanced classes. In contrast, other self-training methods (Li et al., 2018; Sun et al., 2020; Zhou et al., 2019) often select a fixed number of nodes per class as the expanded labeled set. However, a forced equal amount per class can add low-accurate predictions to certain classes;
- A rank aggregation approach is proposed to exploit and fuse the information from distinct GCN models and training executions. Based on the fused confidence which consider different models, an aggregated and more effective ranking of nodes is computed in order to obtain a more accurate expanded labeled set. To the best of our knowledge, it is the first approach which fuses information from distinct GCN models in a self-training setting.

The proposed approach was evaluated for semi-supervised classification tasks on citation datasets broadly used as benchmark in the literature (Bai et al., 2019; Bianchi et al., 2019; Kipf & Welling, 2017; Klicpera, Bojchevski, & Günnemann, 2019; Li et al., 2018; Sun et al., 2020; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019). The experiments indicate significant accuracy gains of the

proposed self-training method applied to four different GCN models. Accuracy results on semi-supervised classification outperform most of state-of-art methods.

The remainder of this paper is organized as follows. Section 2 discusses related work and a formal definition of the problem setting. Section 3 presents the proposed self-training approach. Section 4 describes the conducted experimental evaluation and, finally, Section 5 discusses the conclusions.

2. Problem definition and preliminaries

2.1. Graph-based semi-supervised learning

In this section, we first discuss a formal definition of the semi-supervised learning classification task using graph convolution networks, mostly following (Kipf & Welling, 2017; Li et al., 2018).

Let $\mathcal{G}$ denotes an undirected graph represented by $\mathcal{G} = (\mathcal{V}, \mathcal{E}, \mathbf{X})$, where $\mathcal{V}$ is the node set, $\mathcal{E}$ is the edge set and $\mathbf{X}$ is a feature matrix. The node set is given by $\mathcal{V} = \{v_1, v_2, \dots, v_n\}$ and the edge set is defined by a set of pairs $(v_i, v_j) \in \mathcal{E}$, which can be represented by a non-negative adjacency matrix $\mathbf{A} = [a_{ij}] \in \mathbb{R}^{n \times n}$. The feature matrix is defined as $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n]^T \in \mathbb{R}^{n \times d}$, where $\mathbf{x}_i$ is a d -dimensional feature vector which represents the node v_i .

Let $\mathcal{Y} = \{y_1, y_2, \dots, y_c\}$ be a set of labels which can be assigned to nodes $v_i \in \mathcal{V}$. In this way, the node set can be more specifically defined as $\mathcal{V} = \{v_1, v_2, \dots, v_L, v_{L+1}, \dots, v_n\}$, which denotes a partially labeled dataset, where $\mathcal{V}_L = \{v_i\}_{i=1}^L$ is the labeled data items subset and $\mathcal{V}_U = \{x_i\}_{i=L+1}^n$ is the unlabeled data items subset. For semi-supervised classification, as a general rule, we have $|\mathcal{V}_L| \ll |\mathcal{V}_U|$. Formally, the training set can be seen as a labeling function $l : \mathcal{V}_L \rightarrow \mathcal{Y}$, where $y_i = l(v_i) \forall v_i \in \mathcal{V}_L$. The goal is to learn a function $\hat{l} : \mathcal{V}_U \rightarrow \mathcal{Y}$ to predict the labels of unlabeled nodes in $\mathcal{V}_U$.

2.2. Graph convolutional networks

Recently, much effort has been made on exploiting deep learning approaches for graph data (Cai et al., 2018). In this context, Graph Convolutional Networks (GCN) represent a relevant graph-based neural network model, introduced in Kipf and Welling (2017). In a simplified way, GCN learns the embedding (representation) of each node by iteratively aggregating the embeddings of its neighbors, encoding the graph structure directly on a neural network model. A two-layer GCN model is used for semi-supervised node classification in Kipf and Welling (2017), taking into account a graph represented by a symmetric adjacency matrix $\mathbf{A}$.

The network model can be depicted as a function both on the feature data $\mathbf{X}$ and on the adjacency matrix $\mathbf{A}$, as:

$$\mathbf{Z} = f(\mathbf{X}, \mathbf{A}), \quad (1)$$

where $\mathbf{Z}$ denotes an embedding matrix, such that $\mathbf{Z} = [\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n]^T \in \mathbb{R}^{n \times c}$ and $\mathbf{z}_i$ is a c -dimensional embedded representation learned for the node v_i .

The degree matrices are computed as a pre-processing step, defined as $\hat{\mathbf{A}} = \tilde{\mathbf{D}}^{-1/2} \tilde{\mathbf{A}} \tilde{\mathbf{D}}^{-1/2}$, where $\tilde{\mathbf{A}} = \mathbf{A} + \mathbf{I}$ and $\tilde{\mathbf{D}}$ is the degree matrix of $\tilde{\mathbf{A}}$. Then, the function $f(\cdot)$ which represents the two-layer GCN model assumes the form:

$$\mathbf{Z} = \log(\text{softmax}(\hat{\mathbf{A}} \text{ReLU}(\hat{\mathbf{A}} \mathbf{X} \mathbf{W}^{(0)})) \mathbf{W}^{(1)})) \quad (2)$$

The matrix $\mathbf{W}^{(0)} \in \mathbb{R}^{d \times H}$ defines the neural network weights for an input-to-hidden layer with H feature maps, while $\mathbf{W}^{(1)} \in \mathbb{R}^{H \times c}$ is a hidden-to-output matrix. Both matrices $\mathbf{W}^{(0)}$ and $\mathbf{W}^{(1)}$ are trained using gradient descent, considering the cross-entropy error over all labeled nodes $v_i \in \mathcal{V}_L$. The softmax activation function is applied row-wise and yields the probability distribution over the c class labels for each row, i.e., the probability values sum up to 1 for each row. After log function, the label assigned to a node v_i is defined according to the class with the less negative value in the embedded representation $\mathbf{z}_i$.

Mostly grounded by the success of the GCN (Kipf & Welling, 2017), various related graph convolutional network models have been recently proposed (Bai et al., 2019; Bianchi et al., 2019; Klicpera, Bojchevski, & Günnemann, 2019; Klicpera, Weissenberger, & Günnemann, 2019; Li et al., 2018; Velickovic, et al., 2018; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019). While some approaches focus in the structure of network models (Bai et al., 2019; Bianchi et al., 2019; Klicpera, Bojchevski, & Günnemann, 2019; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019), others present contributions involving training steps and manifold information (Klicpera, Weissenberger, & Günnemann, 2019; Li et al., 2018). Different network models (Bianchi et al., 2019; Kipf & Welling, 2017; Klicpera, Bojchevski, & Günnemann, 2019; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) can be used in conjunction with the proposed self-training approach. The only condition is that a log-soft-max operation is kept as the last layer.

