

Learning to Complete Knowledge Graphs with Deep Sequential Models (Postprint)

Authors: Guo, Lingbing, Zhang, Qingheng, Hu, Wei, Sun, Zequn, Qu, Yuzhong, Hu, Wei

Date: 2022-11-27T00:00:00+00:00

Abstract

Knowledge graph (KG) completion aims at filling the missing facts in a KG, where a fact is typically represented as a triple in the form of (head, relation, tail). Traditional KG completion methods compel two-thirds of a triple provided (e.g., head and relation) to predict the remaining one. In this paper, we propose a new method that extends multi-layer recurrent neural networks (RNNs) to model triples in a KG as sequences. It obtains state-of-the-art performance on the common entity prediction task, i.e., giving head (or tail) and relation to predict the tail (or the head), using two benchmark data sets. Furthermore, the deep sequential characteristic of our method enables it to predict the relations given head (or tail) only, and even predict the whole triples. Our experiments on these two new KG completion tasks demonstrate that our method achieves superior performance compared with several alternative methods.

Full Text

Preamble

RESEARCH PAPER

Learning to Complete Knowledge Graphs with Deep Sequential Models

Lingbing Guo, Qingheng Zhang, Wei Hu[†], Zequn Sun & Yuzhong Qu
State Key Laboratory for Novel Software Technology, Nanjing University, Nanjing 210023, China

Keywords: Knowledge graph; Entity prediction; Triple prediction; Recurrent neural network

Citation: L. Guo, Q. Zhang, W. Hu, Z. Sun & Y. Qu. Learning to complete knowledge graphs with deep sequential models. Data Intelligence 1(2019), 224-243. doi: 10.1162/dint_a_00016}

Received: December 19, 2018; **Revised:** April 22, 2019; **Accepted:** May 5, 2019

ABSTRACT

Knowledge graph (KG) completion aims at filling the missing facts in a KG, where a fact is typically represented as a triple in the form of (head, relation, tail). Traditional KG completion methods compel two-thirds of a triple provided (e.g., head and relation) to predict the remaining one. In this paper, we propose a new method that extends multi-layer recurrent neural networks (RNNs) to model triples in a KG as sequences.

It obtains state-of-the-art performance on the common entity prediction task, i.e., giving head (or tail) and relation to predict the tail (or the head), using two benchmark data sets. Furthermore, the deep sequential characteristic of our method enables it to predict the relations given head (or tail) only, and even predict the whole triples. Our experiments on these two new KG completion tasks demonstrate that our method achieves superior performance compared with several alternative methods.

1. INTRODUCTION

Knowledge graphs (KGs), such as DBpedia [1] and Freebase [2], often use triples, in the form of (h, r, t), to record billions of real-world facts, where h, t denote entities and r denotes a relation between h and t. Despite the enormous effort put into KG creation and maintenance, current KGs are still far from complete.

KG completion is proposed to deal with this problem. Previous methods focus on a general task called entity prediction (also known as link prediction) [3, 4], which requires one to complete a triple in a KG by predicting t given (h, r, ?) or predicting h given (?, r, t). The left-most part of Figure 1 [Figure 1: see original paper] shows a commonly-used model for tail entity prediction. Input h, r is firstly projected by some vectors or matrices, where the bold letters denote the corresponding embeddings of h and r, and then combined to a continuous representation o_t for predicting t.

Although existing methods have shown good performance on entity prediction, in real world they may still be inadequate to complete a KG. Let us assume that there is a method which can effectively complete an entity h given a relation r explicitly. But if we do not provide any relations, this method would be incompetent to enrich h, since it is incapable of knowing which relation should be used to complete this entity. An alternative way is to iterate over all relations in the KG, but such process is time-consuming and error-prone. We argue that this may prevent the existing methods from being useful in the real world.

The recurrent neural network (RNN) is a neural sequence model, which has achieved state-of-the-art performance on many natural language processing (NLP) tasks, such as language modeling, speech recognition and machine

translation [5, 6]. A triple in a KG can be considered to be a sequence of length 3, which inspires us to use RNNs to model KGs. However, we are still challenged by the following two obstacles: (i) triples are not natural language. They model complex structures of KGs with a fixed expression (h, r, t). Such short sequences may be insufficient to provide adequate context for prediction, while it is time-consuming and difficult to construct valuable long sequences from a huge number of paths in KGs; and (ii) relations and entities are elements of two different types, and they appear in triples in a fixed order. It is over-simplified to treat them as the same type.

In this paper, we propose a new method, which employs RNNs to model triples in a KG as sequences. We first design BSKG, a basic sequential model for KGs, as an initial version to illustrate our method (Figure 1). We denote an RNN cell by c , which receives its previous hidden state and the current element as input to predict the next. The cell in the entity layer processes entity embeddings like h , while the cell in the relation layer processes relation embeddings like r . In BSKG, we use one cell to sequentially process all input elements, and thus h , r are fed to the same cell c to obtain their respective output. Then, we use o_h to predict relations of h and o_r to predict tails of $h \rightarrow r$.

However, BSKG lacks effectiveness to model complex structures, because it only uses a single RNN cell to process all input sequences. To improve the performance on complex KGs, a few existing methods like TransH [7] and TransR [8] suggest projecting entities by some vectors or matrices before combination. Different from them, we do not choose to extend BSKG by adding such a layer, since it requires creation of variables for each relation, which consumes considerable resources and may also make the models hard to converge. Instead, we propose DSKG in this paper, a deep sequential model that extends multi-layer RNNs to model KGs. DSKG can smoothly model complex structures, since each cell in DSKG does not need to output an explicit prediction result, but gives its own comprehension by adding or removing information from the hidden state and conveys this processed hidden state to its next cell. Therefore, the original input is continually refined by each cell, and the complex features can be fluently processed.

Figure 1. Different models for tail entity prediction. Note: White and black circles denote input and output, respectively. For BSKG and DSKG, light gray circles denote cells in the entity layer and dark gray circles denote cells in the relation layer.

Specifically, we design three different strategies to integrate the RNN cells in DSKG. The right-most part of Figure 1 shows their two-layer examples.

The **Joint** strategy is the most prevalent way used in the NLP area. RNN cells are reused in different layers like BSKG, but the output hidden state of each cell is not only recurrently fed to itself next time, but also sequentially conveyed to its next cell.

