

# Graph Component Contrastive Learning for Concept Relatedness Estimation

Yueen Ma<sup>1</sup>, Zixing Song<sup>1</sup>, Xuming Hu<sup>2</sup>, Jingjing Li<sup>1</sup>, Yifei Zhang<sup>1</sup>, Irwin King<sup>1</sup>

<sup>1</sup>The Chinese University of Hong Kong,

<sup>2</sup>Tsinghua University

{yema21, zxsong, lij, yfzhang, king}@cse.cuhk.edu.hk, hxm19@mails.tsinghua.edu.cn

## Abstract

Concept relatedness estimation (CRE) aims to determine whether two given concepts are related. Existing methods only consider the pairwise relationship between concepts, while overlooking the higher-order relationship that could be encoded in a concept-level graph structure. We discover that this underlying graph satisfies a set of intrinsic properties of CRE, including *reflexivity*, *commutativity*, and *transitivity*. In this paper, we formalize the CRE properties and introduce a graph structure named ConcreteGraph. To address the data scarcity issue in CRE, we introduce a novel data augmentation approach to sample new concept pairs from the graph. As it is intractable for data augmentation to fully capture the structural information of the ConcreteGraph due to a large amount of potential concept pairs, we further introduce a novel Graph Component Contrastive Learning framework to implicitly learn the complete structure of the ConcreteGraph. Empirical results on three datasets show significant improvement over the state-of-the-art model. Detailed ablation studies demonstrate that our proposed approach can effectively capture the high-order relationship among concepts.

## Introduction

Concept relatedness estimation (CRE) is the task of determining whether two concepts are related. In CRE, a concept can be a Wikipedia entry, a news article, a social media post, etc. Table 1 includes a pair of related concepts and an unrelated concept. In this example, when given the first two concepts, “Open-source software” and “GNU General Public License”, the goal of CRE is to label them as a related pair; but for “Open-source software” and “Landscape architecture”, one should label them as unrelated.

Existing CRE methods (Liu et al. 2019a) only consider the provided concept pairs, which are low-order pairwise relationships. However, we discover that there exist higher-order relationships in CRE that can be encoded by a concept-level graph structure, where the original pairs correspond to immediate neighbors (Ein-Dor et al. 2018; Liu et al. 2019a). The higher-order relationships are manifested as the multi-hop neighbors in the graph, which are validated by three types of intrinsic properties of CRE: *reflexivity*, *commutativity*, and *transitivity*. To construct the graph, we

|           |                                                                                                                                                                                                                                                          |
|-----------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Related   | <b>Open-source software</b> (OSS) is computer software that is released under a license in which the copyright holder grants users the rights to use, study, change, and distribute the software and its source code to anyone and for any purpose . . . |
|           | The <b>GNU General Public License</b> (GNU GPL or simply GPL) is a series of widely used free software licenses that guarantee end users the freedom to run, study, share, and modify the software . . .                                                 |
| Unrelated | <b>Landscape architecture</b> is the design of outdoor areas, landmarks, and structures to achieve environmental, social-behavioural, or aesthetic outcomes . . .                                                                                        |

Table 1: Examples of concept relatedness estimation from the WORD dataset.

treat each concept as a node and add edges between any related concept pairs. In this paper, we name it ConcreteGraph (**Concept relatedness Graph**). It enables us to obtain new concept pairs that are not limited to low-order relationships. For example, one of the *transitivity* properties states that if concepts  $(x_A, x_B)$  are related and concepts  $(x_B, x_C)$  are related, then we can extract a new related concept pair  $(x_A, x_C)$ , where  $x_C$  is a 2-hop neighbor of  $x_A$ . Therefore, when we sample new concept pairs from the ConcreteGraph, the CRE properties are implicitly utilized.

The main challenge is how to take full advantage of the structural information of the ConcreteGraph. The most straightforward approach is to explicitly add new concept pairs to the dataset, which can be seen as a data augmentation method for alleviating the data scarcity issue. However, this approach exhibits a problem: there is an excessive amount of potential concept pairs. Assuming that we extract all possible related concept pairs by adding edges between any two connected concepts in the ConcreteGraph, we essentially turn every component of the ConcreteGraph into a complete subgraph. Because the edges in the original ConcreteGraph are sparse, the number of new related concept pairs grows quadratically with the number of concepts. It would become impractical for a model to learn all of them and, thus, we must only sample a subset of all possible concept pairs. If we sample completely at random, the quality of

the new concept pairs cannot be guaranteed. That is, when the path length between two concepts is long, the quality of their relationship tends to be low because the probability of the existence of a noisy edge on the path becomes high. This phenomenon is later confirmed by our experiments. A simple solution is to extract concept pairs from the local  $k$ -hop neighborhood. But this limits the ConcreteGraph-based data augmentation to contain only the local graph structure.

To this end, we introduce a novel Graph Component Contrastive Learning (GCCL) framework that can implicitly capture the complete structural information of the ConcreteGraph. Rather than explicitly learning new concept pairs, the objective of GCCL is that a concept treats its own component as positive, whereas it forms negative pairs with all other components. The representation of a concept or a component is achieved using a shared backbone encoder. Thus, GCCL gives the encoder access to both local and global structural information, which would be intractable if we used the ConcreteGraph-based data augmentation method. GCCL is inspired by Prototypical Contrastive Learning (PCL) (Li et al. 2021), but the difference is that it does not use  $k$ -means to find clusters. Each component of the ConcreteGraph corresponds to a cluster, and therefore, components in GCCL translate to prototypes in PCL.

Specifically, GCCL involves three main steps: (1) we use a Transformer model as the backbone encoder to embed all concepts and components before each epoch; (2) for each concept, we train the encoder to distinguish the concept's own component from all other components, which is defined to be the GC-NCE objective; (3) along with the main GC-NCE target, we also use a momentum encoder as in the Momentum Contrast (MoCo) framework (He et al. 2020; Chen et al. 2020b) to optimize an InfoNCE loss (van den Oord, Li, and Vinyals 2018). This InfoNCE loss contrasts a concept with other concepts. It is added to the main GC-NCE target to preserve the local smoothness of the overall GCCL loss function (Li et al. 2021). We conduct comprehensive experiments with three different Transformer models on three datasets. The empirical results show that the GCCL framework is capable of significantly improving their performance and, when combined with the ConcreteGraph-based data augmentation method, it can sometimes bring even more improvement. We also conduct detailed ablation studies to show the effectiveness of each module in our method. Our code is available on Github<sup>1</sup>.