3. Rank-based self-training

We propose a self-training approach focused on better exploiting the unlabeled data by taking into account the information encoded in the embedding computed by a first stage semi-supervised classification. The method post-processes the predicted labels computed through a GCN classification, by identifying high-confidence predictions to be used for pseudo-labeled data expansion.

A challenging task for different self-training methods (Li et al., 2018; Sun et al., 2020; Zhou et al., 2019) is how to identify high-accurate label predictions. This task is a central problem of self-training methods. Most approaches consider the soft-max score of the predicted class to estimate the confidence of prediction, which is often not sufficient. In this regard, we present a margin prediction confidence, inspired by active learning approaches (Scheffer, Decomain, & Wrobel, 2001; Settles, 2009), which consider

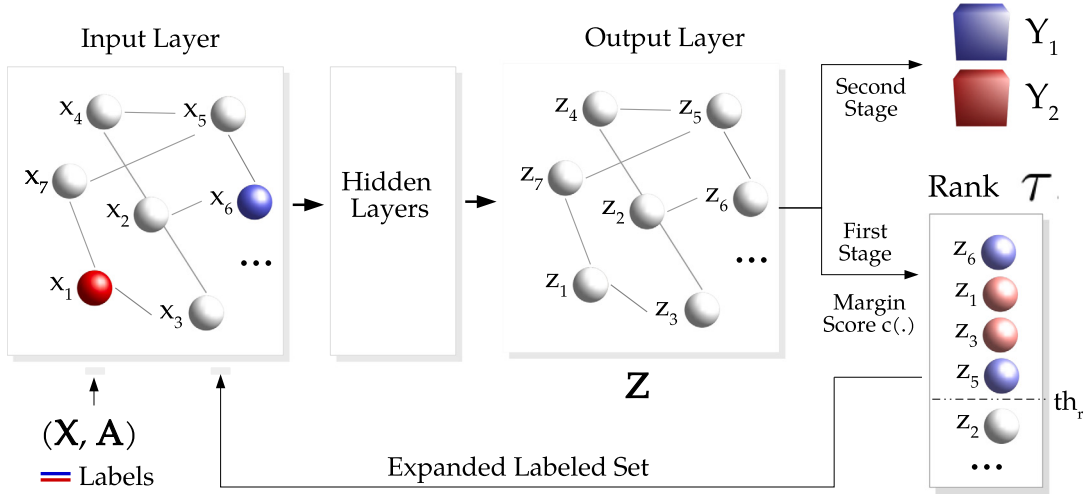


Fig. 1. Overview of the proposed rank-based self-training for graph convolutional networks.

the gap between the first and second higher scores. While active learning methods aims at identifying the most informative samples, we are interested in the most confident labels. Our score is based on the conjecture that the log-soft-max output tends to provide similar values to both first and second classes when the classification is not correct.

Different from other self-training methods, which keep a confidence score update at each epoch (Zhou et al., 2019) or multi-stage (Sun et al., 2020), our method employs a simple two-stage approach, i.e., we train the GCN twice, first with original label set, and then with the augmented label set. A direct advantage is improved efficiency and simplicity. A two-stage training approach requires the computation of the margin score and the global ranking of vertices. Considering a two-stage approach, such steps are computed only once, reducing substantially the computational efforts required. Hence a two-stage approach keeps a favorable trade-off between effectiveness and efficiency.

The proposed approach is illustrated in Fig. 1. Firstly, a first-stage classification is performed by a GCN model. Subsequently, the proposed approach can be divided into three main steps:

1. *Margin Confidence Score*: a margin-based score is computed for each graph node aiming to estimate the confidence on the computed embedding;
2. *Labeled Set Expansion*: the nodes are represented through a rank model, defined according to the margin score. Subsequently, some nodes are selected for an expanded labeled set;
3. *Second Stage Semi-Supervised Classification*: the GCN model is re-trained by taking into account the expanded labeled set.

The three main steps of the proposed method are detailed and formally defined in next sub-sections. A general outline of the proposed method is presented in Algorithm 1. Line 1 performs the first-stage classification. Lines 2–4 define the computation of the margin confidence, discussed in Section 3.1. The labeled set expansion step is defined in lines 5–12 and detailed in Section 3.2. Line 13 performs the second-stage classification, discussed in Section 3.3.

Algorithm 1 Rank-based Self-Training for GCNs

Require: Feature matrix $\mathbf{X}$, adjacency matrix $\mathbf{A}$, labeled data $\mathcal{V}_L, \alpha$

Ensure: Embedding matrix $\mathbf{Z}$, learned function $\hat{l}^{(2)}$

```

1:  $\mathbf{Z} = f_{GCN}(\mathbf{X}, \mathbf{A}, \mathcal{V}_L)$                                      # first-stage classification: learned function  $\hat{l}^{(1)}$ 
2: for all  $v_i \in \mathcal{V}$  do
3:    $c_i = c(v_i, \mathbf{Z})$                                            # margin score computation
4: end for
5:  $\tau = \text{sort}(\mathcal{V}, \mathbf{c})$                                            # ranked list computation
6:  $th_r = |\mathcal{V}| \times \alpha$                                            # threshold definition
7:  $\mathcal{V}_E = \mathcal{V}_L$                                                  # labeled set expansion
8: for all  $v_e \in \mathcal{V}$  do
9:   if  $\tau(v_e) \leq th_r$  then
10:     $\mathcal{V}_E = \mathcal{V}_E \cup v_e$ 
11:   end if
12: end for
13:  $\mathbf{Z} = f_{GCN}(\mathbf{X}, \mathbf{A}, \mathcal{V}_E)$                                      # second-stage classification: learned function  $\hat{l}^{(2)}$ 

```

3.1. Margin confidence score

The margin confidence score is defined based on the learned embeddings. Given a c -dimensional vector which defines the embedded representation $\mathbf{z}_i$ for a node v_i , a set S_i is computed containing the $\mathbf{z}_i$ absolute values. The set S_i is formally defined as:

$$S_i = \{|\mathbf{Z}_{ij}| : j \in \{0, 1, \dots, c\}\} \quad (3)$$

Once a log function is applied to the soft-max layer and the set S_i contains its absolute values, we are interested in the smallest values in S_i . Such values are associated with the most likely classes. Therefore, we consider a function $m(S_i, k)$ that returns the k th smallest element in S_i . Formally, the function $m : S_i \times \mathbb{N} \rightarrow \mathbb{R}$ can be defined as:

$$m(S_i, k) = \begin{cases} \min(S_i), & k = 1 \\ \min(\{v : v \in S_i, v > m(S_i, k-1)\}), & k > 1. \end{cases} \quad (4)$$

More specifically, the value of interest is given by the normalized difference between the smallest and the second smallest values. Therefore, the confidence estimation $c(v_i)$ of the predicted class for a given node v_i is defined as:

$$c(v_i, \mathbf{Z}) = \frac{m(S_i, 2) - m(S_i, 1)}{\sum_{v \in S_i} v}. \quad (5)$$

3.2. Labeled set expansion

Once a confidence estimation is defined, the nodes are ranked according to the function $c(\cdot)$ in order to obtain a ranked list τ . The ranked list τ can be formally defined as a permutation $(v_1, v_2, \dots, v_n)$ of the node set $\mathcal{V}$. A permutation τ is a bijection from the set $\mathcal{V}$ onto the set $[N] = \{1, 2, \dots, n\}$.