The **Respective** strategy uses independent RNN cells for the entity layer and

the relation layer, that is to say, c_1 , c_2 , c_3 and c_4 are all different. We want this strategy to achieve better performance when relations are diverse and complex.

The **Cascade** strategy also uses different RNN cells for the entity layer and the relation layer. Instead of parallel connecting the two layers, we decide to link the end cell of the entity layer to the first cell of the relation layer. As a result, each cell in this strategy can concentrate more on its current work.

The main contributions of this paper are listed as follows:

- We introduce a new method for KG completion, which extends multi-layer RNNs to model triples in a KG as sequences of length 3. Three distinct strategies are proposed to integrate RNN cells and show their different characteristics in our experiments.
- We design two new KG completion tasks, namely relation prediction and triple prediction, as complements to the entity prediction task. The relation prediction task aims to predict relations only given a head (or tail) entity as input. The triple prediction task is to predict the whole triples only given a head entity. Note that the relation prediction task is different from the task to predict relations given a head entity and a tail entity. The latter one is sometimes called relationship prediction.
- Our experimental results show that our method achieves state-of-the-art performance for entity prediction on the benchmark data sets based on Freebase and WordNet. It also achieves promising results on the new relation prediction and triple prediction data sets.

The rest of this paper is organized as follows. Section 2 summarizes the related work. Section 3 describes the details of our method. Section 4 presents the experimental results. Finally, we conclude this paper in Section 5.

2. RELATED WORK

We divide existing models into two sub-areas: translational models and non-translational models. We summarize several representative models in Table 1. Our models are also added for comparison.

Table 1. Comparison of existing models and ours.

Models	Preprepared	Energy functions	Notation explanations
TransE	embeddings	$\ h + r - t\ _2^2$	$\mathbf{h}, \mathbf{r}, \mathbf{t}$: embeddings
TransH	embeddings	$\ h_{\perp} + r - t_{\perp}\ _2^2$	$\mathbf{w}_r$: relation-specific normalization vector
TransR	embeddings	$\ h_r + r - t_r\ _2^2$	$\mathbf{W}_r$: relation-specific matrix
STransE	embeddings	$\ \mathbf{W}_{r,1}h + r - \mathbf{W}_{r,2}t\ _2^2$	$\mathbf{W}_{r,1}, \mathbf{W}_{r,2}$: relation-specific matrices

Models	Preprepared	Energy functions	Notation explanations
PTransE	embeddings; paths	$\ h + r - t\ _2^2 + \ h +$ $P(h, t) - t\ _2^2$	$P(h, t)$: path set of ($h, \dots, t$)
DISTMULTembeddings		$\mathbf{r}^T(\mathbf{h} \odot \mathbf{t})$	$\mathbf{W}_r$: relation-specific diagonal matrix
NLFeat	Node and link features	$\tanh(\mathbf{h}^T \mathbf{W}_r \mathbf{t})$	$\mathbf{i}, \mathbf{j}, \mathbf{k}$: feature vector; $\mathbf{H}$: weight vector
ConvE	embeddings	$g(\text{concat}(g([\mathbf{h}; \mathbf{r}] * \mathbf{w})) \mathbf{W}) \mathbf{t}$	g : activation function; $\mathbf{w}$: convolutional layer
ConvKB	embeddings	$g(\text{concat}(g([\mathbf{h}; \mathbf{r}; \mathbf{t}] * \mathbf{w})) \mathbf{W})$	g : activation function; $\mathbf{w}$: convolutional layer
DSKG (Joint)	embeddings	$\text{softmax}(\mathbf{W}_2 \mathbf{o}_t)$	$\mathbf{o}_t$: RNN output
DSKG (Respec- tive)	embeddings	$\text{softmax}(\mathbf{W}_2 \mathbf{o}_t)$	$\mathbf{o}_t$: RNN output
DSKG (Cas- cade)	embeddings	$\text{softmax}(\mathbf{W}_2 \mathbf{o}_t)$	$\mathbf{o}_t$: RNN output

2.1 Translational Models

Perhaps, TransE [4] is the most well-known embedding model for KG completion. It represents entities and relations as k-dimensional vectors in a unified space, and models a triple (h, r, t) as $\mathbf{h} + \mathbf{r} \approx \mathbf{t}$. TransE works well for one-to-one relationship, but it fails to model more complex (e.g., one-to-many) relationships.

TransH [7] resolves this problem by regarding each relation r as a vector on a hyperplane whose normalization vector is $\mathbf{w}_r$. It projects entity embeddings $\mathbf{h}, \mathbf{t}$ to this hyperplane by $\mathbf{w}_r$, and uses the same energy function as TransE for training. TransR [8] uses relation-specific matrices to project entities. It creates a project matrix $\mathbf{W}_r$ for each relation r and projects $\mathbf{h}, \mathbf{t}$ by $\mathbf{W}_r$. TransR also employs the same energy function for training.

To extend the above models, STransE [9] learns two project matrices $\mathbf{W}_{r,1}$ and $\mathbf{W}_{r,2}$ for each relation r , where $\mathbf{W}_{r,1}$ is used for projecting $\mathbf{h}$ and $\mathbf{W}_{r,2}$ is for projecting $\mathbf{t}$. TransSparse [10] is similar to STransE, but it uses a more complex method to project $\mathbf{h}, \mathbf{t}$. PTransE [11] is also a TransE-like model. It uses additional path information for training. For example, if there exist two triples $(e_1, r_1, e_2), (e_2, r_2, e_3)$, which can be regarded as a path in a KG, and another triple (e_1, r_x, e_3) holds simultaneously, then the path $e_1 \rightarrow r_1 \rightarrow e_2 \rightarrow r_2 \rightarrow e_3$ is a valuable path recorded as (e_1, r_1, r_2, e_3) . However, preparing desirable paths needs to iterate over all possible paths, and thus this process may consume quadratic resources. We argue that PTransE and many path-based models may be inefficient to model large KGs.

All the aforementioned models choose to minimize an energy function that is used in or similar to TransE. Moreover, except TransE and TransH, the remaining models require pre-trained entity and relation embeddings from TransE as initial input, which increases the training expense and also blurs their actual performance.