The main contributions of our paper are as follows:

- For CRE, we formally summarize its intrinsic properties that can be categorized into three types: *reflexivity*, *commutativity*, and *transitivity*. On their basis, we construct a concept-level graph structure, ConcreteGraph, which encodes not only the provided concept pairs but also the higher-order relationships between concepts. To the best of our knowledge, we are the first to find such a graph structure for the CRE task.
- We propose a novel Graph Component Contrastive Learning (GCCL) framework for taking full advantage

of the local and global structural information of the ConcreteGraph. GCCL can be complemented by a novel ConcreteGraph-based data augmentation method that explicitly provides local neighborhood information.

- Our method significantly improves over the state-of-the-art Concept Interaction Graph (CIG) model (Liu et al. 2019a) on two datasets, CNSE and CNSS. We are also the first to apply deep neural networks to the WORD dataset (Ein-Dor et al. 2018).

## Related Work

### Concept Relatedness Estimation

The concept relatedness estimation (CRE) task stemmed from the concept similarity matching (CSM) task in the area of formal concept analysis (FCA). In FCA, a concept is formally defined as a pair of sets: a set of objects and a set of attributes in a given domain (Formica 2006). Methods for assessing concept similarity include ontology-based methods (Formica 2006, 2008), Tversky's-Ratio-based methods (Lombardi and Sartori 2006), rough-set-based methods (Wang and Liu 2008), and semantic-distance-based methods (Ge and Qiu 2008; Li and Xia 2011). The definition for concepts in FCA is not suitable for the CRE task because CRE concepts are long text documents and, thus, the CSM methods cannot be applied to the CRE task. Inspired by the giant success of introducing deep neural networks into natural language processing applications (Li et al. 2019; Gao et al. 2020; Sun et al. 2022; Li et al. 2020), we adopt Transformer models (Vaswani et al. 2017) to address the CRE task.

CRE is also related to tasks such as semantic textual similarity (Cer et al. 2017; Zhang and Zhu 2019, 2020), text similarity (Thijs 2019), text relatedness (Tsatsaronis, Varlamis, and Vazirgiannis 2014), text matching (Jiang et al. 2019), and text classification (Hu et al. 2020, 2021a,b; Liu et al. 2022; Li et al. 2022). The current state-of-the-art model for semantic textual similarity is XLNet (Yang et al. 2019), but there is still no work that applies deep neural networks to the WORD dataset. With the advancement of Graph Neural Networks (GNNs) (Song and King 2022; Song, Zhang, and King 2022; Song et al. 2021, 2022; Zhang et al. 2022a), Liu et al. (2019a) introduced the Concept Interaction Graph (CIG) method to match news article pairs along with two new datasets, CNSE and CNSS. CIG is the current state-of-the-art model on these two datasets.

CRE can play an important role in a wide range of applications, such as information retrieval (Busch et al. 2012; Teevan, Ramage, and Morris 2011), document clustering (Aswani Kumar and Srinivas 2010), plagiarism detection (Muangprathub, Kajornkasirat, and Wanichsombat 2021), etc. Therefore, CRE has been attracting much interest lately.

### Contrastive Learning

Since the introduction of contrastive learning (Chopra, Hadsell, and LeCun 2005), many variants have been developed. InfoNCE (van den Oord, Li, and Vinyals 2018) aims to find a positive in a group of noise samples. MoCo (He et al. 2020; Chen et al. 2020b) provides a queue-based framework for utilizing data from previous batches with the addition of a

<sup>1</sup>Github: <https://github.com/Panmani/GCCL>

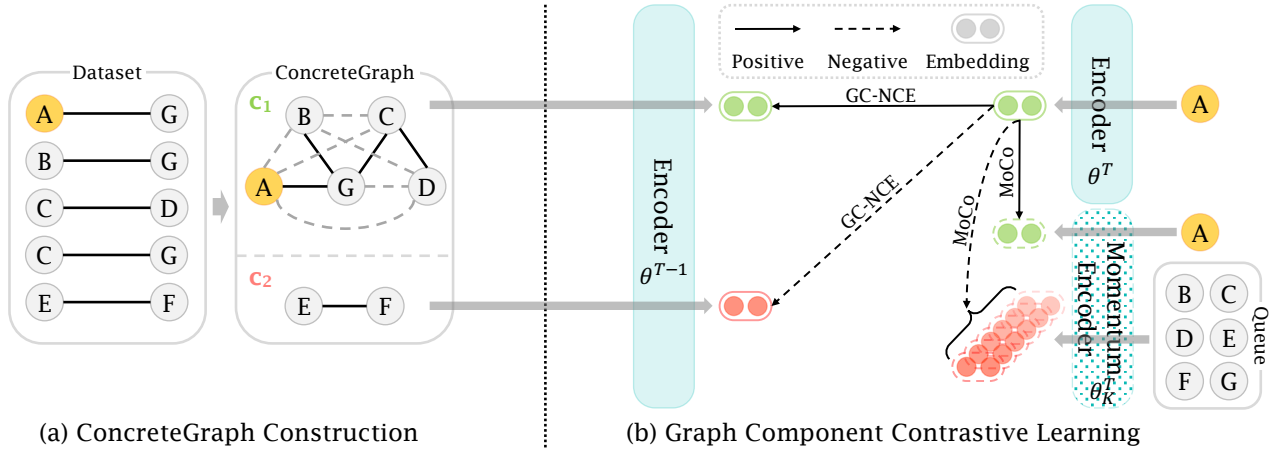


Figure 1: Overview of our proposed framework for concept relatedness estimation. (a) After collecting the concept pairs into the ConcreteGraph, the concepts are divided into multiple components, e.g.,  $c_1$  and  $c_2$ . For simplicity, concepts are represented with their indices, e.g., the node with index A represents concept  $x_A$ ; edges connect related concepts: solid black edges are true annotations, and dashed gray edges correspond to new related concept pairs from ConcreteGraph-based data augmentation; disconnected pairs are unrelated ones, such as  $(x_A, x_F)$ . (b) Under the Graph Component Contrastive Learning (GCCL) framework, the GC-NCE objective aims to distinguish each concept’s own component from other components (in this example, the concept is  $x_A$ , highlighted in yellow); the MoCo objective is to contrast the momentum embedding of each concept with the momentum embeddings of a queue of other concepts. The positive embeddings are colored in green and the negative ones are in red.  $\theta^{T-1}$  and  $\theta^T$  are both the parameters of the main encoder but from two different epochs;  $\theta_K^T$  is the parameters of the momentum encoder from the MoCo framework.

momentum encoder. Due to the use of large Transformer models (Devlin et al. 2019; Liu et al. 2019b; Yang et al. 2019) in our method, the batch size is limited. MoCo provides a solution to decouple the batch size from the number of negatives. Prototypical Contrastive Learning (PCL) (Li et al. 2021) treats prototypes as latent variables and brings clustering into the contrastive learning paradigm. The prototypes are found using the  $k$ -means clustering algorithm and its ProtoNCE loss contrasts different prototypes based on the clustering results. However, the CRE task naturally exhibits a cluster structure, and, therefore, we do not need to rely on a clustering algorithm. Our method is not to be confused with Graph Contrastive Learning (GCL) (You et al. 2020; Zhang et al. 2022b) whose main inputs are graphs, whereas the inputs are text documents in this work.