For a permutation τ , we interpret $\tau(v_i)$ as the position (or rank) of node v_i in the ranked list τ . If v_i is ranked before v_j in the ranked list τ , i.e., $\tau(v_i) < \tau(v_j)$, then $c(v_i, \mathbf{Z}) \geq c(v_j, \mathbf{Z})$. The rank $\tau(v_i)$ is used to decide if the node v_i is included in the expanded training set.

Let $\alpha \in [0, 1]$ denote a hyper-parameter that regulates the extend of labeled set expansion. Based on α , a ranking threshold th_r is defined as:

$$th_r = |\mathcal{V}| \times \alpha \quad (6)$$

The threshold th_r defines a position in the ranked list τ , until which the nodes are included in the expanded training set. Let $\mathcal{V}_E$ denotes the expanded labeled set, it is formally defined as:

$$\mathcal{V}_E = \{v_e \in \mathcal{V} | \tau(v_e) \leq th_r\} \cup \mathcal{V}_L \quad (7)$$

3.3. Second stage semi-supervised classification

Once an expanded labeled set $\mathcal{V}_E$ is defined, a second training stage is performed. Let $\hat{l}^{(1)}$ denotes the labeled function learned by the first training stage. The second stage training uses the same network model and same hyper-parameters of the first stage. However, the second stage training considers the set $\mathcal{V}_E$ as the labeled set to learn a new function $\hat{l}^{(2)}$. Finally, the function $\hat{l}^{(2)}$ is used for performing the definitive classification.

3.4. Complexity analysis

This section presents a brief discussion about the complexity of proposed rank-based self-training approach. The first and second stage classifications are directly associated with the GCN model used. Therefore, the complexity of our method is given by the computation of the margin confidence and the labeled set expansion procedure.

The margin score is computed based on the c -dimensional vector which defines the embedded representation, where c denotes the number of classes. Therefore, the complexity is $O(c)$. The labeled set expansion requires a sorting procedure of the vertices, which has $O(n \log n)$ complexity. Since $c \ll n$, the general complexity of the proposed model is given by $O(n \log n)$.

3.5. Rank aggregation for prediction confidence fusion

Effectively exploiting the useful information encoded in the unlabeled data is a central issue in semi-supervised learning. More specifically in self-training, the use of unlabeled data is closely associated with the identification of high-accurate predictions.

A true power of the proposed self-training approach becomes unlocked when a few diverse GCN models are exploited. The labeled set expansion is defined through the ranked list of vertices, which is computed based on the margin score. Once the margin score varies according to the GCN model, distinct GCN models yield different expanded labeled sets. Therefore, ranked lists from different GCN models present diversity, which can be exploited and aggregated for a more accurate label expansion step.

Additionally, even the same GCN model can give rise to distinct expanded labeled sets through distinct training executions, due to the stochastic characteristics of optimization procedures and random parameters initialization involved. In this scenario, we propose a rank aggregation approach to exploit such diversity. A fused prediction confidence score is computed by aggregating the ranks in two dimensions: distinct GCN models and different training executions of each model. The fused scores gives rise to a more effective aggregated ranking and, therefore, a more accurate expanded labeled set. Based on the expanded labeled set, a second stage classification is performed.

More formally, each training execution of each GCN model assigns a different prediction score to given pair (node, class). As a result, a different margin-based confidence score is computed for each execution. Let $\mathcal{M}$ denotes a set of GCN models, such that a model $M_j \in \mathcal{M}$ and $|\mathcal{M}| = m_f$. Let t denotes the current training execution, such that $t \in \{1, 2, \dots, n_f\}$. Let $c_{t,j}(v_i)$ denotes a confidence score computed by a GCN model M_j on the training execution t for a node v_i . A fused confidence score $c_f(v_i)$ is computed based on a multiplicative rank aggregation formulation (Pedronette & Torres, 2013), as:

$$c_f(v_i) = \prod_{j=1}^{m_f} \prod_{t=1}^{n_f} (1 + c_{t,j}(v_i)) \quad (8)$$

The fused confidence score $c_f(v_i)$ combines information from different GCN models and training executions and can be used in place of function $c(\cdot)$ (Eq. (5)) to define the ranked list τ . In this way, the aggregated ranking τ is exploited to define a more accurate expanded labeled set for a second-stage classification. The second stage classification can be performed by any model $M_j \in \mathcal{M}$. In order to ensure a consistent aggregation, the fused score $c_f(v_i)$ can only be considered if the same label is assigned to node v_i by all GCN models on all training executions. Otherwise, we set $c_f(v_i) = 0$.

4. Experimental evaluation

In this section, we describe the experiments conducted to assess the accuracy of the proposed method in the task of semi-supervised node classification.

4.1. Graph convolutional network models

As previously discussed, the proposed self-training approach allows different GCN models and variants. The margin-based score requirement is restricted only to the last layer, expected to be a log-soft-max operation. We validated the proposed method on four GCN models with a log-soft-max as the last layer, described in the following:

- *GCN*: Graph Convolution Network (Kipf & Welling, 2017), a seminal GCN model broadly defined as an efficient variant CNNs on graphs (detailed in Section 2.2);
- *SGC*: Simple Graph Convolution (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019), a simplification of GCN models obtained by removing nonlinearities and collapsing weight matrices between consecutive layers;
- *APPNP*: Approximate Personalized Propagation of Neural Predictions (Klicpera, Bojchevski, & Günnemann, 2019), an algorithm which exploits the relationship between GCNs and PageRank, deriving a propagation strategy based on personalized PageRank;
- *ARMA*: ARMA Filter Convolutions (Bianchi et al., 2019), a GCN variant which defines a convolutional layer based on Auto-Regressive Moving Average (ARMA) filters.

All four network models are used in the rank aggregation fusion to boost the correctness of the extended label sets. Diverse combinations of pairs and all models are considered on the experimental evaluation.

