

Type-aware Embeddings for Multi-Hop Reasoning over Knowledge Graphs

Zhiwei Hu¹, Víctor Gutiérrez-Basulto², Zhiliang Xiang², Xiaoli Li³, Ru Li^{1*}, Jeff Z. Pan^{4*}

¹School of Computer and Information Technology, Shanxi University, China

²School of Computer Science and Informatics, Cardiff University, UK

³Institute for Infocomm Research/Centre for Frontier AI Research, A*STAR, Singapore

⁴ILCC, School of Informatics, University of Edinburgh, UK

zhiwei.hu@whu.edu.cn, {gutierrezbasulto, xiangz6}@cardiff.ac.uk, xlli@i2r.a-star.edu.sg, liru@sxu.edu.cn, j.z.pan@ed.ac.uk

Abstract

Multi-hop reasoning over real-life knowledge graphs (KGs) is a highly challenging problem as traditional subgraph matching methods are not capable to deal with noise and missing information. To address this problem, it has been recently introduced a promising approach based on jointly embedding logical queries and KGs into a low-dimensional space to identify answer entities. However, existing proposals ignore critical semantic knowledge inherently available in KGs, such as type information. To leverage type information, we propose a novel **TypE-aware Message Passing (TEMP)** model, which enhances the entity and relation representations in queries, and simultaneously improves generalization, deductive and inductive reasoning. Remarkably, TEMP is a plug-and-play model that can be easily incorporated into existing embedding-based models to improve their performance. Extensive experiments on three real-world datasets demonstrate TEMP’s effectiveness.

1 Introduction

In recent years, the multi-hop reasoning problem of answering first-order logic queries (FOL) on large-scale incomplete knowledge graphs (KGs) [Pan *et al.*, 2016] has gained a lot of attention in the AI community. A major challenge for traditional subgraph matching methods for query answering is that KGs are inevitably incomplete and noisy. Indeed, when schema [Wiharja *et al.*, 2020] and triples are incomplete in the KG, correct answers are not guaranteed to be found under normal deductive reasoning, leading to empty or wrong answers. Another problem is their intrinsic high computational complexity as they need to keep track of all intermediate entities occurring on reasoning paths, leading to an exponential blow-up. For instance, to answer the query “*List the presidents of Asian countries that have held the Summer Olympics*” shown in Fig. 1, we require two traversing-steps (many more for other queries): one to identify countries that have held the summer Olympics and another one to identify Asian countries, each producing intermediate countries.

To address these challenges, a *query embedding (QE)* approach to query answering has been recently introduced as an alternative to subgraph matching methods. The main idea is to embed entities and queries into a joint low-dimensional vector space such that entities that answer the query are close to the embedding of the query. Several QE models for query answering, showing very promising performance, have been proposed so far [Hamilton *et al.*, 2018; Ren *et al.*, 2020; Ren and Leskovec, 2020; Zhang *et al.*, 2021; Choudhary *et al.*, 2021b; Luus *et al.*, 2021]. However, these models *fail to leverage semantic knowledge inherently available in KGs*, such as entity description [Yao *et al.*, 2019; Daza *et al.*, 2021] or entity type information [Niu *et al.*, 2020; Pan *et al.*, 2021]. Advantages of introducing type information are that: 1) it can enhance the representation of entities or relations; e.g., the types *sports* and *event* can enrich the representation of the entity *Summer Olympics* in the context of sport events (cf. Fig. 1). 2) It can also help tackling the *inductive query answering* problem where entities used in test queries cannot be observed at training time; e.g., consider the queries in Figure 1: “*List the presidents of Asian countries that have held the Summer Olympics*” and “*List the presidents of European countries that have held the Winter Olympics*”, which are generated from two KGs with disjoint sets of entities: *Train KG* and *Test KG*, respectively. Even if the entities *Summer Olympics* and *Winter Olympics* are different, they have similar type information, such as *sports* and *event*. Consequently, after using type information to represent entities, the model associated to the query generated from *Train KG* is also effective on the query generated from *Test KG*.