2.2 Non-Translational Models

There also exist a number of models that are very different from the translational models. A neural tensor network (NTN) [12] defines a different loss function and creates many variables for each relation. This makes NTN unsuitable for KGs having plenty of relations. Furthermore, NTN requires word embeddings as initial input, which severely increases the training expense. DISTMULT [13] is as simple as TransE, but it employs a completely different energy function. More specifically, it is based on the Bilinear model [14] and represents each relation as a diagonal matrix. Node+LinkFeat (abbr. NLFeat) [15] can also be regarded as a path-based model like PTransE, but it only needs to extract paths of length 1 for constructing node and link features. For example, if there exists a triple (e_i, r_k, e_j) in a KG, and (e_i, r', e_j) holds at the same time, then a binary feature is constructed as $1(r' \& r_k)$. Although using paths of length 1 to construct features is much easier than using longer paths, it still consumes considerable resources for large KGs.

Recently, the deep neural networks have received much attention in KG completion. Instead of using simple algebraic operations, the deep neural models stack a group of different neural layers to model complex patterns in KGs. R-GCN [16] leverages the conventional graph convolutional networks [17] and extends it to multi-relational data like KGs. ConvE [18] applies convolutional neural networks (CNNs) [19] for KG completion. ConvKB [20] is also based on CNNs but implemented with a novel fashion.

Other methods like [21, 22] use extra data that cannot be extracted from the original training data, such as text corpora or entity descriptions. Due to the focus of this paper, we do not consider them currently.

3. BASIC AND DEEP SEQUENTIAL MODELS

In this section, we first describe our basic sequential model BSKG. Then, we present our deep sequential model DSKG with three different integrating strategies. Finally, two optimization methods useful for accelerating convergence and preventing over-fitting are introduced.

3.1 Basic Sequential Model

We start with the basic sequential model, which has only one single RNN cell. Let $\mathcal{T}, \mathcal{E}, \mathcal{R}$ be the sets of triples, entities and relations in a KG, respectively. Given a triple $(h, r, t) \in \mathcal{T}$, we represent $h, t \in \mathcal{E}$ and $r \in \mathcal{R}$ all as k-dimensional

vectors $\mathbf{h}, \mathbf{t}, \mathbf{r}$, and use the tail entity prediction task for example to write the basic RNN layer as follows:

$$\begin{aligned}\mathbf{s}_h &= c(\mathbf{s}_0, \mathbf{h}) \\ \mathbf{s}_r &= c(\mathbf{s}_h, \mathbf{r})\end{aligned}$$

where c denotes an RNN cell, which receives its previous hidden state and the current element as input, and outputs the processed hidden state. $\mathbf{s}_h$ and $\mathbf{s}_r$ denote the processed hidden states of input h, r , respectively. $\mathbf{s}_0$ denotes the zero hidden state for initialization.

There are a variety of candidate RNN cells that can be used to implement our model, e.g., the gated recurrent unit (GRU) cell [23] or the long short-term memory (LSTM) cell [24]. For simplicity, we do not discuss the details here, but use $c(\cdot)$ to abstractly represent the operations for calculating the output of RNN cells.

Then, we can respectively use $\mathbf{o}_h$ and $\mathbf{o}_r$ to predict the relations of h and the tail entities of $h \rightarrow r$ as follows:

$$\begin{aligned}\mathbf{p}_r &= \mathbf{W}_1 \mathbf{o}_h + \mathbf{b}_1 \\ \mathbf{p}_t &= \mathbf{W}_2 \mathbf{o}_r + \mathbf{b}_2\end{aligned}$$

where $\mathbf{W}_1$ denotes the weight matrix of relation prediction, which has a shape of $|\mathcal{R}| \times k$. $\mathbf{b}_1$ denotes the bias vector. $\mathbf{W}_2, \mathbf{b}_2$ are defined similarly. Therefore, $\mathbf{p}_r$ and $\mathbf{p}_t$ can be directly regarded as the unscaled probabilities for relation prediction and tail entity prediction, respectively.

We respectively calculate the sampled softmax cross-entropy relation loss L_r and entity loss L_t [25] for the unscaled probabilities $\mathbf{p}_r$ and $\mathbf{p}_t$ as follows:

$$\begin{aligned}L_r &= -\log(n_r) - \sum_{i \in \mathcal{N}_r} \log(1 - n_i) \\ L_t &= -\log(n_t) - \sum_{i \in \mathcal{N}_t} \log(1 - n_i)\end{aligned}$$

where n_i denotes the normalized probability of the i -th relation with softmax function and y_i is the label for n_i . Due to the fact that the r -th relation (the t -th entity) is the only correct one in this case, we can remove the zero components in the sum. $\mathcal{N}_r$ and $\mathcal{N}_t$ denote the sets of sampled negative labels for relation prediction and entity prediction, respectively.

There may exist multiple correct labels for both relation prediction and entity prediction, so in this paper we only use the current input triple to provide the correct labels for both relation prediction and entity prediction. For example,

let (h_0, r_0, t_0) be the current training triple and $t_1, \dots, t_m$ be other correct labels for $(h_0, r_0, ?)$, as these triples $(h_0, r_0, t_i), i = 1, \dots, m$ also exist in the training data. But we only consider t_0 as the exclusively right label and avoid adding it during sampling negative labels. We propose this method since it makes training much faster and does not need to prepare the label information for multi-label classification. This is also because the number of possible negative labels is much larger than that of correct ones.

Finally, we can jointly minimize both two losses L_r and L_t , or just minimize L_t , as the final loss for training:

$$\begin{aligned} L_j &= L_r + L_t \\ L_e &= L_t \end{aligned}$$

where L_j denotes the joint loss, and L_e denotes the entity-only loss. Our experimental results show that minimizing L_j performs better on the entity prediction task, and also makes our models capable of predicting relations and triples.