4.2. Datasets

The experimental evaluation was conducted on three citation network datasets¹: Cora (McCallum, Nigam, Rennie, & Seymore, 2000; Sen, et al., 2008) Citeseer (Giles, Bollacker, & Lawrence, 1998; Sen, et al., 2008) and Pubmed (Yang, Cohen, & Salakhutdinov, 2016). Such datasets have been largely used in the literature (Bai et al., 2019; Chen, Zhu, & Song, 2018; Fey & Lenssen, 2019; Velickovic, et al., 2018, 2019; Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) as benchmark for semi-supervised classification tasks, including some representative (Kipf & Welling, 2017; Yang et al., 2016) and recent (Bai et al., 2019; Velickovic, et al., 2019) works. Table 1 presents some statistics of the three datasets, briefly detailed in the following.

The three datasets are composed by textual documents and a list of citation links between them. A graph is defined, where each node represents a document and each edge represents a citation link. Additionally, a sparse bag-of-words feature vector is associated to each document. The Cora dataset contains scientific publications divided into 7 categories. Each publication is described by a binary bag-of-word representation, where 0 (or 1) indicates the absence (or presence) of the corresponding word from the dictionary. The dictionary consists of 1433 unique words (features). The Citeseer dataset employs an analogous representation but of 3, 703 dimensions. The Pubmed dataset is divided into 3 classes and uses a vectorial representation using Term Frequency-Inverse Document Frequency (TF-IDF), based on a dictionary of 500 terms. For all datasets, each document has a single class label.

¹ <https://linqs.soe.ucsc.edu/data>.

Table 1
Citation network datasets statistics.

Dataset	Nodes	Edges	Classes	Features	Train/Val/Test
Cora	2708	5429	7	1433	140/500/1000
CiteSeer	3327	4732	6	3703	120/500/1000
PubMed	19717	44338	3	500	60/500/1000

4.3. Experimental protocol and implementation details

The experiments were conducted according to the protocol initially used in Yang et al. (2016) and followed by other works (Fey & Lenssen, 2019; Kipf & Welling, 2017). The training is performed using 20 labels per class, but all feature vectors (labeled and unlabeled data). The accuracy prediction accuracy is evaluated on a test set of 1000 nodes. The dataset splits follows (Fey & Lenssen, 2019; Kipf & Welling, 2017; Yang et al., 2016), with a validation set of 500 labeled samples used by Kipf and Welling (2017) and Fey and Lenssen (2019). The labels of the validation set are not used for training.

The implementation of the proposed Rank-based Self-Training was made upon PyTorch Geometric (PyG) (Fey & Lenssen, 2019), a geometric deep learning extension library for PyTorch (Paszke, et al., 2017). We also used the PyG (Fey & Lenssen, 2019) implementation of the four network models previously discussed: Graph Convolution Network (GCN) (Kipf & Welling, 2017); Simple Graph Convolution (SGC) (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019); Approximate Personalized Propagation of Neural Predictions (APPNP) (Klicpera, Bojchevski, & Günnemann, 2019); and ARMA Filter Convolutions (ARMA) (Bianchi et al., 2019).

All the models were trained for 200 epochs using Adam (Kingma & Ba, 2015) optimization. During the training process, for both first and second stages, the model is selected according to the lowest validation loss. The hyperparameters and network configurations for each model followed the default values given by the benchmark provided in PyG (Fey & Lenssen, 2019).² The learning rate is defined as 0.01 for all networks except SGC, which used 0.1. The dropout parameter is set to 0.5 also for all networks except SGC, which does not employ dropout. The early stop window size was defined to 10 for all networks except ARMA, which used 100. The rank-based Self Training approach has only one hyper-parameter α , which, is discussed in next sub-section.

Regarding feature normalization, the PyG implementation includes a normalization procedure, which impacts positively on most of network models. However, our self-training approach performs better without the normalization. Therefore, we report results before self-training on both situations: with and without feature normalization.

4.4. Parameter space analysis

This section presents an analysis of the parameter space and definition of parameter settings. The proposed rank-based Self Training requires only one hyper-parameter α , which defines a rank threshold. The impact of α on accuracy was evaluated on Cora dataset considering an average of 20 executions. Fig. 2 presents the results for SGC and GCN networks.

We can observe that very low and high values of α lead to small gains in comparison with the network model in isolation (orange line). However, a large intermediary region tends to produce higher accuracy gains, indicating the robustness of our approach to small parameter variations. In all experiments we used $\alpha = 0.4$, which approximates the beginning of intermediary values with high associated accuracy gains.

The aggregation of network models also considers a parameter n_f , which defines the aggregation of executions on first stage and the average of executions on second stage classifications. We varied n_f in the interval [5,50] and evaluated the impact on accuracy. Fig. 3 shows the results, which are very stable to different settings. The value of $n_f = 20$ was empirically defined and used in all aggregation experiments.

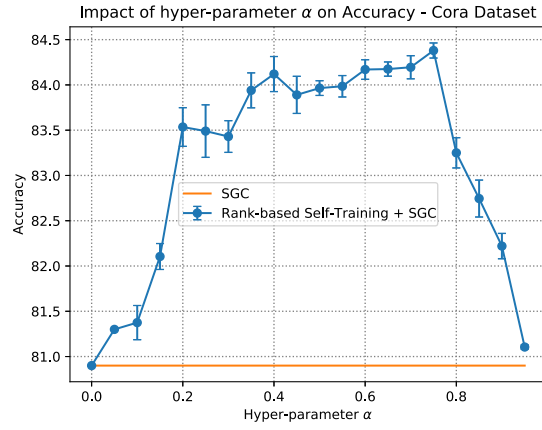
4.5. Results

This section presents the accuracy results on semi-supervised classification tasks for the citation datasets. Firstly, we report the accuracy results of the four network models in isolation, considering feature normalization. Table 2 presents the results reported according to PyG (Fey & Lenssen, 2019), with higher accuracy in bold.

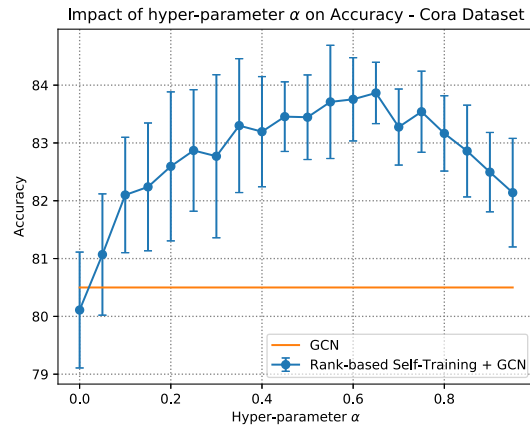
Table 3 presents the accuracy results for the proposed Rank-based Self-Training method. We reported the results of each network model in isolation and jointly with the proposed self-training approach. All results consider an average of 100 executions and do not use feature normalization. The best accuracy results are highlighted in bold. We can observe significant gains obtained for most network models and datasets. All the highest accuracy results (in bold) are achieved by the Rank-based Self-Training, even considering results of Table 2.