The goal of this paper is to introduce a *type-aware plug-and-play* model which makes full use of type information in the knowledge graph, and can be easily embedded into existing QE-based models. To this aim, we propose a novel **TypE-aware Message Passing (TEMP)** model, which contains two key components. 1) Type-aware Entity Representations (TER), aggregating type information of entities to strengthen their vector representation (cf. Section 4.1). 2) Type-aware Relation Representations (TRR), using entity type information to construct a global *type graph* to enhance the relation representation, and simultaneously integrate it with its type representation and existing entity type information (cf. Section 4.2). Importantly, some queries have *variable nodes* in the query paths (see Figure 1), which increase

*Contact Authors

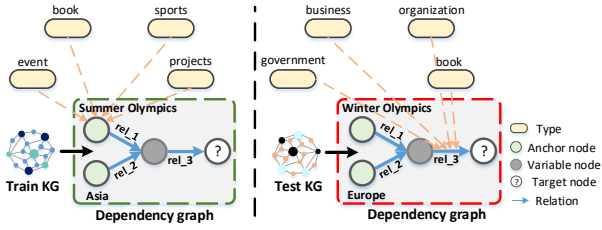


Figure 1: Inductive setting for FOL queries. The left part shows the query “List the presidents of Asian countries that have held the Summer Olympics”. The right part shows the query “List the presidents of European countries that have held the Winter Olympics”. *rel.1*, *rel.2*, and *rel.3* represent relations: *Hold*, *Locate*, and *President of*, respectively.

the difficulty of subsequent reasoning steps in the chain, as variable nodes are unknown. To address this, the TRR component uses a bidirectional mechanism for the *anchor node* to supervise the relations in the query path, and vice versa. Furthermore, as mentioned, after using type information to represent entities and relations, the model becomes inherently inductive as the occurrence of new entities or relations will not affect the type-based representations.

Our main contributions can be summarized as follows:

- We propose TEMP, a novel *type-aware plug-and-play* model for multi-hop reasoning over KGs, that can be easily incorporated into the existing QE-based models.
- We design a new bidirectional integration mechanism that incorporates the pairwise dependencies among {entity, relation, type} information, even in the absence of schema axioms like domain and range.
- Extensive experiments demonstrate that after incorporating TEMP into four state-of-the-art baselines, their *generalization*, *deductive* and *inductive reasoning* abilities are significantly improved across three benchmark datasets consistently.

Data, code, and an extended version with appendix is available at <https://github.com/SXUNLP/QE-TEMP>

2 Related Work

Query Embeddings. QE models first embed entities and FOL queries into a joint low-dimensional vector space, and subsequently compute a similarity score between the *entity representation* and *query representation* in the latent embedding space. In general, according to the type of embedding spaces, QE-based methods can be divided into four categories: (i) geometric-based methods, such as GQE [Hamilton et al., 2018], Q2B [Ren et al., 2020], HypE [Choudhary et al., 2021b], and ConE [Zhang et al., 2021], where logical queries and KG entities are embedded into a geometric vector space as points, boxes, hyperbolic, and cone shapes, respectively; (ii) distribution-based methods, such as BETAE [Ren and Leskovec, 2020], embedding queries to beta distributions with bounded support, and PERM [Choudhary et al., 2021a], using Gaussian densities to reason over KGs; (iii) logic-based methods, relating so-called set logic with FOL [Luus et al., 2021]; (iv) neural-based methods, e.g., EMQL [Sun et al., 2020] using neural retrieval to implement logical operations.

Considering QE-based methods are the mainstream in the current CQA field, we mainly focus on how to construct a *plug-and-play model* to embed the type information for existing QE-based methods.

Other Methods. Besides the QE-based approach, the path-based approach is another method for CQA, but it faces an exponential growth of the search space with the number of hops. For instance, CQD [Arakelyan et al., 2021] uses a beam search method to explicitly track intermediate entities, and repeatedly combines scores from a pretrained link predictor via t-norms to search answers while tracking intermediaries. However, CQD does not support the full set of FOL queries.

Inductive KG Completion (KGC). In the context of KGC, there have been some works on inductive settings where test entities are not seen in the training stage. Based on the source of information used, they can be split into two categories: Using graph structure information, e.g., subgraph or topology structures [Teru et al., 2020; Chen et al., 2021; Wang et al., 2021], or using external information, e.g., textual descriptions of entities [Daza et al., 2021]. However, *all these methods focus on the inductive KGC task, which can be seen as answering simpler one-hop queries.*

Type-aware Tasks. Type information was previously used in other tasks such as KGC or entity typing [Yao et al., 2019; Zhao et al., 2020; Daza et al., 2021; Niu et al., 2020; Pan et al., 2021]. However, these works cannot be directly used for answering FOL queries because this requires multi-hop reasoning, producing intermediate uncertain entities.

3 Background

In this paper, a knowledge graph [Pan et al., 2016] is represented in a standard format for graph-structured data such as RDF. A *knowledge graph* $\mathcal{G}$ is a tuple $(\mathcal{E}, \mathcal{R}, \mathcal{C}, \mathcal{T})$, where $\mathcal{E}$ is a set of entities, $\mathcal{R}$ is a set of relation types, $\mathcal{C}$ is a set of entity types, and $\mathcal{T}$ is a set of triples. Triples in $\mathcal{T}$ are either relation assertions (h, r, t) , where $h, t \in \mathcal{E}$ are respectively the *head* and *tail* entities of the triple, and $r \in \mathcal{R}$ is the *edge* of the triple connecting head and tail, or entity type assertions $(e, type, c)$, where $e \in \mathcal{E}$ is an entity, $c \in \mathcal{C}$ is an entity type and *type* is the instance-of relation [Pan, 2009].