## Our Method

### Concept Relatedness Estimation

The concept relatedness estimation (CRE) task is to predict whether two given concepts are related or unrelated. Thus, it is a binary classification task with two labels “related” and “unrelated”. Formally, given a set of concepts  $X = \{x_i \mid i \in \mathbb{N}_n\}$  with an index set  $\mathbb{N}_n = \{1, \dots, n\}$  and a set of known pairwise binary labels  $Y = \{y_{i,j} \mid (i, j) \subset \mathbb{N}_n \times \mathbb{N}_n\}$ , the CRE task is to learn a parameterized encoder  $f_\theta$  that can estimate the true relatedness  $y_{i,j}$  of two concepts  $x_i$  and  $x_j$ . In this paper, we focus on concepts that are long documents.

**CRE Properties.** The CRE task exhibits some unique intrinsic properties that are rarely present in typical NLP tasks.

To state these properties formally, we assume three concepts  $x_A$ ,  $x_B$ , and  $x_C$ . The similarity symbol “ $\sim$ ” is used to denote that two concepts are “related”, while the dissimilarity symbol “ $\approx$ ” is used to connect two unrelated concepts.

**Property 1 (Reflexivity).** A concept  $x_A$  is related to itself:  $x_A \sim x_A$ .

**Property 2 (Commutativity of Relatedness).** If  $x_A$  and  $x_B$  are related, then  $x_B$  and  $x_A$  are related:  $x_A \sim x_B \iff x_B \sim x_A$ .

**Property 3 (Commutativity of Unrelatedness).** If  $x_A$  and  $x_B$  are unrelated, then  $x_B$  and  $x_A$  are unrelated:  $x_A \approx x_B \iff x_B \approx x_A$ .

**Property 4 (Transitivity of Relatedness).** If  $x_A$  is related to  $x_B$  and  $x_B$  is related to  $x_C$ , then  $x_A$  and  $x_C$  are related:  $(x_A \sim x_B) \wedge (x_B \sim x_C) \implies x_A \sim x_C$ .

**Property 5 (Transitivity of Unrelatedness).** If  $x_A$  is related to  $x_B$  but  $x_B$  is unrelated to  $x_C$ , then  $x_A$  and  $x_C$  are unrelated:  $(x_A \sim x_B) \wedge (x_B \approx x_C) \implies x_A \approx x_C$ .

**ConcreteGraph.** To make use of the CRE properties in practice, we can build a graph to encode both the pairwise relationships from the dataset and the higher-order relationships based on those properties. We name it ConcreteGraph (**C**oncept **r**elatedness **G**raph). An illustration of how to construct the ConcreteGraph is shown in Figure 1(a). Each node in ConcreteGraph represents a concept, and edges are only added between related concept pairs. The construction of the ConcreteGraph takes  $O(|Y|)$  time as we need to iterate over all labeled concept pairs. Since it is a binary classification task, we do not need to set the edge weights. By the

nature of CRE, the ConcreteGraph usually has multiple connected components, where each component contains a set of related concepts. To identify the components of the ConcreteGraph, we can use breadth-first search (BFS), which takes  $O(V + E)$  time where  $V$  is the number of nodes and  $E$  is the number of edges.

### Data Augmentation

The ConcreteGraph enables a straightforward data augmentation method for addressing the data scarcity problem. The CRE properties automatically come into play when we sample new concept pairs from the ConcreteGraph. *Commutativity* is used if we sample two nodes that are already provided by the dataset but in a different order. For example, if  $(x_A, x_B)$  is provided and we sample  $(x_B, x_A)$ , then this new pair is validated by the *commutativity of relatedness* property (Property 2) or the *commutativity of unrelatedness* property (Property 3). *Transitivity* is utilized when the path between the two sampled concepts has at least two edges or when there is no path between them. For example, if there is a path  $x_A - x_B - x_C$ , then the new concept pair  $(x_A, x_C)$  is justified by the *transitivity of relatedness* property (Property 4). When two concepts are sampled from two different components, the *transitivity of unrelatedness* property (Property 5) is employed to prove that they are unrelated.

In theory, any two concepts within the same graph component are related; any two concepts from two different components are unrelated. But it is intractable to use all possible concept pairs. For example, the largest component of the ConcreteGraph of the WORD dataset (Ein-Dor et al. 2018) has 4,301 nodes. We could sample up to 9,247,150 related concept pairs from it, which is impractical for a Transformer model to learn. Therefore, we must only use a subset of those concept pairs. An intuitive heuristic is to pick a concept pair with a short path length. If the path length is too long, the connection between the concept pair becomes “risky”. Namely, when there are more edges on the path, the probability of the existence of a noisy edge becomes higher, which could degrade the quality of the path. We later prove this theory with an empirical ablation study in Sub-section “Ablation Study of Data Augmentation”. In practice, we obtain the  $k$ -hop neighborhood for every node, where  $k$  is a small integer, such as 2 or 3. We use a subset of the nodes in this local neighborhood as related concepts by setting a target augmentation ratio. Unrelated concept pairs can be produced by simply sampling concepts from other components.

### Graph Component Contrastive Learning

We introduce Graph Component Contrastive Learning (GCCL) to take full advantage of local and global structural information of the ConcreteGraph, which is neglected by previous CRE methods. Our ConcreteGraph-based data augmentation only provides concept pairs from a local  $k$ -hop neighborhood, and it is intractable to include all concept pairs that satisfy the CRE properties. In fact, it is harmful to explicitly extract all possible concept pairs as augmented data due to the quality problem caused by long paths. GCCL captures global structural information by learning the representations of the components in the ConcreteGraph and

contrasting a concept’s own component against other components. In this way, all of the related and unrelated concept pairs are preserved implicitly because each component embedding aggregates a group of related concept embeddings. Our method also has a connection to clustering methods: if we add edges for all possible related pairs, the ConcreteGraph itself would become a cluster graph because all the components become complete subgraphs. For that reason, a ConcreteGraph component in GCCL is equivalent to a cluster (prototype) in Prototypical Contrastive Learning (PCL) (Li et al. 2021). The difference is that GCCL does not rely on clustering algorithms to find clusters.