3.2 Deep Sequential Model

RNN can be considered to be a deep neural network with indefinite layers, since it can recurrently process input in arbitrary times. However, only using a single RNN cell to process all information may be inefficient for KGs. Multi-layer RNNs have shown promising performance on modeling complex hierarchical architectures in the NLP area [26], and KGs happen to have such architectures. Therefore, we propose DSKG, which uses multi-layer RNNs to model complex KGs. Three different integrating strategies, namely Joint, Respective and Cascade, are designed to integrate the RNN cells in DSKG. Figure 1 illustrates a two-layer version of our DSKG model with the three strategies. We describe them in detail below:

3.3 Joint

This strategy uses two distinct RNN cells c_1 and c_2 to process both entities and relations. The output hidden state of each cell needs to be fed to itself next time. Note that c_1 's output hidden state is also conveyed to c_2 . By doing this, we do not need c_1 to give an explicit result for each input, but enable it to convey its own comprehension to c_2 . c_2 performs prediction based on its previous hidden state and c_1 's processed hidden state, which include both sequential and hierarchical information. We formalize this process as follows:

$$\begin{aligned} \mathbf{s}_h^1 &= c_1(\mathbf{s}_0, \mathbf{h}) \\ \mathbf{s}_r^1 &= c_1(\mathbf{s}_h^1, \mathbf{r}) \\ \mathbf{s}_h^2 &= c_2(\mathbf{s}_0, \mathbf{s}_h^1) \end{aligned}$$

$$\mathbf{s}_r^2 = c_2(\mathbf{s}_r^1, \mathbf{s}_r^1)$$

We still use Equation (2) to calculate the output of RNN cells, and get two output pairs $(\mathbf{o}_h^1, \mathbf{o}_h^2)$ and $(\mathbf{o}_r^1, \mathbf{o}_r^2)$. For each pair, we can either combine its two outputs by weights or just use its last output to predict relations or tail entities. The former one has an advantage of modeling long hierarchical sequences, such as multiple brace-matching; while the latter one usually performs better when the sequences are short [26].

In this paper, we only use the output of the last cells for prediction, since triples in KGs are short.

3.4 Respective

Because entities and relations have very different characteristics, we believe that building multi-layer RNNs for them respectively may help our model capture very complex structures. Based on this intuition, we provide the relation layer with independent RNN cells. As shown in Figure 1, we still use c_1 and c_2 to process input h , but assign independent RNN cells c_3 and c_4 to process r :

$$\mathbf{s}_h^1 = c_1(\mathbf{s}_0, \mathbf{h})$$

$$\mathbf{s}_r^1 = c_3(\mathbf{s}_h^1, \mathbf{r})$$

$$\mathbf{s}_h^2 = c_2(\mathbf{s}_0, \mathbf{s}_h^1)$$

$$\mathbf{s}_r^2 = c_4(\mathbf{s}_r^1, \mathbf{s}_r^1)$$

This is the only difference compared with Equation (7). Our experiments show that this strategy improves the performance on completing KGs with more complex structures. For example, FB15K [4] is considered to be a complex data set since it has more than 1,000 different relations, while WN18 [4] only has 18 types of relations.

3.5 Cascade

We propose this strategy because we believe that longer sequential RNN cells may model knowledge structures better, and the cells in the entity layer or the relation layer can concentrate more on its own current work. For comparison, c_1 and c_2 in Respective pass their output hidden states to c_3 and c_4 , respectively. Differently, Cascade only feeds c_2 's output hidden state to c_3 . So, Cascade cuts down the conflicts between the entity layer and the relation layer. We believe that this can help improve performance on triple prediction. We formalize this strategy as follows:

$$\mathbf{s}_h^1 = c_1(\mathbf{s}_0, \mathbf{h})$$

$$\mathbf{s}_h^2 = c_2(\mathbf{s}_0, \mathbf{s}_h^1)$$

$$\mathbf{s}_r^1 = c_3(\mathbf{s}_h^2, \mathbf{r})$$

$$\mathbf{s}_r^2 = c_4(\mathbf{s}_r^1, \mathbf{s}_r^1)$$

However, h needs to be sequentially conveyed four times to obtain the final output. Such long sequence may help our model process h 's features, but may also lose some information of h during conveying.

In addition to the above three integrating strategies, we also design a simple variant, called Entity-only, for comparing with Joint. The only difference between them is that Entity-only only optimizes the entity-only loss, which is more like previous entity prediction models.

3.6 Batch Normalization and Dropout

To accelerate convergence and prevent over-fitting, we also employ two optimization methods in our models.

3.7 Batch Normalization

Batch normalization is widely used to alleviate the impact of improperly-initialized neural networks [27]. It enforces the input of next layers to a uniform Gaussian distribution. In our models, the batch normalization layers are placed before and after both the entity layer and the relation layer.

3.8 Dropout

Dropout is a simple but effective method to prevent over-fitting [28]. It is implemented by keeping a neuron active with probability p_D during training. In our models, we add the dropout layers before and after each RNN cell.

4. EXPERIMENTS

We implemented our method with TensorFlow, and conducted three different experiments, namely entity prediction, relation prediction and triple prediction, to evaluate it. In this section, we first introduce the data sets and experiment settings. Then, we describe the procedure of each experiment and report the corresponding results.

By carrying out these experiments, we want to answer the following three questions:

1. Can our method achieve state-of-the-art performance on some benchmark data sets?
2. How does our method perform on the new KG completion tasks such as relation prediction and triple prediction?
3. What are the strengths and weaknesses of each integrating strategy in our deep sequential model?

4.1 Data Sets and Experiment Settings

In all our experiments, we chose FB15K and WN18 as our data sets, which were proposed in [4] and used by a large number of previous methods. FB15K has 1,345 different relations, while WN18 contains 18 distinct relations. The detailed statistical data of these two data sets are listed in Table 2.

Table 2. Statistics of the experimental data sets.

Data set	Entities	Relations	Training triples	Validation triples	Testing triples
FB15K	14,951	1,345	483,142	50,000	59,071
WN18	40,943	18	141,442	5,000	5,000

For each data set, we used the Adam [29] optimizer to train one model for all the evaluation tasks and stopped training when the entity prediction result on the validation data is optimized. Thus, the relation prediction and triple prediction results may not be optimal. For both FB15K and WN18, we set the parameters as follows: learning rate $l = 0.001$, embedding dimension $k = 512$, negative sample number $n_s = 512$, batch size $n_B = 2,048$, and the keep-probability for dropout layers $p_D = 0.5$. We employed the LSTM cells in all our models and used two such cells in DSKG (Joint), four in DSKG (Respective) and DSKG (Cascade), just as shown in Figure 1. Adding more cells for DSKG may improve the performance, but would cause slower training speed. Also, our experiments would demonstrate the effectiveness of the two-layer DSKG model.