We consider FOL queries that use existential quantification ($\exists$), conjunction ($\wedge$), disjunction ($\vee$) and negation ($\neg$) operations. We will work with FOL queries in Disjunctive Normal Form, i.e. represented as a disjunction of conjunctions. To introduce FOL queries, we assume that $\mathcal{V}_a \subset \mathcal{E}$ represents a set of non-variable *input anchor entities*, $V_1, \dots, V_m$ denote *existentially quantified variables* and $V_?$ is the *target variable*. A FOL query Q is a formula of the following form:

$$Q[V_?] = V_? \cdot \exists V_1, \dots, V_m : c_1 \vee c_2 \vee \dots \vee c_n$$

where $c_i = e_{i1} \wedge \dots \wedge e_{ik}$, $k \leq m$ such that each e_{ij} is of one of the following forms: $r(V_a, V)$, $\neg r(V_a, V)$, $r(V, V')$ or $\neg r(V, V')$, with $V_a \in \mathcal{V}_a$, $V \in \{V_?, V_1, \dots, V_m\}$, $V' \in \{V_1, \dots, V_m\}$, $V \neq V'$.

The *dependency graph (DG)* of a query Q is a graphical representation of Q , where nodes correspond to variable or non-variable arguments in Q and edges correspond to relations in Q . Figure 1 shows an example of a DG.

projections may cause exponential growth in the search space, further increasing the complexity of the model.

We start by observing that the types of a relation are correlated with its representation in the KG. For example, assuming a relation r has types *government/political appointer* and *organization/role*, then we can plausibly infer that the relation r represents *President of*. In long-chain queries, type-enhancement on relations can help to reduce the answer entity space and cascading errors caused by multiple projections. However, in most existing KGs, relations lack specific type annotations (such as domain and range). We address this problem by building, based on the original KG, a novel *type graph* with types as nodes and relations as edges (see bottom left of Fig. 2). In a subsequent step, we aggregate the type information on the type graph to obtain the type embedding corresponding to a specific relation. Finally, we integrate the entity representation, the aggregated type information of a relation and its representation by a bidirectional attention mechanism, so that the intermediate variable nodes can perceive the message of anchor or target nodes and of the relations in the chain of reasoning (see right of Fig. 2). This will help to avoid the weakening of the connection between anchor and target entity caused by long-chain reasoning.

Step 1: Type Graph Construction

We formally define a *type graph* $\mathcal{G}_{tp}$. Let $\mathcal{G} = (\mathcal{E}, \mathcal{R}, \mathcal{C}, \mathcal{T})$ be a KG. For a relation $r \in \mathcal{R}$, $\mathcal{T}_r \subseteq \mathcal{T}$ denotes the set of relation assertions in which r occurs. For a relation assertion $t \in \bigcup_{r \in \mathcal{R}} \mathcal{T}_r$, $hd(t)$ and $tl(t)$ respectively denote the head and tail entities of t , and $tp_t(hd(t)) = \{c \mid (hd(t), type, c) \in \mathcal{T}\}$ denotes the set of types of the head of t ; $tp_t(tl(t))$ is defined analogously. Since r may occur in multiple relation assertions, we will compute the type information of r by taking the intersection of the types of the head and tail entities of relation assertions in which r occurs. For $r \in \mathcal{R}$, we define

$$tp_r^{hd}(\mathcal{G}) = \bigcap_{t \in \mathcal{T}_r} tp_t(hd(t)), \quad tp_r^{tl}(\mathcal{G}) = \bigcap_{t \in \mathcal{T}_r} tp_t(tl(t)).$$

In addition, we define $\mathcal{G}_{tp} = (V, E, T)$ by setting $V = \bigcup_{r \in \mathcal{R}} tp_r^{hd}(\mathcal{G}) \cup tp_r^{tl}(\mathcal{G})$, $E = \mathcal{R}$, and $(v, r, v') \in T$ if there exists $t \in \mathcal{T}_r$ such that $v = tp_t(hd(t))$ and $v' = tp_t(tl(t))$.