The overview of our GCCL framework is illustrated in Figure 1(b). The main objective of GCCL is to treat the embedding of a concept’s own component embedding as the positive and the embeddings of all other components as negatives. This allows the model to learn representations for both individual concepts and their components, which implicitly encode the global graph structure. Formally, we assume that the ConcreteGraph has a set of  $r$  components  $C = \{c_1, c_2, \dots, c_r\}$ , where each component itself is a set of concepts  $c_s = \{x_{s_1}, x_{s_2}, \dots, x_{s_p}\}$  and  $p$  is the number of concepts in  $c_s$ . Same as PCL, we assume that  $f_\theta(x_i)$  follows an isotropic Gaussian distribution around the component embedding:

$$p(x_i; c_s, \theta) = \frac{\exp\left(-\left(f_\theta(x_i) - f_\theta(c_s)\right)^2 / 2\sigma_s^2\right)}{\sum_{j=1}^r \exp\left(-\left(f_\theta(x_i) - f_\theta(c_j)\right)^2 / 2\sigma_j^2\right)}, \quad (1)$$

where  $x_i \in c_s$  and  $\sigma^2$  is the variance. We also overload  $f_\theta$  with the input of a component, such that we get the representation of a component using a shared backbone encoder. In this paper, the component embedding is simply the element-wise average of the embeddings of its concepts:  $f_\theta(c_s) = \frac{1}{p} \sum_{k=1}^p f_\theta(x_{s_k})$ . Theoretically, we could add an extra Graph Neural Network (GNN) layer (Wu et al. 2021) on top of the concept embeddings to get the component embedding. However, because the components are complete subgraphs, GNN layers would quickly over-smooth the concept embeddings, which is close to their average. Thus, we do not use GNN layers to obtain the component embeddings.

After applying  $\ell_2$ -normalization to both the concept embeddings and the component embeddings, we have  $(f_\theta(x_i) - f_\theta(c_s))^2 = 2 - 2f_\theta(x_i) \cdot f_\theta(c_s)$ . Then the maximum log-likelihood estimation for  $\theta$  optimizes a novel Graph-Component-NCE (GC-NCE) loss  $\mathcal{L}_{GC-NCE}$ :

$$\begin{aligned} \theta^* &= \arg \max_{\theta} \sum_{i=1}^n \log p(x_i; C, \theta) \\ &= \arg \min_{\theta} \sum_{i=1}^n -\log \frac{\exp(f_\theta(x_i) \cdot f_\theta(c_s) / \phi_s)}{\sum_{j=1}^r \exp(f_\theta(x_i) \cdot f_\theta(c_j) / \phi_j)} \\ &\triangleq \arg \min_{\theta} \mathcal{L}_{GC-NCE}, \end{aligned} \quad (2)$$

where

$$\phi_s = \frac{\sum_{k=1}^p \|f_\theta(x_{s_k}) - f_\theta(c_s)\|_2}{p \log(p + \alpha)} \quad (3)$$

| Model                            | WORD          |               |               | CNSE          |               |               | CNSS          |               |               |
|----------------------------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|
|                                  | Acc           | F1            | AUC           | Acc           | F1            | AUC           | Acc           | F1            | AUC           |
| GCNII (Chen et al. 2020a)        | 54.62         | 44.55         | 57.03         | 62.81         | 56.22         | 66.40         | 58.40         | 61.06         | 63.00         |
| ShadDow-GNN (Zeng et al. 2021)   | 62.27         | 65.96         | 67.39         | 60.78         | 57.34         | 64.75         | 60.40         | 67.63         | 64.85         |
| JacobiConv (Wang and Zhang 2022) | 62.13         | 60.30         | 66.89         | 68.32         | 59.69         | 76.42         | 62.94         | 66.43         | 67.87         |
| GCN / CIG (Liu et al. 2019a)     | 62.81         | 67.75         | 67.00         | 83.95         | 82.05         | 91.92         | 89.69         | 90.05         | 95.97         |
| (Corrected CIG)                  | -             | -             | -             | (77.57)       | (75.13)       | -             | (82.87)       | (83.43)       | -             |
| BERT (Devlin et al. 2019)        | 76.85         | 76.12         | 84.86         | 83.97         | 81.67         | 92.10         | 90.63         | 91.09         | 97.07         |
| w. Aug (Ours)                    | 77.42         | 76.81         | 85.75         | 85.17         | 83.39         | 93.08         | 91.32         | 91.71         | 97.50         |
| w. GCCL (Ours)                   | 77.83         | 77.59         | 85.86         | <b>85.79</b>  | 84.18         | 92.96         | <b>93.67</b>  | <b>93.78</b>  | <b>98.37</b>  |
| w. Aug + GCCL (Ours)             | <b>78.18</b>  | <b>78.25</b>  | <b>*86.60</b> | 85.76         | <b>84.84</b>  | <b>93.46</b>  | 93.32         | 93.54         | 98.16         |
| RoBERTa (Liu et al. 2019b)       | 77.25         | <b>78.11</b>  | 85.55         | 84.52         | 82.79         | 92.88         | 92.84         | 92.74         | 98.20         |
| w. Aug (Ours)                    | 77.88         | 76.34         | 86.30         | 85.35         | 84.20         | 93.97         | 93.08         | 93.19         | 98.11         |
| w. GCCL (Ours)                   | <b>78.61</b>  | 77.93         | 86.26         | 86.21         | 84.40         | 93.94         | <b>93.82</b>  | <b>93.93</b>  | <b>98.50</b>  |
| w. Aug + GCCL (Ours)             | 78.14         | 76.69         | <b>86.45</b>  | <b>87.00</b>  | <b>86.05</b>  | <b>94.39</b>  | 93.46         | 93.62         | 98.05         |
| XLNet (Yang et al. 2019)         | 77.18         | <b>*78.42</b> | 86.10         | 87.00         | 84.65         | 95.21         | 93.70         | 93.84         | 98.53         |
| w. Aug (Ours)                    | 77.75         | 76.51         | 86.00         | 88.41         | 86.87         | 95.29         | 93.79         | 93.84         | 98.40         |
| w. GCCL (Ours)                   | 77.98         | 77.56         | 85.72         | 88.06         | 86.47         | 95.21         | <b>*94.60</b> | <b>*94.74</b> | <b>*98.80</b> |
| w. Aug + GCCL (Ours)             | <b>*79.15</b> | 78.32         | <b>86.25</b>  | <b>*88.82</b> | <b>*87.70</b> | <b>*95.53</b> | 94.15         | 94.34         | 98.68         |

Table 2: Performance comparison with various baselines for our ConcreteGraph-based data augmentation method and Graph Component Contrastive Learning (GCCL) framework. The best results are highlighted with bold text for each model-dataset combination; the best results for each of the three datasets are marked with \*.

is a concentration estimation that replaces the variance in Eq. 1, which is also introduced in PCL;  $\alpha$  is a smoothing hyper-parameter for preventing  $\phi$  from becoming excessively large when the component size  $p$  is small. When  $\phi$  is small, the concept embeddings within the same component are pulled closer together and, thus, have a high concentration level.