Additionally, we added the reversed relations in the training data. This strategy is helpful to model relation pairs reversed to each other. There are many methods using reversed relations, such as PTransE [11] and ConvE [15]. We directly used the results reported in their papers. Specifically, for each triple (h, r, t) in the training data, we constructed a reversed triple (t, r^{-1}, h) and added it into the training data. Adding the reversed relations also enabled our models to predict head and tail entities in an integrated fashion, which means that our models can predict tail entities with input $(h, r, ?)$, and predict head entities with $(t, r^{-1}, ?)$ simultaneously. The reversed relations provide a convenient way to evaluate the head prediction. Also, they enhance the connectivity of KGs. Adding them can slightly improve the performance, while significantly boosting the convergence speed.

4.2 Entity Prediction

Entity prediction aims to predict h (or t) given an incomplete triple $(?, r, t)$ (or $(h, r, ?)$). We evaluated our models with the same method used in [4]. For each triple (h, r, t) in the testing data, we constructed two incomplete triples $(h, r, ?)$, $(t, r^{-1}, ?)$ for tail prediction and head prediction, respectively.

Following [4] and many others, two evaluation metrics were used: the mean rank of correct entities (MR) and the percentage of correct entities in ranked top-10 (Hits@10). Besides, an incomplete triple like $(h, r, ?)$ may have multiple correct tails. Entity t in the current testing triple (h, r, t) is only one of the correct entities. Thus, we also employed the filtered mean rank (FMR) and filtered Hits@10 (FHits@10) [4], which removed all the correct entities except t during ranking.

The experimental results on FB15K and WN18 are shown in Table 3. We can observe that:

1. DSKG outperformed the other models on FB15K and also achieved superior performance on WN18 for Hits@10 and FHits@10.
2. Compared with BSKG, DSKG significantly improved the performance of FHits@10 on FB15K, which has a more complex structure than WN18.
3. DSKG (Joint) outperformed DSKG (Entity-only) on both FB15K and WN18, which proved that jointly optimizing relation prediction and entity prediction is helpful for predicting entities.
4. The three strategies in DSKG performed similarly on FB15K, but diverged for MR and FMR on WN18. For example, DSKG (Cascade) achieved a better FHits@10 than DSKG (Joint) on WN18, but it performed worse on MR and FMR. We argue that WN18 only has 5,000 triples for testing, but it has about 40,000 different entities. So, if a model fails on one triple in the testing data, its MR and FMR would drop 8.0. However, for FB15K, its MR and FMR would only drop 0.25 at most.
5. DSKG (Respective) outperformed DSKG (Joint) on FB15K, which may show that using distinct RNN cells for the entity layer and the relation layer is helpful for modeling complex KGs.

Table 3. Entity prediction results.

Models	FB15K			WN18				
	MR	FMR	Hits@10	FHits@10	MR	FMR	Hits@10	FHits@10
TransE	243	125	34.9	47.1	263	251	75.4	89.2
TransH	212	87	45.7	64.4	401	388	73.0	82.3
TransD	194	91	53.4	77.3	224	212	79.6	92.5
TransR	198	77	48.2	68.7	332	315	78.8	91.7
CTransR	231	89	51.8	70.2	-	-	-	-
PTransE	200	54	51.8	77.7	-	-	-	-
DISTMULT	254	117	49.1	71.5	655	628	82.2	93.6
NLFeat	181	78	49.4	73.8	-	-	-	-
STransE	217	69	52.8	79.7	257	245	79.7	93.4
ConvE	246	51	55.8	83.1	374	356	82.6	95.5
DSKG	198	48	56.2	83.5	342	328	83.1	95.8
(Entity-only)								

Models	FB15K				WN18			
DSKG (Joint)	192	45	57.1	84.2	338	324	83.4	96.1
DSKG (Re-spective)	189	43	57.8	84.9	341	327	83.2	95.9
DSKG (Cascade)	195	46	56.9	84.0	345	331	83.5	96.2

Note: “-” indicates the result unreported in literature. Models are ordered according to their publishing years.

Table 4. Detailed entity prediction results on FB15K by relationship categories.

Models	Head prediction (FHits@10)				Tail prediction (FHits@10)			
	1:1	1:M	M:1	M:N Overall	1:1	1:M	M:1	
TransE	43.7	65.7	18.2	47.2 47.1	43.7	18.2	65.7	
TransH	66.8	73.8	31.1	57.9 64.4	66.8	31.1	73.8	
TransD	77.3	89.1	38.8	70.2 77.3	77.3	38.8	89.1	
TransR	78.8	89.2	34.1	69.2 77.2	78.8	34.1	89.2	
CTransR	80.9	89.4	38.6	70.8 78.0	80.9	38.6	89.4	
PTransE	78.6	91.3	51.4	78.2 77.7	78.6	51.4	91.3	
STransE	80.8	94.0	48.5	78.6 79.7	80.8	48.5	94.0	
DSKG (Entity-only)	82.1	94.3	52.4	80.1 83.5	82.1	52.4	94.3	
DSKG (Joint)	82.8	94.7	53.1	80.9 84.2	82.8	53.1	94.7	
DSKG (Re-spective)	83.4	95.1	54.2	81.6 84.9	83.4	54.2	95.1	
DSKG (Cascade)	82.5	94.5	53.0	80.5 84.0	82.5	53.0	94.5	

We can observe that BSKG and DSKG outperformed the others on the complex relationship categories (i.e., 1:M, M:1 and M:N). The reasons are: (1) Our models are probability-based rather than margin-based, so they would not suffer from the problem of modeling the complex relationships in the margin-based models; (2) RNN cells can properly model triples that have complex relationships, since they can transfer and model entities or relations alternately.

4.3 Relation Prediction

Relation prediction aims to predict relations given a head (or tail) entity only. For each triple (h, r, t) in the same testing data as used in the entity prediction experiment, we took h as input for predicting r , and t for r^{-1} . Note that the testing data used here can be redundant. For example, (h_1, r_1, t_1) and (h_1, r_1, t_2) are two different triples for entity prediction, while relation prediction only considers h_1 and r_1 . We did not remove the redundant testing triple (h_1, r_1, t_2) in the relation prediction and triple prediction experiments, since the occurrence of (h_1, r_1) can be regarded as its weight, reflecting its prevalence and importance in a KG.