Step 2: Relation Type Aggregation

For a given relation $r \in E$, we define the types associated with a relation as $tp_r(\mathcal{G}_{tp}) = tp_r^{hd}(\mathcal{G}) \cup tp_r^{tl}(\mathcal{G})$. We fix an arbitrary linear order on the elements of $tp_r(\mathcal{G}_{tp})$, and denote by $tp_r^i(\mathcal{G}_{tp})$ the i -th type, for all $1 \leq i \leq |tp_r(\mathcal{G}_{tp})|$. Note that not all types in $tp_r(\mathcal{G}_{tp})$ are relevant for answering a given query. For example, assume that the relation *has part* contains the types $\{vehicle, animal, universe\}$. For the query “*What organs are parts of a cat?*”, we should give type *animal* more attention, but for the query “*What components are parts of a car?*” we should concentrate on the type *vehicle*. So, instead of simply concatenating (or averaging) all the type information associated to a relation, we model the relation type aggregation as an attention neural network, defined as:

$$\mathcal{H}^s = \sum_i a_i \odot \mathcal{H}_s^i \quad (5)$$

$$a_i = \frac{\exp(\mathbf{MLP}(\mathcal{H}_r^i))}{\sum_j \exp(\mathbf{MLP}(\mathcal{H}_r^j))} \quad (6)$$

$\mathcal{H}^s$ is the vector representation of the aggregated type information $tp_r(\mathcal{G}_{tp})$; $\mathcal{H}_s^i \in \mathbb{R}^{d \times 1}$ is the vector representation of the i -th type $tp_r^i(\mathcal{G}_{tp})$, which is initialized to a uniform distribution with dimension d , $1 \leq i \leq |tp_r(\mathcal{G}_{tp})|$; $a_i \in \mathbb{R}^{d \times 1}$ is a positive weight vector that satisfies $\sum_{i=1}^n [a_i]_j = 1$ for all $1 \leq j \leq d$; and $\mathbf{MLP}(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a multi-layer perceptron network.

Step 3: Pairwise Representation Integration

When embedding queries, integrating the information of entities, relations, and types can help to smooth decision boundaries, but this needs to be done in a way that the intended match of the query into the KG is captured. For example, for the query “*Which countries have held the Summer Olympics?*”, we need to concentrate on *Held* connections from *Summer Olympics*, rather than e.g., *Watch* connections. Analogously, we should only consider *Held* connections starting at *Summer Olympics*, rather than e.g., at *World Cup*. To properly capture this restriction in the triple $\{\mathcal{H}^e, \mathcal{H}^r, \mathcal{H}^s\}$ ($\mathcal{H}^e$ and $\mathcal{H}^s$ defined as in Equations (4) and (5), and $\mathcal{H}^r$ is the initialization relation vector), we introduce a *bidirectional attention mechanism* [Zhang *et al.*, 2020] to integrate each state of pairwise representation pairs: *entity-relation*, *entity-type*, and *relation-type*. Here, we show how to do this for entity-relation pair. Bidirectional integration representation between $\mathcal{H}^e$ and $\mathcal{H}^r$ can be calculated as follows:

$$\mathcal{G}^{er} = \text{Relu}(W_1 \begin{bmatrix} \mathcal{H}^e \ominus \mathcal{H}^r \\ \mathcal{H}^e \otimes \mathcal{H}^r \end{bmatrix} + b_1) \quad (7)$$

$$\mathcal{G}^{re} = \text{Relu}(W_2 \begin{bmatrix} \mathcal{H}^r \ominus \mathcal{H}^e \\ \mathcal{H}^r \otimes \mathcal{H}^e \end{bmatrix} + b_2) \quad (8)$$

$\{W_1, W_2\} \in \mathbb{R}^{2d \times 2d}$ and $\{b_1, b_2\} \in \mathbb{R}^{2d \times 1}$ are learnable parameters. We use element-wise subtraction $\ominus$ and multiplication $\otimes$ to build better matching representations [Tai *et al.*, 2015]. $\mathcal{G}^{er} \in \mathbb{R}^{2d \times 1}$ is the result of integrating entity relation information. Through bidirectional integration of entities and relations, we simultaneously get a relation-aware entity representation and an entity-aware relation representation, capturing the interaction between entities and relations.

We then use a gated mechanism to combine the results produced by bidirectional fusion as it better regulates the information flow [Srivastava *et al.*, 2015]. Take the entity fusion representation as an example, using the relation-aware entity $\mathcal{G}^{er}$ and type-aware entity $\mathcal{G}^{es}$ representations as input, the final representation of entity is computed as

$$g = \sigma(W_3 \mathcal{G}^{er} + W_4 \mathcal{G}^{es} + b_3 + b_4) \quad (9)$$

$$\widetilde{\mathcal{G}}^e = g * \mathcal{G}^{er} + (1 - g) * \mathcal{G}^{es} \quad (10)$$

$\{W_3, W_4\} \in \mathbb{R}^{2d \times 2d}$ and $\{b_3, b_4\} \in \mathbb{R}^{2d \times 1}$ are the parameters to learn. g is the reset gate, and $\widetilde{\mathcal{G}}^e \in \mathbb{R}^{2d \times 1}$ is the final entity representation.