Results for fusion of confidence prediction are reported on Table 4. Different pairs combinations and the aggregation of all network models are considered. SGC and APPNP network models are considered for the second stage classification, once have achieved the best results on individual results reported in Table 3. As it can be observed, the results of fusion reached the high accuracy results for all datasets. The highest results for each segment of the table is highlighted in bold.

² https://github.com/rusty1s/pytorch_geometric/tree/master/benchmark.



(a) SGC



(b) GCN

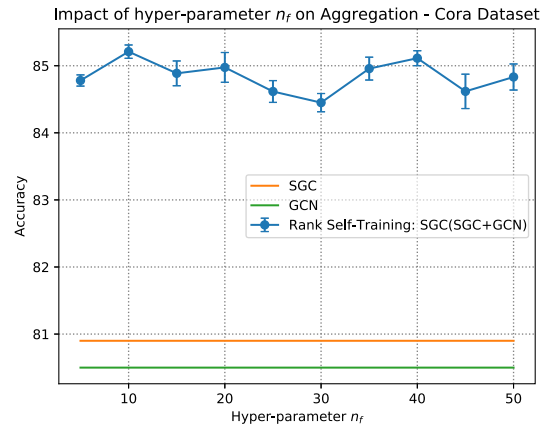
Fig. 2. Evaluation of the impact of hyper-parameter α on accuracy.Fig. 3. Evaluation of the impact of hyper-parameter n_f on aggregation.

Table 2

Accuracy of semi-supervised classification on citation datasets as reported in (Fey & Lenssen, 2019) with feature normalization.

Method	Cora	CiteSeer	PubMed
GCN (Kipf & Welling, 2017)	81.5 $\pm$ 0.6	71.1 $\pm$ 0.7	79.0 $\pm$ 0.6
SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019)	81.7 $\pm$ 0.1	71.3 $\pm$ 0.2	78.9 $\pm$ 0.1
ARMA (Bianchi et al., 2019)	82.8 $\pm$ 0.6	72.3 $\pm$ 1.1	78.8 $\pm$ 0.3
APPNP (Klicpera, Bojchevski, & Günnemann, 2019)	83.3 $\pm$ 0.5	71.8 $\pm$ 0.5	80.1 $\pm$ 0.2

Table 3

Accuracy of Rank-based Self-Training on citation datasets. Results without feature normalization.

Method	Cora	CiteSeer	PubMed
GCN (Kipf & Welling, 2017)	80.5 $\pm$ 0.6	66.9 $\pm$ 0.7	78.9 $\pm$ 0.3
Rank-based Self-Trainig + GCN	83.3 $\pm$ 0.9	69.4 $\pm$ 1.9	80.3 $\pm$ 0.4
SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019)	80.9 $\pm$ 0.0	69.3 $\pm$ 0.0	78.9 $\pm$ 0.0
Rank-based Self-Trainig + SGC	84.1 $\pm$ 0.2	73.1 $\pm$ 0.2	75.9 $\pm$ 0.0
ARMA (Bianchi et al., 2019)	80.1 $\pm$ 1.0	64.6 $\pm$ 2.2	78.2 $\pm$ 0.4
Rank-based Self-Trainig + ARMA	82.9 $\pm$ 1.0	68.0 $\pm$ 4.4	79.5 $\pm$ 0.6
APPNP (Klicpera, Bojchevski, & Günnemann, 2019)	82.8 $\pm$ 0.7	70.1 $\pm$ 0.7	79.9 $\pm$ 0.2
Rank-based Self-Trainig + APPNP	84.7 $\pm$ 0.6	71.2 $\pm$ 0.7	81.2 $\pm$ 0.7

Table 4

Accuracy results for confidence prediction fusion based on Rank Aggregation Self-Training.

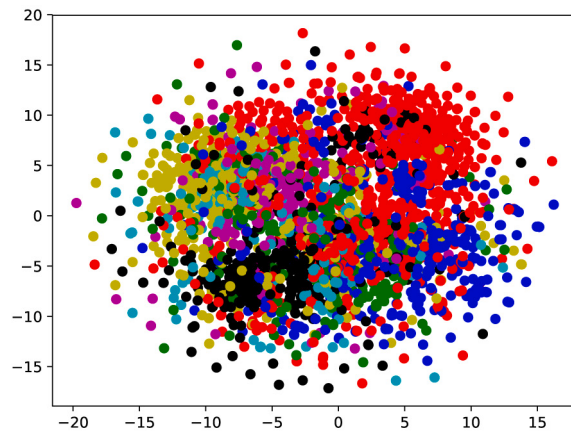
Method	Cora	CiteSeer	PubMed
Rank-based Self-Trainig + GCN	83.3 $\pm$ 0.9	69.4 $\pm$ 1.9	80.3 $\pm$ 0.4
Rank-based Self-Trainig + SGC	84.1 $\pm$ 0.2	73.1 $\pm$ 0.2	75.9 $\pm$ 0.0
Rank-based Self-Trainig + ARMA	82.9 $\pm$ 1.0	68.0 $\pm$ 4.4	79.5 $\pm$ 0.6
Rank-based Self-Trainig + APPNP	84.7 $\pm$ 0.6	71.2 $\pm$ 0.7	81.2 $\pm$ 0.7
Rank Aggregation Self-Training: SGC for second stage classification			
GCN+SGC	84.2 $\pm$ 0.2	73.2 $\pm$ 0.0	76.9 $\pm$ 0.0
GCN+ARMA	83.7 $\pm$ 0.0	73.3 $\pm$ 0.0	79.1 $\pm$ 0.0
GCN+APPNP	84.2 $\pm$ 0.0	73.3 $\pm$ 0.1	78.9 $\pm$ 0.0
SGC+ARMA	84.2 $\pm$ 0.2	73.4 $\pm$ 0.0	77.8 $\pm$ 0.0
SGC+APPNP	84.4 $\pm$ 0.2	72.8 $\pm$ 0.1	77.8 $\pm$ 0.0
ARMA+APPNP	84.1 $\pm$ 0.0	73.5 $\pm$ 0.1	79.0 $\pm$ 0.0
GCN+SGC+ARMA+APPNP	84.3 $\pm$ 0.0	74.1 $\pm$ 0.0;	78.3 $\pm$ 0.0
Rank Aggregation Self-Training: APPNP for second stage classification			
GCN+SGC	84.0 $\pm$ 0.4	71.4 $\pm$ 0.4	79.3 $\pm$ 0.3
GCN+ARMA	84.6 $\pm$ 0.3	71.6 $\pm$ 0.5	80.3 $\pm$ 0.4
GCN+APPNP	84.9 $\pm$ 0.5	71.4 $\pm$ 0.5	80.8 $\pm$ 0.2
SGC+ARMA	84.7 $\pm$ 0.4	72.1 $\pm$ 0.3	79.3 $\pm$ 0.2
SGC+APPNP	85.1 $\pm$ 0.4	72.5 $\pm$ 0.5	79.6 $\pm$ 0.2
ARMA+APPNP	85.0 $\pm$ 0.3	71.9 $\pm$ 0.5	80.9 $\pm$ 0.2
GCN+SGC+ARMA+APPNP	85.0 $\pm$ 0.5	72.1 $\pm$ 0.5	79.9 $\pm$ 0.2