In addition to the main GC-NCE objective of GCCL, we also use the InfoNCE (van den Oord, Li, and Vinyals 2018) objective under the MoCo framework (He et al. 2020; Chen et al. 2020b). The idea of InfoNCE is to distinguish the embedding of a concept from the embeddings of a set of noise concepts. MoCo introduces a dynamic dictionary based on a queue data structure for storing the embeddings from a momentum encoder  $f_{\theta_k}$  in a first-in-first-out (FIFO) fashion, which enables the decoupling between the batch size and the amount of negatives. This is especially useful since we use large Transformer models as encoders. The MoCo objective is constructed by contrasting the query embedding  $f_{\theta}(x)$  from the main encoder and the key embeddings  $f_{\theta_k}(x)$  from the momentum encoder. The momentum encoder  $f_{\theta_k}$  is the moving average of the main encoder  $f_{\theta}$  with a momentum coefficient  $m \in [0, 1]$ :  $\theta_k \leftarrow m\theta_k + (1 - m)\theta_q$ .

Similar to InfoNCE, the MoCo loss treats the query and key embeddings of a concept as the positive pair, and the key embeddings of other concepts in the MoCo queue are negatives to the query embedding of the current concept. That leads to the MoCo objective:

$$\mathcal{L}_{\text{MoCo}} = \sum_{i=1}^n -\log \frac{\exp(f_{\theta}(x_i) \cdot f_{\theta_k}(x_i)/\tau)}{\sum_{j=0}^Q \exp(f_{\theta}(x_i) \cdot f_{\theta_k}(x_j)/\tau)}, \quad (4)$$

where the denominator includes one positive and  $Q$  negatives in the MoCo queue;  $\tau$  is a temperature hyper-parameter. The purpose of adding this additional MoCo objective is to preserve local smoothness of the loss function

(Li et al. 2021): the GC-NCE objective (Eq. 2) and the CRE objective (Eq. 5) both pull the embeddings of related concepts closer, those concept embeddings might become very similar after finetuning, which can cause poor generalization; the MoCo objective preserves the distinction of individual concept embeddings by contrasting a concept against all other concepts.

In this paper, we use Transformer (Vaswani et al. 2017) as the encoder  $f_{\theta}$ . After passing the concepts into  $f_{\theta}$  and obtaining their concept embeddings, the concept relatedness is estimated as follows:

$$\begin{aligned} h(x_i, x_j; \theta) &= f_{\theta}(x_i) \oplus f_{\theta}(x_j) \oplus |f_{\theta}(x_i) - f_{\theta}(x_j)| \\ p(x_i, x_j) &= \sigma(\mathbf{W}h(x_i, x_j; \theta)) \\ \mathcal{L}_{\text{CRE}} &= L_{\text{CE}}(p(x_i, x_j), y_{i,j}), \end{aligned} \quad (5)$$

where  $\oplus$  is the concatenation operation;  $|f_{\theta}(x_i) - f_{\theta}(x_j)|$  is the element-wise absolute value of the difference of the two vectors;  $\sigma$  is a sigmoid activation function;  $\mathbf{W}$  is a trainable parameter matrix;  $L_{\text{CE}}$  is the binary cross-entropy loss.

The overall objective  $\mathcal{L}$  combines the GC-NCE loss (Eq. 2), the MoCo loss (Eq. 4) and the CRE loss (Eq. 5):

$$\mathcal{L} = \mathcal{L}_{\text{CRE}} + \beta \mathcal{L}_{\text{GCCL}} = \mathcal{L}_{\text{CRE}} + \beta(\mathcal{L}_{\text{GC-NCE}} + \mathcal{L}_{\text{MoCo}}),$$

where  $\beta$  is the hyper-parameter that adjusts the effect of the GCCL loss  $\mathcal{L}_{\text{GCCL}}$ .

## Experiments

We conduct experiments on three CRE datasets: WORD (Ein-Dor et al. 2018), CNSE and CNSS (Liu et al. 2019a). We compare the performance of three Transformer models, including BERT (Devlin et al. 2019), RoBERTa (Liu et al. 2019b) and XLNet (Yang et al. 2019). We also include four latest Graph Neural Network (GNN) models (Chen et al. 2020a; Zeng et al. 2021; Wang and Zhang 2022; Liu et al.

| Objective                            | WORD         |              |              | CNSE         |              |              | CNSS         |              |              |
|--------------------------------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|
|                                      | Acc          | F1           | AUC          | Acc          | F1           | AUC          | Acc          | F1           | AUC          |
| BERT (Devlin et al. 2019) (CRE only) | 76.85        | 76.12        | 84.86        | 83.97        | 81.67        | 92.10        | 90.63        | 91.09        | 97.07        |
| w. MoCo                              | 77.10        | 76.79        | 84.82        | 84.07        | 81.78        | 92.25        | 92.57        | 92.88        | 97.96        |
| w. GC-NCE                            | 77.58        | <b>78.06</b> | 85.81        | 84.93        | 83.17        | 92.68        | 92.93        | 93.06        | 98.15        |
| w. GCCL (GC-NCE & MoCo)              | <b>77.83</b> | 77.59        | <b>85.86</b> | <b>85.79</b> | <b>84.18</b> | <b>92.96</b> | <b>93.67</b> | <b>93.78</b> | <b>98.37</b> |

Table 3: The ablation study of the GCCL objective.

2019a) as baselines. Detailed ablation studies show the effect of different modules of our method.

## Datasets

The Wikipedia Oriented Relatedness Dataset (WORD) dataset (Ein-Dor et al. 2018) collects English concepts from Wikipedia, including 19,176 concept pairs. The Chinese News Same Event (CNSE) dataset and the Chinese News Same Story (CNSS) dataset were both created by Liu et al. (2019a), containing 29,063 and 33,503 news article pairs from the major internet news providers in China, respectively. We use the official dataset split for WORD whose train-test ratio is approximately 2:1. Since CNSE and CNSS do not provide an official dataset split, they are split randomly with a train-dev-test ratio of 7:2:1. These dataset splits are fixed throughout different experiments across all models.

## Experimental Setup

We use the AdamW optimizer (Loshchilov and Hutter 2019) with learning rate  $= 1 \times 10^{-5}$  and  $\epsilon = 1 \times 10^{-8}$ , following a linear schedule. The Transformer models are trained for 5 epochs. For GC-NCE, we use  $\alpha = 10$ . For MoCo (He et al. 2020; Chen et al. 2020b), we use queue size  $Q = 32$ , momentum coefficient  $m = 1 - 1 \times 10^{-4}$ , and temperature  $\tau = 0.1$ . We use  $\beta = 0.1$  for the overall loss. CIG (Liu et al. 2019a) is trained with its default hyper-parameters. Experiments are conducted on four Nvidia TITAN V GPUs.