To transform the fused feature to the original vector size, we use one linear layer: $\widetilde{\mathcal{H}}^e = W_5 \widetilde{\mathcal{G}}^e + b_5$, where $W_5 \in \mathbb{R}^{d \times 2d}$ and $b_5 \in \mathbb{R}^{d \times 1}$ are learnable parameters. $\widetilde{\mathcal{H}}^e$ is the final entity representation enhanced by relation and type.

(a) Generalization on FB15k-237 (Q2B datasets)											FB15k	NELL
Method	1p	2p	3p	2i	3i	pi	ip	2u	up	Avg	Avg	Avg
GQE	41.3	21.5	15.2	26.5	38.5	16.7	8.8	17.1	15.8	22.4	40.1	23.5
+TEMP	47.6	29.6	24.7	36.3	48.4	25.5	13.4	30.2	21.0	30.7	56.6	38.2
Q2B	47.1	24.9	19.4	33.2	46.4	21.8	11.3	25.3	19.3	27.6	51.2	31.1
+TEMP	45.7	27.8	23.4	36.9	49.6	22.9	11.7	27.6	18.9	29.4	55.4	37.3
BETAE	42.6	25.4	21.6	30.2	43.3	20.7	9.2	24.2	18.3	26.2	50.6	33.4
+TEMP	43.3	27.2	22.7	35.3	47.5	24.3	10.6	26.7	18.2	28.4	53.6	34.5
LOGICE	45.6	27.8	24.1	34.7	46.5	23.5	12.0	27.1	20.8	29.1	54.2	39.1
+TEMP	46.6	29.2	25.0	35.8	47.9	24.9	13.5	28.7	21.0	30.3	55.7	39.1
(b) Deductive Reasoning on FB15k-237 (Q2B datasets)											FB15k	NELL
GQE	56.4	30.1	24.5	35.9	51.2	25.1	13.0	25.8	22.0	31.6	43.7	49.8
+TEMP	76.3	48.6	39.0	49.7	60.4	36.9	22.1	59.0	36.3	47.6	71.4	75.5
Q2B	58.5	34.3	28.1	44.7	62.1	23.9	11.7	40.5	22.0	36.2	43.7	51.1
+TEMP	87.2	59.6	47.9	67.2	72.7	49.1	29.8	78.6	43.4	59.5	69.6	90.4
BETAE	77.9	52.6	44.5	59.0	67.8	42.2	23.5	63.7	35.1	51.8	60.6	80.2
+TEMP	84.7	58.3	49.4	62.3	68.8	45.3	28.5	74.5	41.1	57.0	60.6	81.5
LOGICE	81.5	54.2	46.0	58.1	67.1	44.0	28.5	66.6	40.8	54.1	65.5	85.3
+TEMP	84.5	59.8	51.9	59.3	68.1	47.0	33.4	70.8	45.4	57.8	67.1	84.6

Table 1: Hits@3 results on the Q2B datasets testing generalization and deductive reasoning. Please see appendix for full results on FB15k and NELL.

5 Experiments

Our aim is to analyse the impact of adding TEMP to existing QE models on their generalization, deductive and inductive reasoning abilities.

Datasets and Queries. For generalization and deductive reasoning, we use three standard benchmark KGs with official training/validation/test splits: FB15k [Bordes *et al.*, 2013], FB15k-237 [Toutanova and Chen, 2015], and NELL995 (NELL) [Xiong *et al.*, 2017], and two query datasets: one with 9 query structures without negation from Query2Box (Q2B) [Ren *et al.*, 2020] and another with 14 (9 positive + 5 with negation) from BETAE [Ren and Leskovec, 2020]. To test inductive reasoning, we use the datasets FB15k-237-V2 and NELL-V3 provided by GraIL [Teru *et al.*, 2020], ensuring that the test and training sets have a disjoint set of entities. Note that we generate queries over these datasets as done for BETAE datasets. We choose Hit@K and Mean Reciprocal Rank (MRR) as two evaluation metrics².

Baselines. We embed TEMP on four state-of-the-art baselines for complex QA on KGs: Q2B, BETAE, GQE [Hamilton *et al.*, 2018], and LOGICE [Luus *et al.*, 2021].