Since there are no previous models designed for this task, we proposed a specific relation prediction model for comparison, which has two fully-connected layers. Each layer in this model has a weight matrix and a bias vector, and uses ReLU as the activating function. Note that DSKG (Entity-only) may not suit this experiment task, because it does not minimize the relation prediction loss during training. We used the same evaluation metrics in this experiment.

The relation prediction results are shown in Table 5. For FB15K, we can observe the following:

1. Predicting forward relations was more accurate than predicting backward relations for FHits@10 and FMR, due to the facts that KGs are usually constructed by humans, and (h, r, t) is a more natural representation than (t, r^{-1}, h) .
2. DSKG outperformed the fully-connected two-layer model, which verified that jointly optimizing relation prediction and entity prediction can also help predict relations.
3. Both BSKG and DSKG achieved an FMR of 1.5 on predicting forward relations, which indicated that our models predicted the correct forward relations of an entity with a very high probability on average.

We also evaluated our models on WN18. Because this data set has only 18 distinct relations, it is not surprising that all the models achieved similar performance. Still, DSKG showed a slight advantage for Hits@10 and FHits@10. It is worth noting that our models can also predict r given both h and t using the same method as in [11] (i.e., the aforementioned relationship prediction task).

Table 5. Relation prediction results.

Data sets	Models	Forward (h →)			Backward (t →)		
		Hits@10	FHits@10	FMR	Hits@10	FHits@10	FMR
FB15K	Fully-connected	89.2	91.4	2.1	1.8	87.3	89.5
	DSKG (Joint)	91.8	93.7	1.8	1.5	89.1	91.2

Data sets	Models	Forward (h →)				Backward (t →)			
WN18	DSKG (Re- spec- tive)	92.1	94.0	1.7	1.5	89.4	91.5	2.1	
	DSKG (Cas- cade)	91.5	93.5	1.9	1.6	88.8	90.9	2.3	
	Fully- connected	98.7	99.1	1.1	1.0	98.5	98.9	1.2	
	DSKG (Joint)	99.1	99.4	1.0	1.0	98.9	99.2	1.1	
	DSKG (Re- spec- tive)	99.2	99.5	1.0	1.0	99.0	99.3	1.1	
	DSKG (Cas- cade)	99.0	99.3	1.0	1.0	98.8	99.1	1.1	

4.4 Triple Prediction

DSKG is capable of not only predicting entities and relations, but also predicting the whole triples given an entity. So, triple prediction can be considered to be a more integrated task for KG completion. For each triple (h, r, t) in the same testing data as used in the entity prediction experiment, we only used h as input to predict the triples about h . Each model was first asked to predict the relation r_p of h with the highest probability, and then used this relation to predict the most probable tail t_p . As a result, we got one output triple (h, r_p, t_p) for each input h . Note that the models in this experiment only predicted triples in the forward direction, because: (i) we have shown that the backward relation prediction results on FMR were worse than the forward relation prediction results; and (ii) predicting forward triples of an entity is more natural to KG modeling.

We evaluated the output triples with the following method. Let $\mathcal{S}_R$ denote the set of all triples in a KG, $\mathcal{S}_O$ denote the set of output triples, $\mathcal{S}_p$ denote the union of testing and validation triples. $\mathcal{S}_R$ and $\mathcal{S}_p$ are referred to as the representation and prediction triple sets, respectively. We calculate the correct representation number as $n_R = |\mathcal{S}_R \cap \mathcal{S}_O|$, and the correct prediction number as $n_p = |\mathcal{S}_p \cap \mathcal{S}_O|$.

So, the representation precision and the prediction precision are calculated as follows:

$$\text{Representation precision} = \frac{n_R}{|\mathcal{S}_O|}$$

$$\text{Prediction precision} = \frac{n_p}{|\mathcal{S}_O|}$$

We slightly modified the evaluation procedures of TransE, TransR and ConvE for comparison, since they are unable to predict relations. We provided relations for them using four different providers, i.e., DSKG (Joint), DSKG (Respective), DSKG (Cascade) and correct relations.

The experimental results are shown in Table 6 . We can observe the following:

1. DSKG outperformed TransE and TransR, especially on WN18, even though providing the correct relations for them. This may be caused by the fact that both TransE and TransR are margin-based, and thus comparing distances for prediction may mislead them.
2. ConvE achieved similar performance compared with BSKG, but performed slightly weaker than DSKG.
3. Even we did not provide the correct relations, and DSKG still achieved good results, especially DSKG (Cascade), which outperformed all the others on both FB15K and WN18.
4. For all the integrating strategies in DSKG, using their own providers to provide relations performed even better for representation precision, which is probably because correct relations are actually extracted from testing data, while their own providers preferred to give more relations that have already been trained.

Table 6. Triple prediction results.

Relation providers	Models	FB15K	WN18
		Representation precision	Prediction precision
Correct relations	TransE (Joint)	0.342	0.128 0.156 0.087
	TransE (Respective)	0.341	0.127 0.155 0.086
	TransE (Cascade)	0.340	0.126 0.154 0.085
	TransR (Joint)	0.351	0.132 0.162 0.089
	TransR (Respective)	0.350	0.131 0.161 0.088

Relation providers	Models	FB15K	WN18		
Own providers	TransR	0.349	0.130	0.160	0.087
	(Cascade)				
	ConvE	0.378	0.141	0.184	0.092
	(Joint)				
	ConvE	0.377	0.140	0.183	0.091
	(Respec-				
	tive)				
	ConvE	0.376	0.139	0.182	0.090
	(Cascade)				
	BSKG	0.381	0.143	0.186	0.093
	(Joint)				
	BSKG (Re-	0.382	0.144	0.187	0.094
	spective)				
	BSKG	0.383	0.145	0.188	0.095
	(Cascade)				
	DSKG	0.392	0.151	0.195	0.098
	(Joint)				
	DSKG	0.394	0.153	0.197	0.099
	(Respec-				
	tive)				
	DSKG	0.398	0.156	0.201	0.102
	(Cascade)				
	TransE	0.328	0.121	0.148	0.082
	TransR	0.336	0.125	0.152	0.084
	ConvE	0.365	0.135	0.172	0.088
	DSKG	0.392	0.151	0.195	0.098
	(Joint)				
	DSKG	0.394	0.153	0.197	0.099
	(Respec-				
	tive)				
	DSKG	0.398	0.156	0.201	0.102
	(Cascade)				

5. CONCLUSION AND FUTURE WORK

In this paper, we proposed a new KG completion method that extends multi-layer RNNs to model triples in KGs as sequences. Our experimental results on FB15K and WN18 showed that our method achieved superior performance not only on the benchmark entity prediction task, but also on two new KG completion tasks to predict relations and triples given one single entity.