4.6. Visual analysis

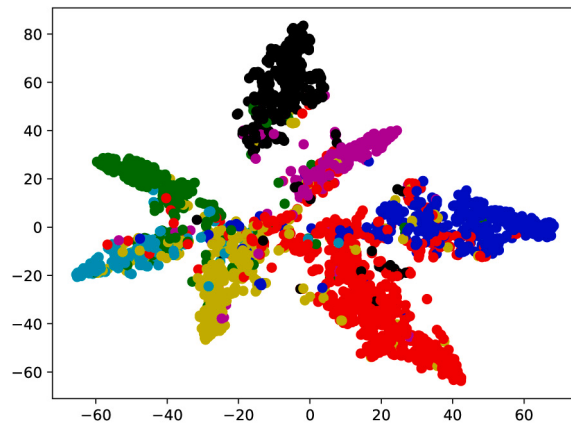
In order to enrich the discussion about the impact of the proposed method, we present a qualitative visualization of the feature space. The visual analysis illustrates the outcome of the method through a 2-D projection of data features and their respective computed embeddings. The analysis is conducted on three datasets with the 2-D projected space computed by the t-SNE (van der Maaten & Hinton, 2008) algorithm.

Fig. 4 shows three t-SNE visualizations on the Cora dataset: (a) depicting the raw dataset features; (b) the embeddings computed by SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) model; and (c) the embeddings computed by the proposed self-training approach based on SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019). Despite of the subjectivity associated to the visualization, we can observe a higher cohesion of data items according to the seven topic classes of Cora and better separation of different classes for the proposed self-training embedding.

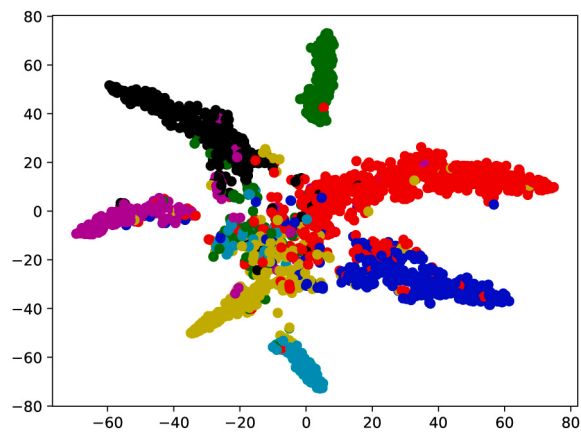
Fig. 5 illustrates analogous t-SNE visualizations on CiteSeer and PubMed datasets. The impact of self-training is more pronounced on the CiteSeer dataset, where the embeddings representations of each class are thinned. On the PubMed dataset, the difference is slight, but a better separation between blue and black classes can be observed.



(a) Cora Raw features

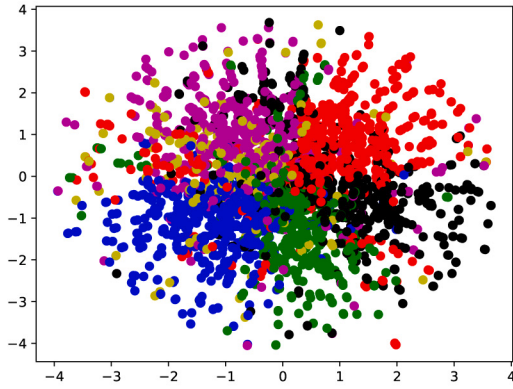


(b) Cora SGC embeddings

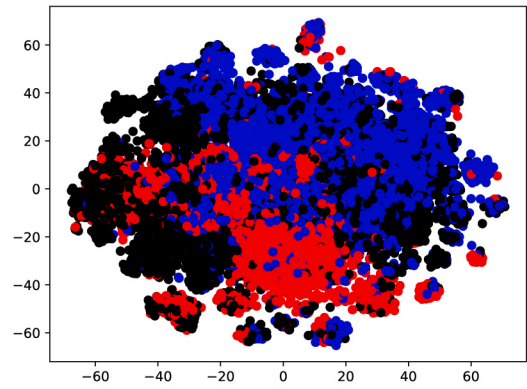


(c) Cora Self-Training SGC embeddings

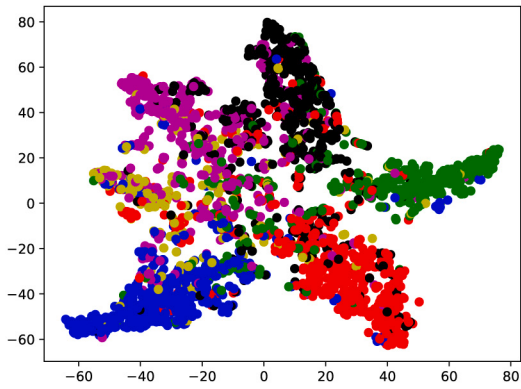
Fig. 4. t-SNE (van der Maaten & Hinton, 2008) visualization of initial features and computed embeddings for the Cora dataset.