Generalization. The goal is to find *non-trivial* answers to FOL queries over incomplete KGs, i.e. answers cannot get using subgraph matching. We follow the evaluation protocol in [Ren and Leskovec, 2020]. Briefly, three KGs are built: $\mathcal{G}_{train}$ containing only training triples, $\mathcal{G}_{valid}$ containing $\mathcal{G}_{train}$ plus validation triples, and $\mathcal{G}_{test}$ containing $\mathcal{G}_{valid}$ as well as test triples. The models are trained using $\mathcal{G}_{train}$ to evaluate the generalization ability because queries have at least one edge to predict to find an answer.

²Refer to appendix for details on datasets, queries, and metrics.

Deductive Reasoning. The goal is to test the ability of finding answers only using standard reasoning. Following [Sun *et al.*, 2020], we train models using the full KG ($\mathcal{G}_{train} \cup \mathcal{G}_{valid} \cup \mathcal{G}_{test}$), so only inference (not generalization) is used to get correct answers.

Inductive Reasoning. All baseline models have inductive ability at the query level as they can answer queries with structures that are not seen during training. For example, the Q2B and BETAE datasets consider five query structures during the training and validation phase and four ‘unseen’ structures are used during testing. However, it is not known whether they have *entity-level* inductive ability, i.e. during testing, the query structure has *entities that do not appear* in the training phase. We will analyse this for the first time.

5.1 Main Results

We compare the performance of the four baseline models with their counterparts after adding our TEMP model in four different aspects: 1) generalization, 2) deductive reasoning, 3) inductive reasoning, and 4) queries with negation.

Generalization. Table 1(a) shows that for long-chain queries *2p* and *3p*, the improvement brought by TEMP exceeds that of short-chain *1p*, confirming the suitability of type information for dealing with long-chain queries. In addition, TEMP-enhanced models also achieve improvements on queries *ip/pi/2u/up*, which do not occur in the training KG, showing that type information is also helpful to improve inductive ability at the query-level structure. Notably, for GQE, adding type information can shorten the gap or even surpass the state-of-the-art baseline models (without TEMP).

Deductive Reasoning. Table 1(b) shows that after adding type information, the reasoning ability of the baselines are significantly improved on all datasets consistently. Specifi-

Data	Method	1p	2p	3p	2i	3i	pi	ip	2u	up	Avg
FB15k-237-V2	GQE	0.5	5.5	1.8	0.5	0.6	7.1	13.1	2.1	6.1	4.1
	+TEMP	14.6	22.1	14.1	13.9	14.4	15.7	22.0	9.7	20.1	16.3
	Q2B	0.5	5.5	1.7	0.7	0.7	7.8	12.9	2.7	6.1	4.3
	+TEMP	12.9	12.7	12.3	10.0	10.4	13.7	10.7	7.8	9.5	11.1
	BETAE	0.9	0.5	0.4	0.7	0.3	0.7	0.4	1.0	0.3	0.6
	+TEMP	10.8	15.0	10.6	10.8	10.1	11.3	13.6	5.5	13.6	11.3
	LOGICE	0.7	2.5	1.1	1.0	0.9	1.0	3.1	0.3	1.6	1.4
	+TEMP	15.7	17.1	15.1	14.6	13.8	13.6	16.3	7.0	14.2	14.2
NELL-V3	GQE	0.3	2.3	0.6	0.4	0.1	3.2	4.9	1.8	3.2	1.9
	+TEMP	9.6	5.7	6.0	7.2	8.1	4.7	6.2	4.0	4.0	6.2
	Q2B	0.2	2.2	0.5	0.3	0.2	2.6	4.8	1.8	2.8	1.7
	+TEMP	8.0	5.6	6.0	7.6	7.1	4.6	5.3	4.1	3.2	5.7
	BETAE	0.4	0.1	0.1	0.5	0.1	0.2	0.1	0.3	0.2	0.2
	+TEMP	8.4	4.7	5.7	5.6	5.8	4.1	4.5	4.5	3.0	5.1
	LOGICE	0.2	0.5	0.3	0.1	0.3	0.2	0.3	0.2	0.4	0.3
	+TEMP	10.7	5.0	5.8	7.6	8.1	5.3	5.5	4.7	3.4	6.2

Table 2: Hits@10 results on queries generated from the FB15k-237-V2 and NELL-V3 inductive datasets from GraIL.