For future work, we plan to explore the following three directions:

- **Integrating the attention mechanism** [30] in our method. While the attention mechanism has shown its power in the NLP area, applying it

to KG completion has not been well studied. In future, we would like to extend our method with this mechanism to improve its inference ability.

- **Using a provided KG to complete another KG.** Recently, several methods start to leverage extra textual data for improving KG completion. However, textual data like entity descriptions and text corpora are written in natural language. Due to the ambiguity and heterogeneity of textual data, they may bring mistakes in prediction. Therefore, we think that taking existing KGs as another kind of extra data may improve performance.
- **Embedding KGs for entity alignment.** We also want to consider the problem of learning KG embeddings for entity alignment. Particularly, we look forward to learning embeddings of different KGs in a unified space and employing RNNs to explore neighboring context for entity alignment.

AUTHOR CONTRIBUTIONS

L. Guo (lbguo.nju@gmail.com) designed the model. Q. Zhang (qhzhang.nju@gmail.com) conducted main experiments. W. Hu (whu@nju.edu.cn) assembled the manuscript. Z. Sun (zqsun.nju@gmail.com) and Y. Qu (yzqu@nju.edu.cn) revised the whole paper.

ACKNOWLEDGEMENTS

This work was supported by the National Natural Science Foundation of China (No. 61872172).

REFERENCES

- [1] S. Auer, C. Bizer, G. Kobilarov, J. Lehmann, R. Cyganiak, & Z.G. Ives. DBpedia: A nucleus for a Web of open data. In: K. Aberer et al. (eds.) The Semantic Web. Berlin: Springer, 2007, pp. 722-735. doi: 10.1007/978-3-540-76298-0_52.
- [2] K.D. Bollacker, C. Evans, P. Paritosh, T. Sturge, & J. Taylor. Freebase: A collaboratively created graph database for structuring human knowledge. In: Proceedings of the 2008 ACM SIGMOD International Conference on Management of Data, 2008, pp. 1247-1250. doi: 10.1145/1376616.1376746.
- [3] A. Bordes, J. Weston, R. Collobert, & Y. Bengio. Learning structured embeddings of knowledge bases. In: Proceedings of the 25th AAAI Conference on Artificial Intelligence, 2011, pp. 301-306. Available at: <https://www.aaai.org/ocs/index.php/AAAI/AAAI11/paper/view/3659/3898>.
- [4] A. Bordes, N. Usunier, A. Garcia-Durán, J. Weston, & O. Yakhnenko. Translating embeddings for modeling multi-relational data. In: Proceedings of the 26th International Conference on Neural Information Processing Systems

(NIPS), 2013, pp. 2787–2795. Available at: <http://papers.nips.cc/paper/5071-translating-embeddings-for-modeling-multi-relational-data.pdf>.

[5] O. Vinyals, L. Kaiser, T. Koo, S. Petrov, I. Sutskever, & G.E. Hinton. Grammar as a foreign language. In: Proceedings of the 28th International Conference on Neural Information Processing Systems, 2015, pp. 2773–2781. Available at: <http://papers.nips.cc/paper/5635-grammar-as-a-foreign-language.pdf>.

[6] R. Józefowicz, O. Vinyals, M. Schuster, N. Shazeer, & Y. Wu. Exploring the limits of language modeling. In: Proceedings of the 4th International Conference on Learning Representations, 2016, pp. 1–11.

[7] Z. Wang, J. Zhang, J. Feng, & Z. Chen. Knowledge graph embedding by translating on hyperplanes. In: Proceedings of the 28th AAAI Conference on Artificial Intelligence, 2014, pp. 1112–1119. Available at: <https://www.aaai.org/ocs/index.php/AAAI/AAAI14/paper/view/8531/8546>.

[8] Y. Lin, Z. Liu, M. Sun, Y. Liu, & X. Zhu. Learning entity and relation embeddings for knowledge graph completion. In: Proceedings of the 29th AAAI Conference on Artificial Intelligence, 2015, pp. 2181–2187. doi: 10.1016/j.procs.2017.05.045.

[9] D.Q. Nguyen, K. Sirts, L. Qu, & M. Johnson. STransE: A novel embedding model of entities and relationships in knowledge bases. In: Proceedings of the 15th Annual Conference of the North American Chapter of the Association for Computational Linguistics, 2016, pp. 460–466. doi: 10.18653/v1/N16-1054.

[10] G. Ji, K. Liu, S. He, & J. Zhao. Knowledge graph completion with adaptive sparse transfer matrix. In: Proceedings of the 30th AAAI Conference on Artificial Intelligence, 2016, pp. 985–991. Available at: <https://aaai.org/ocs/index.php/AAAI/AAAI16/paper/view/11982/11693>.

[11] Y. Lin, Z. Liu, H.-B. Luan, M. Sun, S. Rao, & S. Liu. Modeling relation paths for representation learning of knowledge bases. In: Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics, 2015, pp. 705–714. doi: 10.18653/v1/D15-1082.

[12] R. Socher, D. Chen, C. Manning, D. Chen, & A. Ng. Reasoning with neural tensor networks for knowledge base completion. In: Proceedings of the 26th International Conference on Neural Information Processing Systems, 2013, pp. 926–934. doi: 10.1109/ICICIP.2013.6568119.

[13] B. Yang, S.W. Yih, X. He, J. Gao, & L. Deng. Embedding entities and relations for learning and inference in knowledge bases. In: Proceedings of the 3rd International Conference on Learning Representations, 2015, pp. 1–13. Available at: https://www.microsoft.com/en-us/research/wp-content/uploads/2016/02/ICLR2015_{updated}.pdf.

[14] M. Nickel, V. Tresp, & H.-P. Kriegel. A three-way model for collective learning on multi-relational data. In: L. Getoor, & T. Scheffer (eds.) Proceedings

of the 28th International Conference on Machine Learning (ICML-11), 2011, pp. 809-816.