## Performance Comparison

The performance comparison of different models is summarized in Table 2. Three performance metrics are used: accuracy, F1, and AUC. We compare our method with different GNNs and Transformers.

We use four recent GNNs as baselines. For GCNII (Chen et al. 2020a), ShaDow-GNN (Zeng et al. 2021) and Jacobi-Conv (Wang and Zhang 2022), we build a complete graph for each concept, where each node is a token and the node features are BERT token embeddings. For CNSE and CNSS, Concept Interaction Graph (CIG) (Liu et al. 2019a) is the current state-of-the-art model. CIG extracts a concept graph for each article, where its concepts are different from CRE concepts in that they are more similar to named entities. One major problem with this model is that the size of the concept graphs is limited. If a concept graph exceeds the limit, the model simply discards the article pair. Therefore, CIG practically uses an easier subset of the original dataset, causing inaccurate measurement of its performance. Thus, we add a ‘‘Corrected CIG’’ row in the table to include the corrected

accuracy and F1. Because CIG is based on a graph convolutional network (GCN) (Fey and Lenssen 2019) but it only works on CNSE and CNSS, we developed a similar GCN based on ELMo (Peters et al. 2018) for the WORD dataset. Because most GNNs are designed to process sparse graph inputs and they can suffer from the over-smoothing issue, they usually do not excel at processing text inputs.

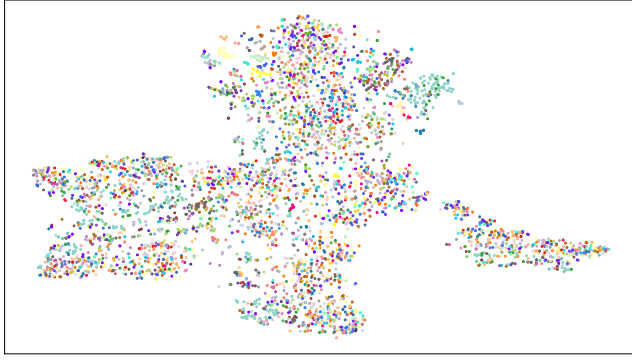
We finetune three Transformer models, BERT (Devlin et al. 2019), RoBERTa (Liu et al. 2019b), and XLNet (Yang et al. 2019) on the three datasets. GCCL is independent of the ConcreteGraph-based data augmentation method. Thus, they can be used either separately or simultaneously. To illustrate their benefits, we report the performance of the base Transformer model, Transformer with data augmentation (‘‘w. Aug’’), Transformer with GCCL (‘‘w. GCCL’’), and Transformer with both modules (‘‘w. Aug + GCCL’’).

When compared with GNNs, Transformer models exhibit a large advantage because the self-attention mechanism automatically optimizes the attention weights between tokens, whereas the adjacency matrices in GNNs are fixed. In addition, the base Transformer models can even outperform the uncorrected CIG and, with the help of our GCCL framework, the margin is enlarged. When comparing our methods with the base Transformer models, we can see that the GCCL framework is able to bring more improvement than our data augmentation method in most cases. This shows that the global structural information from GCCL is indeed more effective for the CRE task.

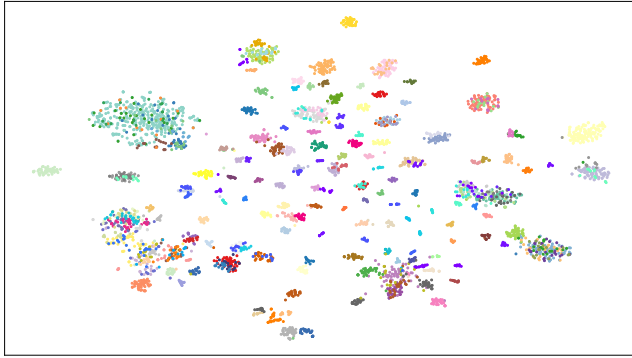
When ‘‘Aug’’ and ‘‘GCCL’’ are combined, the performance is further enhanced in most metrics on WORD and CNSE; however, applying only GCCL produces better results on the CNSS dataset. ConcreteGraph-based data augmentation is confined to a local  $k$ -hop neighborhood, while GCCL captures the information of the entire component. They can complement each other in most cases but their difference may also cause disagreement as to whether two concepts are related if they are connected by a long path. That is, data augmentation does not sample any concept pairs with path lengths longer than  $k$ , but GCCL retains those connections. Such contradiction sometimes worsens the performance.

## Ablation Study of GCCL

The GCCL loss has two parts: the main GC-NCE loss and the auxiliary MoCo loss. We conduct an ablation study on their individual effect on BERT (Devlin et al. 2019). As shown in Table 3, both the MoCo loss and the GC-NCE loss can improve the model performance. In addition, the main GC-NCE objective produces more performance improvement than MoCo does. When they are combined, the



(a) Original Representations



(b) GCCL Representations

Figure 2: The t-SNE visualization of the original and GCCL concept embeddings. Each point corresponds to a concept embedding, and the concepts in the same component share the same color. Both sub-figures show 200 Components.

performance improvement from the two objectives stacks.

In Figure 2, we also visualize the concept embeddings with t-SNE (Van der Maaten and Hinton 2008) to qualitatively illustrate the effect of GCCL. GCCL better represents the concepts by taking the global structure of the ConcreteGraph into account: the embeddings of the concepts in the same component are closer to each other using GCCL and the embedding space of GCCL is better organized. This gives us an insight into how GCCL helps the main CRE objective: GCCL pulls related concepts together and pushes away unrelated components.

### Ablation Study of Data Augmentation

To understand whether the quality or the quantity of augmented data better helps the model performance, we experiment with different  $k$  values and target augmentation ratios, as summarized in Table 4. When the target augmentation ratio is specified, the sampling process of new related concept pairs stops when the target augmentation ratio is reached, regardless of whether all possible related concept pairs are sampled. If there is no target augmentation ratio ( $\infty$ ), we find all possible related concept pairs for the given  $k$ . In

| $k$ | Target Ratio | Related Pairs Samp. | Pairs Poss. | Acc          | F1           | AUC          |
|-----|--------------|---------------------|-------------|--------------|--------------|--------------|
| 1   | -            | 6.8K                | 6.8K        | 76.85        | 76.12        | 84.86        |
| 2   | 2            | 12.9K               | 107K        | <b>77.42</b> | <b>76.81</b> | <b>85.75</b> |
| 2   | 3            | 19.3K               | 107K        | 75.66        | 74.51        | 83.96        |
|     | 4            | 25.8K               | 107K        | 72.23        | 72.07        | 79.99        |
|     | 5            | 32.2K               | 107K        | 70.99        | 70.78        | 78.95        |
| 2   | $\infty$     | 107K                | 107K        | 63.51        | 58.65        | 70.97        |
| 3   |              | 401K                | 401K        | 56.76        | 51.38        | 63.69        |
| 3   |              | 12.9K               | 401K        | 75.91        | 75.81        | 83.78        |
| 4   | 2            | 12.9K               | 1.4M        | 74.90        | 75.44        | 83.08        |
| 5   |              | 12.9K               | 3.9M        | 74.06        | 75.17        | 81.56        |

Table 4: The effect of the number of hops  $k$  and the augmentation ratio in the ConcreteGraph-based data augmentation method. ‘‘Samp.’’ means sampled related pairs for training and ‘‘Poss.’’ means all possible related pairs in the  $k$ -hop neighborhood. The number of sampled unrelated pairs is equal to the number of sampled related pairs and omitted.

both cases, we keep sampling unrelated concept pairs until the number of unrelated pairs is equal to that of related pairs. This ensures the balance of the augmented dataset.