Datasets	Model	Our model	2in	3in	inp	pin	pni	Avg
FB15k	BETAE	None +TEMP	14.3 15.2	14.7 15.6	11.5 11.5	6.5 6.8	12.4 13.4	11.8 12.5
	LOGICE	None +TEMP	15.1 15.2	14.2 14.7	12.5 12.7	7.1 7.6	13.4 13.7	12.5 12.8
FB15k-237	BETAE	None +TEMP	5.1 4.3	7.9 8.0	7.4 7.6	3.6 3.5	3.4 2.9	5.4 5.3
	LOGICE	None +TEMP	4.9 5.4	8.2 8.7	7.7 7.9	3.6 4.0	3.5 3.8	5.6 6.0
NELL	BETAE	None +TEMP	5.1 5.1	7.8 7.5	10.0 10.5	3.1 3.1	3.5 3.3	5.9 5.9
	LOGICE	None +TEMP	5.3 5.4	7.5 7.6	11.1 11.3	3.3 3.4	3.8 3.9	6.2 6.3

Table 3: MRR Results on the BETAE Datasets for BETAE and LOGICE on queries with negation. See appendix for results for other query structures.

cally, the improvement of embedding models based on geometric operations (GQE, Q2B) is more significant than that of BETAE or LOGICE. Remarkably, Q2B + TEMP surpasses the state-of-the-art baseline models (without TEMP). The main reason for the modest gain for BETAE and LOGICE is that they impose excessive restrictions on the embedding of entities and relations. For instance, in LOGICE, the logic embeddings with bounded support change the type-enriched vector representations, thus affecting the effect of type information.

Inductive Reasoning. Table 2 shows that the addition of TEMP significantly outperforms all baselines in the (entity-level) inductive setting, by 12.2%, 6.8%, 10.7%, and 12.8%, in terms of Hits@10 in the FB15k-237-V2 dataset. The main reason is that in TEMP entity and relation representations are mainly characterized by the type information. So, newly emerging entities or relations can be captured through their type information, making unnecessary the coupling between entities and entity set, and relations and relation set.

Models	TER	TRR	FB15k-237	FB15k	NELL
GQE	✗	✗	19.4	32.8	19.3
	✓	✗	23.5	42.1	26.8
	✗	✓	26.3	50.8	33.2
	✓	✓	28.2	50.9	34.5
Q2B	✗	✗	24.2	43.4	25.7
	✓	✗	24.7	44.3	29.1
	✗	✓	25.8	47.7	33.4
	✓	✓	26.8	49.7	33.9

Table 4: Average MRR results on the Q2B datasets with TER or TRR. See appendix for detailed results.

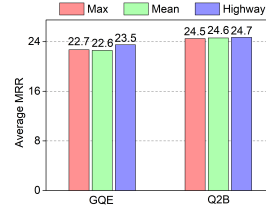


Figure 3: MRR results with different entity type aggregators.

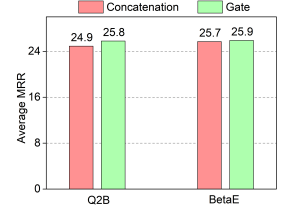


Figure 4: MRR results with different relation type aggregators.

Queries with Negation. Table 3 shows the results of BETAE and LOGICE on queries with negation in the BETAE dataset. The main reason for the small gains is that, unlike BETAE and LOGICE, TEMP does not have specific mechanisms to deal with negation. Specifically, TEMP lacks mechanisms to associate type information to the negation of a relation, i.e., a way to ‘negate’ a type. Boosting queries with negation using type information is left as an interesting future work.

5.2 Ablation Studies

We select GQE and Q2B, as they benefit the most by adding type information, and conduct ablation experiments on the three datasets to study the effect of separately using type-enhancement on entities or relations, see Table 4. We further study different implementations of entity and relation type aggregation models, see Figure 3 and Figure 4.

Type-enhancement on relations is consistently better than on entities, this is explained by the fact that enhancing relation representations is more helpful for queries with long chains of existentially quantified variables as it better deals with cascading errors introduced by relation projections. We also show that using type-enhancement on both entities and relations usually leads to even better performance.

6 Conclusions

We proposed TEMP, a *type-aware plug-and-play* model for answering FOL queries on incomplete KGs. We experimentally show that TEMP can significantly improve four state-of-the-art models in terms of generalization, deductive and inductive reasoning abilities across three benchmark datasets consistently.

Acknowledgments

This work has been supported by the National Key Research and Development Program of China (No.2020AAA0106100), by the National Natural Science Foundation of China (No.61936012), by the Chang Jiang Scholars Program (J2019032), by a Leverhulme Trust Research Project Grant (RPG-2021-140), and by the Centre for Artificial Intelligence, Robotics and Human-Machine Systems (IROHMS) at Cardiff University.