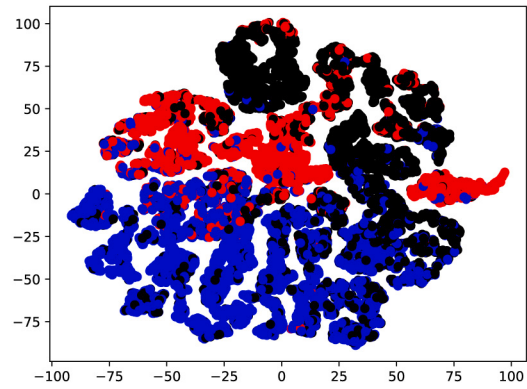
(a) CiteSeer Raw features



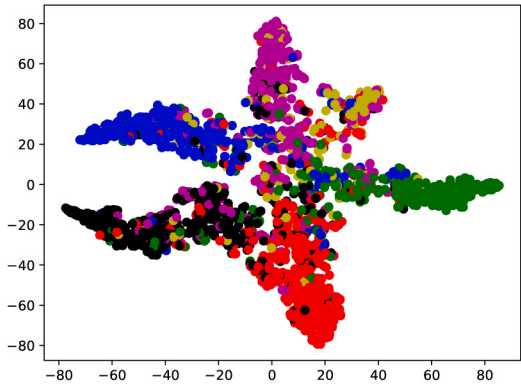
(d) PubMed Raw features



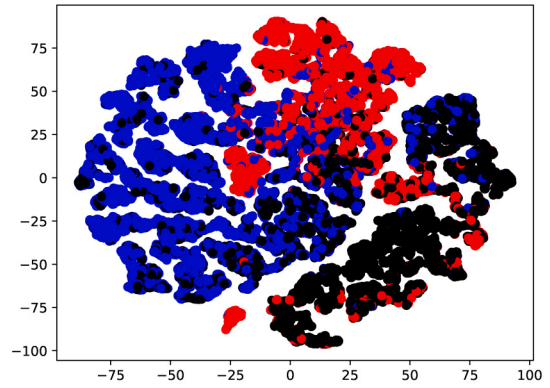
(b) CiteSeer SGC embeddings



(e) PubMed SGC embeddings



(c) CiteSeer Self-Training SGC embeddings



(f) PubMed Self-Training SGC embeddings

Fig. 5. t-SNE (van der Maaten & Hinton, 2008) visualization of initial features and computed embeddings for the CiteSeer and PubMed datasets. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

4.7. Comparison with state-of-the-art

We compared the proposed method with various state-of-the-art algorithms on semi-supervised classification tasks. The comparison was conducted on Cora, Citeseer, and Pubmed datasets. The experimental setting and data splits followed the protocol

Table 5
Comparison with state-of-the-art methods in terms of classification accuracy (%).

Method	Cora	Citeseer	Pubmed
Manifold Regularization (Belkin, Niyogi, & Sindhwani, 2006)	59.5	60.1	70.7
Semi-Supervised Embedding (Weston, Ratle, & Collobert, 2008)	59.0	59.6	71.1
Label Propagation (LP) (Zhu et al., 2003)	68.0	45.3	63.0
DeepWalk (Perozzi et al., 2014)	67.2	43.2	65.3
Iterative Classification Algorithm (ICA) (Lu & Getoor, 2003)	75.1	69.1	73.9
Planetoid (Yang et al., 2016)	75.7	64.7	77.2
Graph Convolution Netork (GCN) (Kipf & Welling, 2017)	81.5	70.3	79.0
Graph Attention Network (GAT) (Velickovic, et al., 2018)	83.0	72.5	79.0
Variance Reduction (Chen et al., 2018)	82.0	72.9	79.0
Self-Training (Li et al., 2018)	80.5	69.9	78.3
Simple Graph Convolution (SGC) (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019)	81.0	71.9	78.9
Deep Graph Infomax (DGI) (Velickovic, et al., 2019)	82.3	71.8	76.8
Hypergraph Convolution and Attention (Bai et al., 2019)	82.7	71.2	78.4
Auto-Regressive Moving Average Filters (ARMA) (Bianchi et al., 2019)	83.4	72.5	78.9
Approx. Pers. Propagation of Neural Pred. (APPNP) (Klicpera, Bojchevski, & Günnemann, 2019)	83.3	71.8	80.1
Rank-based Self-Training + SGC	84.1	73.1	75.9
Rank-based Self-Training + APPNP	84.7	71.2	81.2
Rank Agregation Self-Training: SGC (ARMA+APPNP)	84.1	73.5	79.0
Rank Agregation Self-Training: APPNP (SGC+APPNP)	85.1	72.5	79.6

used in Kipf and Welling (2017) and Yang et al. (2016), which have been established as a standard benchmark protocol in the area. The classification accuracy of the compared methods is directly quoted from the literature. Several methods were considered, from traditional semi-supervised learning methods (Lu & Getoor, 2003; Zhu, Ghahramani, & Lafferty, 2003) to more recent baselines (Perozzi, Al-Rfou, & Skiena, 2014; Yang et al., 2016). Representative works, as GCN (Kipf & Welling, 2017) and GAT (Velickovic, et al., 2018), were also considered in addition to very recent network models (Bai et al., 2019; Velickovic, et al., 2019). Related self-training approaches were also included in the comparison, considering the higher accuracy results reported in Li et al. (2018). Other self-training methods (Sun et al., 2020; Zhou et al., 2019) were not included due to distinct experiments settings and data splits from Yang et al. (2016).

Table 5 presents the accuracy results on the three datasets. We report the results for Rank-based Self-Training considering the network models which presented the best accuracy results in isolation: APPNP (Klicpera, Bojchevski, & Günnemann, 2019) and SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019). For the Rank Agregation Self-Training, we considered two selected combination of pairs reported in Table 4: ARMA (Bianchi et al., 2019) + APPNP (Klicpera, Bojchevski, & Günnemann, 2019) with SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) as second stage classification and SGC (Wu, Souza, Zhang, Fifty, Yu, & Weinberger, 2019) + APPNP (Klicpera, Bojchevski, & Günnemann, 2019) with APPNP (Klicpera, Bojchevski, & Günnemann, 2019) for second stage. As we can observe, the proposed approach achieved the best accuracy results on the three datasets in comparison with all state-of-art methods (in bold).

5. Conclusion

In this paper, we proposed a simple and effective self-training approach for Graph Convolutional Networks. Our approach uses a margin-based score computed over log-soft-max layer of GCNs and analyzed through a rank-based model. Considering the crucial role of unlabeled data on semi-supervised learning, the proposed Rank-based Self Training approach allows an effective labeled set expansion and more accurate results on a second-stage classification. Evaluated on different GCN models, our approach achieved state-of-the-art results on benchmarks widely used in the literature. In future work, we intend to investigate the use of our approach on graph classification tasks, in addition to node classification.

CRedit authorship contribution statement

Daniel Carlos Guimarães Pedronette: Conceptualization, Methodology, Software, Visualization, Writing - original draft. **Longin Jan Latecki:** Conceptualization, Supervision, Formal analysis, Writing - review & editing.

Acknowledgments

The authors are grateful to Fulbright Commission, São Paulo Research Foundation — FAPESP, Brazil (grants #2018/15597-6 and #2017/25908-6), Brazilian National Council for Scientific and Technological Development — CNPq (grant #308194/2017-9) and Microsoft Research, USA. This work was also partly supported by the National Science Foundation, USA Grant No. IIS-1814745.