[15] K. Toutanova, & D. Chen. Observed versus latent features for knowledge base and text inference. In: Proceedings of the 3rd Workshop on Continuous Vector Space Models and Their Compositionality, 2015, pp. 57-66. doi: 10.18653/v1/w15-4007.

[16] M.S. Schlichtkrull, T.N. Kipf, P. Bloem, R. van den Berg, I. Titov, & M. Welling. Modeling relational data with graph convolutional networks. In: A. Gangemi et al. (eds.) The Semantic Web. ISWC 2018. Cham, Switzerland: Springer, 2018, pp. 593-607. doi: 10.1007/978-3-319-93417-4_38}.

[17] D.K. Duvenaud, D. Maclaurin, J. Aguilera-Iparraguirre, R. Gómez-Bombarelli, T. Hirzel, A. Aspuru-Guzik, & R.P. Adams. Convolutional networks on graphs for learning molecular fingerprints. In: Proceedings of the 28th International Conference on Neural Information Processing Systems, 2015, pp. 2224-2232. Available at: <http://dl.acm.org/citation.cfm?id=2969488>.

[18] T. Dettmers, P. Minervini, P. Stenetorp, & S. Riedel. Convolutional 2d knowledge graph embeddings. In: Proceedings of the 32nd AAAI Conference on Artificial Intelligence, 2018, pp. 1811-1818. Available at: <https://aaai.org/ocs/index.php/AAAI/AAAI18/paper/view/17366>.

[19] Y. LeCun, L. Bottou, Y. Bengio, & P. Haffner. Gradient-based learning applied to document recognition. Proceedings of the IEEE 86(11)(1998), 2278-2324. doi: 10.1109/5.726791.

[20] D.Q. Nguyen, T.D. Nguyen, D.Q. Nguyen, & D.Q. Phung. A novel embedding model for knowledge base completion based on convolutional neural network. In: Proceedings of the 17th Annual Conference of the North American Chapter of the Association for Computational Linguistics, 2018, pp. 327-333. doi: 10.18653/v1/N18-2053.

[21] R. Xie, Z. Liu, J. Jia, H. Luan, & M. Sun. Representation learning of knowledge graphs with entity descriptions. In: Proceedings of the 30th AAAI Conference on Artificial Intelligence, 2016, pp. 2659-2665. Available at: <https://aaai.org/ocs/index.php/AAAI/AAAI16/paper/view/12216/12004>.

[22] H. Xiao, M. Huang, L. Meng, & X. Zhu. SSP: Semantic space projection for knowledge graph embedding with text descriptions. In: Proceedings of the 31st AAAI Conference on Artificial Intelligence, 2017, pp. 3104-3110.

[23] K. Cho, B. van Merriënboer, C. Gulcehre, D. Bahdanau, F. Bougares, H. Schwenk, & Y. Bengio. Learning phrase representations using RNN encoder-decoder for statistical machine translation. In: Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing, 2014, pp. 1724-1734. doi: 10.3115/v1/D14-1179.

[24] S. Hochreiter, & J. Schmidhuber. Long short-term memory. Neural Computation 9(8)(1997), 1735-1780. doi: 10.1162/neco.1997.9.8.1735.

- [25] S. Jean, K. Cho, R. Memisevic, & Y. Bengio. On using very large target vocabulary for neural machine translation. In: Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics, 2015, pp. 1-10.
- [26] M. Hermans, & B. Schrauwen. Training and analyzing deep recurrent neural networks. In: Proceedings of the 26th International Conference on Neural Information Processing Systems, 2013, pp. 190-198. Available at: <http://papers.nips.cc/paper/5166-training-and-analysing-deep-recurrent-neural-networks.pdf>.
- [27] S. Ioffe, & C. Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In: Proceedings of the 32nd International Conference on Machine Learning, 2015, pp. 448-456. Available at: <https://dl.acm.org/citation.cfm?id=3045167>.
- [28] N. Srivastava, G. E. Hinton, A. Krizhevsky, I. Sutskever, & R. Salakhutdinov. Dropout: A simple way to prevent neural networks from overfitting. Journal of Machine Learning Research 15(2014), 1929-1958. Available at: <http://jmlr.org/papers/volume15/srivastava14a/srivastava14a.pdf>.
- [29] D.P. Kingma, & J. Ba. Adam: A method for stochastic optimization. In: Proceedings of the 3rd International Conference on Learning Representations, 2015, pp. 1-15.
- [30] D. Bahdanau, K. Cho, & Y. Bengio. Neural machine translation by jointly learning to align and translate. In: Proceedings of the 3rd International Conference on Learning Representations, 2015, pp. 1-15.

AUTHOR BIOGRAPHY

Lingbing Guo is currently a graduate student at the Department of Computer Science and Technology, Nanjing University, China. He received his Bachelor's degree in Computer Science and Technology in 2016 from Henan University, China. His research interest is knowledge graph completion.

Qingheng Zhang is currently a graduate student at the Department of Computer Science and Technology, Nanjing University, China. He received his Bachelor's degree in Computer Science and Technology in 2017 from Hohai University, China. His research interest is knowledge graph embedding.

Wei Hu is currently an Associate Professor at the State Key Laboratory for Novel Software Technology, the Department of Computer Science and Technology, Nanjing University, China. He received his PhD degree in Computer Software and Theory in 2009, and his Bachelor's degree in Computer Science and Technology in 2005, both from Southeast University, China. His main research interests include knowledge graph, data integration and intelligent software.

Zequan Sun is currently a PhD student at the Department of Computer Science and Technology, Nanjing University, China. He received his Bachelor's degree

in Computer Science and Technology in 2016 from Hohai University, China. His research interest is entity alignment.

Yuzhong Qu is currently a Professor at the State Key Laboratory for Novel Software Technology, the Department of Computer Science and Technology, Nanjing University, China. He received his Bachelor' s and Master' s degrees in Mathematics from Fudan University, China, in 1985 and 1988, respectively, and his PhD degree in Computer Software from Nanjing University, China, in 1995. His research interests include semantic Web, question answering and novel software technology for the Web.

Note: Figure translations are in progress. See original paper for figures.

Source: ChinaXiv –Machine translation. Verify with original.