When  $k$  is fixed at 2 and the augmentation ratio increases, the performance worsens. If we take it to the extreme by not setting a target ratio, the performance drops significantly. These results show that a large augmentation ratio does not always help because it shrinks the proportion of annotated concept pairs and the quality of sampled concept pairs never matches that of annotated concept pairs. When the augmentation ratio is fixed at 2 and  $k$  is incremented, the performance also declines. This confirms our theory about how path length can affect the quality of sampled concept pairs. Namely, when the path length between two concepts is long, the existence of a noisy edge becomes more likely, which degrades the overall quality of the connection.

## Conclusion

Concept relatedness estimation is an emerging task that has a wide range of applications. To the best of our knowledge, we are the first to discover the concept-level graph structure that unveils the high-order relationship among concepts. We name it ConcreteGraph and develop a novel data augmentation method based on it. However, because data augmentation cannot capture the global structural information of the ConcreteGraph, we integrate a novel Graph Component Contrastive Learning (GCCL) framework to encode the complete graph structure. Experimental results show that the data augmentation method can improve the performance by  $\sim 1\%$ , whereas the improvement from GCCL can be up to  $\sim 3\%$ . Detailed analysis shows that GCCL better organizes the embedding space where related concepts are close while unrelated concepts are separated. We also conduct an experiment to show that it is difficult for the ConcreteGraph-based data augmentation method to provide as much performance benefit as GCCL by simply sampling more concept pairs, due to the degradation of the data quality caused by the long path length.

## Acknowledgements

The work described here was partially supported by grants from the National Key Research and Development Program of China (No. 2018AAA0100204) and from the Research Grants Council of the Hong Kong Special Administrative Region, China (CUHK 2410021, Research Impact Fund, No. R5034-18).

## References

- Aswani Kumar, C.; and Srinivas, S. 2010. Concept lattice reduction using fuzzy K-Means clustering. *Expert Systems with Applications*, 37(3): 2696–2704.
- Busch, M.; Gade, K.; Larson, B.; Lok, P.; Luckenbill, S.; and Lin, J. 2012. Earlybird: Real-Time Search at Twitter. In *2012 IEEE 28th International Conference on Data Engineering*, 1360–1369.
- Cer, D.; Diab, M.; Agirre, E.; Lopez-Gazpio, I.; and Specia, L. 2017. Semeval-2017 task 1: Semantic textual similarity-multilingual and cross-lingual focused evaluation. *arXiv preprint arXiv:1708.00055*.
- Chen, M.; Wei, Z.; Huang, Z.; Ding, B.; and Li, Y. 2020a. Simple and Deep Graph Convolutional Networks. In *ICML*, volume 119 of *Proceedings of Machine Learning Research*, 1725–1735. PMLR.
- Chen, X.; Fan, H.; Girshick, R. B.; and He, K. 2020b. Improved Baselines with Momentum Contrastive Learning. *CoRR*, abs/2003.04297.
- Chopra, S.; Hadsell, R.; and LeCun, Y. 2005. Learning a Similarity Metric Discriminatively, with Application to Face Verification. In *CVPR (1)*, 539–546. IEEE Computer Society.
- Devlin, J.; Chang, M.; Lee, K.; and Toutanova, K. 2019. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *NAACL-HLT (1)*, 4171–4186. Association for Computational Linguistics.
- Ein-Dor, L.; Halfon, A.; Kantor, Y.; Levy, R.; Mass, Y.; Rinott, R.; Shnarch, E.; and Slonim, N. 2018. Semantic Relatedness of Wikipedia Concepts - Benchmark Data and a Working Solution. In *LREC*. European Language Resources Association (ELRA).
- Fey, M.; and Lenssen, J. E. 2019. Fast Graph Representation Learning with PyTorch Geometric. In *ICLR Workshop on Representation Learning on Graphs and Manifolds*.
- Formica, A. 2006. Ontology-based concept similarity in formal concept analysis. *Information sciences*, 176(18): 2624–2641.
- Formica, A. 2008. Concept similarity in formal concept analysis: An information content approach. *Knowledge-based systems*, 21(1): 80–87.
- Gao, Y.; Wu, C.-S.; Li, J.; Joty, S.; Hoi, S. C.; Xiong, C.; King, I.; and Lyu, M. 2020. Discern: Discourse-Aware Entailment Reasoning Network for Conversational Machine Reading. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2439–2449.
- Ge, J.; and Qiu, Y. 2008. Concept Similarity Matching Based on Semantic Distance. In *SKG*, 380–383. IEEE Computer Society.
- He, K.; Fan, H.; Wu, Y.; Xie, S.; and Girshick, R. B. 2020. Momentum Contrast for Unsupervised Visual Representation Learning. In *CVPR*, 9726–9735. Computer Vision Foundation / IEEE.
- Hu, X.; Ma, F.; Liu, C.; Zhang, C.; Wen, L.; and Yu, P. S. 2021a. Semi-supervised Relation Extraction via Incremental Meta Self-Training. In *Proc. of EMNLP: Findings*.
- Hu, X.; Wen, L.; Xu, Y.; Zhang, C.; and Yu, P. S. 2020. Self-ORE: Self-supervised Relational Feature Learning for Open Relation Extraction. In *Proc. of EMNLP*, 3673–3682.
- Hu, X.; Zhang, C.; Yang, Y.; Li, X.; Lin, L.; Wen, L.; and Yu, P. S. 2021b. Gradient Imitation Reinforcement Learning for Low Resource Relation Extraction. In *Proc. of EMNLP*, 2737–2746.
- Jiang, J.-Y.; Zhang, M.; Li, C.; Bendersky, M.; Golbandi, N.; and Najork, M. 2019. Semantic Text Matching for Long-Form Documents. In *The World Wide Web Conference, WWW '19*, 795–806. New York, NY, USA: Association for Computing Machinery. ISBN 9781450366748.
- Li, J.; Gao, Y.; Bing, L.; King, I.; and Lyu, M. R. 2019. Improving Question Generation With to the Point Context. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, 3216–3226.
- Li, J.; Li, Z.; Ge, T.; King, I.; and Lyu, M. R. 2022. Text Revision by On-the-Fly Representation Optimization. In *Thirty-Sixth AAAI Conference on Artificial Intelligence*, 10956–10964.
- Li, J.; Li, Z.; Mou, L.; Jiang, X.; Lyu, M.; and King, I. 2020. Unsupervised text generation by learning from search. *Advances in Neural Information Processing Systems*, 33: 10820–10831.
- Li, J.; Zhou, P.; Xiong, C.; and Hoi, S. C. H. 2021. Prototypical Contrastive Learning of Unsupervised Representations. In *ICLR*. OpenReview.net.
- Li, W.; and Xia, Q. 2011. A method of concept similarity computation based on semantic distance. *Procedia Engineering*, 15: 3854–3859.
- Liu, B.; Niu, D.; Wei, H.; Lin, J.; He, Y.; Lai, K.; and Xu, Y. 2019a. Matching Article Pairs with Graphical Decomposition and Convolutions. In *ACL (1)*, 6284–6294. Association for Computational Linguistics.
- Liu, S.; Hu, X.; Zhang, C.; Li, S.; Wen, L.; and Yu, P. S. 2022. HiURE: Hierarchical Exemplar Contrastive Learning for Unsupervised Relation Extraction. In *Proc. of NAACL*.
- Liu, Y.; Ott, M.; Goyal, N.; Du, J.; Joshi, M.; Chen, D.; Levy, O.; Lewis, M.; Zettlemoyer, L.; and Stoyanov, V. 2019b. RoBERTa: A Robustly Optimized BERT Pretraining Approach. *CoRR*, abs/1907.11692.
- Lombardi, L.; and Sartori, G. 2006. Concept similarity: An abstract relevance classes approach. In *The 7th International Conference on Cognitive Modeling, Trieste*, 190–195.