References

- Bai, S., Zhang, F., & Torr, P. H. S. (2019). Hypergraph convolution and hypergraph attention. CoRR abs/1901.08150.
- Belkin, M., Niyogi, P., & Sindhwan, V. (2006). Manifold regularization: A geometric framework for learning from labeled and unlabeled examples. *Journal of Machine Learning Research*, 7, 2399–2434.
- Bianchi, F. M., Grattarola, D., Livi, L., & Alippi, C. (2019). Graph Neural Networks with convolutional ARMA filters. CoRR abs/1901.01343.
- Cai, H., Zheng, V. W., & Chang, K. C. (2018). A comprehensive survey of graph embedding: Problems, techniques, and applications. *IEEE Transactions on Knowledge and Data Engineering*, 30(9), 1616–1637.
- Chapelle, O., Schölkopf, B., & Zien, A. (2010). *Semi-supervised learning* (1st ed.). The MIT Press.
- Chen, J., Zhu, J., & Song, L. (2018). Stochastic training of graph convolutional networks with variance reduction. In *International conference on machine learning* (vol. 80). (pp. 941–949).
- Fey, M., & Lenssen, J. E. (2019). Fast graph representation learning with pytorch geometric. CoRR abs/1903.02428.
- Giles, C. L., Bollacker, K. D., & Lawrence, S. (1998). *CiteSeer: An automatic citation indexing system*.
- Kingma, D. P., & Ba, J. (2015). Adam: A Method for Stochastic Optimization. In Bengio, Y. and LeCun, Y. (Eds.), *International conference on learning representations*.
- Kipf, T. N., & Welling, M. (2017). Semi-supervised classification with graph convolutional networks. In *5th International conference on learning representations*.
- Klicpera, J., Bojchevski, A., & Günnemann, S. (2019). Predict then propagate: Graph Neural Networks meet Personalized PageRank. In *International conference on learning representations*.
- Klicpera, J., Weissenberger, S., & Günnemann, S. (2019). Diffusion improves graph learning. In *Advances in neural information processing systems* (pp. 13333–13345).
- Li, Q., Han, Z., & Wu, X. (2018). Deeper insights into graph convolutional networks for semi-supervised learning. In S. A. McIlraith, & K. Q. Weinberger (Eds.), *Proceedings of the thirty-second AAAI conference on artificial intelligence* (pp. 3538–3545). AAAI Press.
- Lu, Q., & Getoor, L. (2003). Link-based classification. In *International conference on international conference on machine learning* (pp. 496–503). AAAI Press.
- van der Maaten, L., & Hinton, G. (2008). Visualizing data using t-SNE. *Journal of Machine Learning Research*, 9, 2579–2605.
- McCallum, S. K., Nigam, K., Rennie, J., & Seymore, K. (2000). Automating the construction of internet portals with machine learning. *Information Retrieval*, 3, 127–163.
- Paszke, A., Gross, S., Chintala, S., Chanan, G., Yang, E., DeVito, Z., et al. (2017). Automatic differentiation in PyTorch. In *NIPS-W*.
- Pedronette, D. C. G., & Torres, R. D. S. (2013). Image re-ranking and rank aggregation based on similarity of ranked lists. *Pattern Recognition*, 46(8), 2350–2360.
- Pei, H., Wei, B., Chang, K. C., Lei, Y., & Yang, B. (2020). Geom-GCN: Geometric graph convolutional networks. In *International conference on learning representations*.
- Perozzi, B., Al-Rfou, R., & Skiena, S. (2014). DeepWalk: Online learning of social representations. In *ACM SIGKDD international conference on knowledge discovery and data mining*. (pp. 701–710).
- Scheffer, T., Decomain, C., & Wrobel, S. (2001). Active hidden Markov models for information extraction. In *International conference on advances in intelligent data analysis* (pp. 309–318). Berlin, Heidelberg: Springer-Verlag.
- Scudder, H. J. (1965). Probability of error of some adaptive pattern-recognition machines. *IEEE Transactions on Information Theory*, 11(3), 363–371.
- Sen, P., Namata, G., Bilgic, M., Getoor, L., Galligher, B., & Eliassi-Rad, T. (2008). Collective classification in network data. *AI Magazine*, 29(3), 93.
- Settles, B. (2009). *Active learning literature survey: Technical report 1648*, University of Wisconsin–Madison, URL <http://axon.cs.byu.edu/~martinez/classes/778/Papers/settles.activelearning.pdf>.
- Sun, K., Lin, Z., & Zhu, Z. (2020). Multi-stage self-supervised learning for graph convolutional networks on graphs with few labeled nodes. In *Proceedings of the thirty-second AAAI conference on artificial intelligence*. AAAI Press.
- Tian, Q., Yu, J., Xue, Q., & Sebe, N. (2004). A new analysis of the value of unlabeled data in semi-supervised learning for image retrieval. In *2004 IEEE International conference on multimedia and expo* (vol. 2). (pp. 1019–1022).
- Triguero, I., García, S., & Herrera, F. (2015). Self-labeled techniques for semi-supervised learning: taxonomy, software and empirical study. *Knowledge and Information Systems*, 42(2), 245–284.
- Velickovic, P., Cucurull, G., Casanova, A., Romero, A., Liò, P., & Bengio, Y. (2018). Graph attention networks. In *6th international conference on learning representations*.
- Velickovic, P., Fedus, W., Hamilton, W. L., Liò, P., Bengio, Y., & Hjelm, R. D. (2019). Deep Graph Infomax. In *International conference on learning representations*.
- Weston, J., Ratle, F., & Collobert, R. (2008). Deep learning via semi-supervised embedding. In *International conference on machine learning*. (pp. 1168–1175).
- Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., & Yu, P. S. (2019). A comprehensive survey on graph neural networks. CoRR abs/1901.00596.
- Wu, F., Souza, A., Zhang, T., Fifty, C., Yu, T., & Weinberger, K. (2019). Simplifying graph convolutional networks. In *International conference on machine learning* (vol. 97). (pp. 6861–6871).
- Yang, Z., Cohen, W. W., & Salakhutdinov, R. (2016). Revisiting semi-supervised learning with graph embeddings. In *International conference on international conference on machine learning* (pp. 40–48). JMLR.org.
- Zhou, Z., Zhang, S., & Huang, Z. (2019). Dynamic self-training framework for graph convolutional networks. arXiv:abs/1910.02684.
- Zhu, X., Ghahramani, Z., & Lafferty, J. (2003). Semi-supervised learning using gaussian fields and harmonic functions. In *Proceedings of the twentieth international conference on international conference on machine learning*. (pp. 912–919).