- Loshchilov, I.; and Hutter, F. 2019. Decoupled Weight Decay Regularization. In *ICLR (Poster)*. OpenReview.net.
- Muangprathub, J.; Kajornkasirat, S.; and Wanichsombat, A. 2021. Document Plagiarism Detection Using a New Concept Similarity in Formal Concept Analysis. *Journal of Applied Mathematics*, 2021: 1–10.
- Peters, M. E.; Neumann, M.; Iyyer, M.; Gardner, M.; Clark, C.; Lee, K.; and Zettlemoyer, L. 2018. Deep Contextualized Word Representations. In *NAACL-HLT*, 2227–2237. Association for Computational Linguistics.
- Song, Z.; and King, I. 2022. Hierarchical Heterogeneous Graph Attention Network for Syntax-Aware Summarization. In *AAAI*, 11340–11348. AAAI Press.
- Song, Z.; Meng, Z.; Zhang, Y.; and King, I. 2021. Semi-supervised Multi-label Learning for Graph-structured Data. In *CIKM*, 1723–1733. ACM.
- Song, Z.; Yang, X.; Xu, Z.; and King, I. 2022. Graph-Based Semi-Supervised Learning: A Comprehensive Review. *IEEE Transactions on Neural Networks and Learning Systems*, 1–21.
- Song, Z.; Zhang, Y.; and King, I. 2022. Towards an Optimal Asymmetric Graph Structure for Robust Semi-supervised Node Classification. In *KDD*, 1656–1665. ACM.
- Sun, X.; Ge, T.; Ma, S.; Li, J.; Wei, F.; and Wang, H. 2022. A Unified Strategy for Multilingual Grammatical Error Correction with Pre-trained Cross-Lingual Language Model. *arXiv preprint arXiv:2201.10707*.
- Teevan, J.; Ramage, D.; and Morris, M. 2011. #TwitterSearch: a comparison of microblog search and web search. In *WSDM '11*.
- Thijs, B. 2019. Paragraph-based intra- and inter- document similarity using neural vector paragraph embeddings. In *ISSI*.
- Tsatsaronis, G.; Varlamis, I.; and Vazirgiannis, M. 2014. Text Relatedness Based on a Word Thesaurus. *CoRR*, abs/1401.5699.
- van den Oord, A.; Li, Y.; and Vinyals, O. 2018. Representation Learning with Contrastive Predictive Coding. *CoRR*, abs/1807.03748.
- Van der Maaten, L.; and Hinton, G. 2008. Visualizing data using t-SNE. *Journal of machine learning research*, 9(11).
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A. N.; Kaiser, L.; and Polosukhin, I. 2017. Attention is All you Need. In *NIPS*, 5998–6008.
- Wang, L.; and Liu, X. 2008. A new model of evaluating concept similarity. *Knowledge-Based Systems*, 21(8): 842–846.
- Wang, X.; and Zhang, M. 2022. How Powerful are Spectral Graph Neural Networks. In *ICML*, volume 162 of *Proceedings of Machine Learning Research*, 23341–23362. PMLR.
- Wu, Z.; Pan, S.; Chen, F.; Long, G.; Zhang, C.; and Yu, P. S. 2021. A Comprehensive Survey on Graph Neural Networks. *IEEE Trans. Neural Networks Learn. Syst.*, 32(1): 4–24.
- Yang, Z.; Dai, Z.; Yang, Y.; Carbonell, J. G.; Salakhutdinov, R.; and Le, Q. V. 2019. XLNet: Generalized Autoregressive Pretraining for Language Understanding. In *NeurIPS*, 5754–5764.
- You, Y.; Chen, T.; Sui, Y.; Chen, T.; Wang, Z.; and Shen, Y. 2020. Graph Contrastive Learning with Augmentations. In *NeurIPS*.
- Zeng, H.; Zhang, M.; Xia, Y.; Srivastava, A.; Malevich, A.; Kannan, R.; Prasanna, V. K.; Jin, L.; and Chen, R. 2021. Decoupling the Depth and Scope of Graph Neural Networks. In *NeurIPS*, 19665–19679.
- Zhang, Y.; and Zhu, H. 2019. Doc2hash: Learning discrete latent variables for documents retrieval. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, 2235–2240.
- Zhang, Y.; and Zhu, H. 2020. Discrete Wasserstein Autoencoders for Document Retrieval. In *ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 8159–8163. IEEE.
- Zhang, Y.; Zhu, H.; Meng, Z.; Koniusz, P.; and King, I. 2022a. Graph-adaptive rectified linear unit for graph neural networks. In *Proceedings of the ACM Web Conference 2022*, 1331–1339.
- Zhang, Y.; Zhu, H.; Song, Z.; Koniusz, P.; and King, I. 2022b. COSTA: Covariance-Preserving Feature Augmentation for Graph Contrastive Learning. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, 2524–2534.