References

- [Arakelyan *et al.*, 2021] Erik Arakelyan, Daniel Daza, Pasquale Minervini, and Michael Cochez. Complex query answering with neural link predictors. In *ICLR*, 2021.
- [Bordes *et al.*, 2013] Antoine Bordes, Nicolas Usunier, Alberto García-Durán, Jason Weston, and Oksana Yakhnenko. Translating embeddings for modeling multi-relational data. In *NeurIPS*, 2013.
- [Chen *et al.*, 2021] Jiajun Chen, Huarui He, Feng Wu, and Jie Wang. Topology-aware correlations between relations for inductive link prediction in knowledge graphs. In *AAAI*, pages 6271–6278, 2021.
- [Choudhary *et al.*, 2021a] Nurendra Choudhary, Nikhil Rao, Sumeet Katariya, Karthik Subbian, and Chandan Reddy. Probabilistic entity representation model for reasoning over knowledge graphs. In *NeurIPS*, 2021.
- [Choudhary *et al.*, 2021b] Nurendra Choudhary, Nikhil Rao, Sumeet Katariya, Karthik Subbian, and Chandan K Reddy. Self-supervised hyperboloid representations from logical queries over knowledge graphs. In *WWW*, 2021.
- [Daza *et al.*, 2021] Daniel Daza, Michael Cochez, and Paul Groth. Inductive entity representations from text via link prediction. In *WWW*, 2021.
- [Hamilton *et al.*, 2018] William L Hamilton, Payal Bajaj, Marinka Zitnik, Dan Jurafsky, and Jure Leskovec. Embedding logical queries on knowledge graphs. In *NeurIPS*, 2018.
- [Luus *et al.*, 2021] Francois Luus, Prithviraj Sen, Pavan Kapani, Ryan Riegel, Ndivhuwo Makondo, Thabang Lebese, and Alexander Gray. Logic embeddings for complex query answering. In *NeurIPS*, 2021.
- [Niu *et al.*, 2020] Guanglin Niu, Bo Li, Yongfei Zhang, Shiliang Pu, and Jingyang Li. Autoeter: Automated entity type representation for knowledge graph embedding. In *EMNLP*, 2020.
- [Pan *et al.*, 2016] J.Z. Pan, G. Vetere, J.M. Gomez-Perez, and H. Wu. *Exploiting Linked Data and Knowledge Graphs for Large Organisations*. Springer, 2016.
- [Pan *et al.*, 2021] Weiran Pan, Wei Wei, and Xianling Mao. Context-aware entity typing in knowledge graphs. In *EMNLP*, 2021.
- [Pan, 2009] Jeff Z. Pan. Resource description framework. In *Handbook on Ontologies*, pages 71–90. 2009.
- [Ren and Leskovec, 2020] Hongyu Ren and Jure Leskovec. Beta embeddings for multi-hop logical reasoning in knowledge graphs. In *NeurIPS*, 2020.
- [Ren *et al.*, 2020] Hongyu Ren, Weihua Hu, and Jure Leskovec. Query2box: Reasoning over knowledge graphs in vector space using box embeddings. In *ICLR*, 2020.
- [Srivastava *et al.*, 2015] Rupesh Kumar Srivastava, Klaus Greff, and Jürgen Schmidhuber. Highway networks. *arXiv preprint arXiv:1505.00387*, 2015.
- [Sun *et al.*, 2020] Haitian Sun, Andrew O. Arnold, Tania Bedrax-Weiss, Fernando Pereira, and William W. Cohen. Faithful embeddings for knowledge base queries. In *NeurIPS*, 2020.
- [Tai *et al.*, 2015] Kai Sheng Tai, Richard Socher, and Christopher D. Manning. Improved semantic representations from tree-structured long short-term memory networks. In *ACL*, pages 1556–1566, 2015.
- [Teru *et al.*, 2020] Komal Teru, Etienne Denis, and Will Hamilton. Inductive relation prediction by subgraph reasoning. In *ICML*, pages 9448–9457, 2020.
- [Toutanova and Chen, 2015] Kristina Toutanova and Danqi Chen. Observed versus latent features for knowledge base and text inference. In *ACL*, pages 57–66, 2015.
- [Wang *et al.*, 2021] Hongwei Wang, Hongyu Ren, and Jure Leskovec. Relational message passing for knowledge graph completion. In *SIGKDD*, pages 1697–1707, 2021.
- [Wiharja *et al.*, 2020] Kemas Wiharja, Jeff Z. Pan, Martin J. Kollingbaum, and Yu Deng. Schema Aware Iterative Knowledge Graph Completion. *Journal of Web Semantics*, 2020.
- [Xiong *et al.*, 2017] Wenhan Xiong, Thien Hoang, and William Yang Wang. Deeppath: A reinforcement learning method for knowledge graph reasoning. In *EMNLP*, 2017.
- [Yao *et al.*, 2019] Liang Yao, Chengsheng Mao, and Yuan Luo. Kg-bert: Bert for knowledge graph completion. *arXiv preprint arXiv:1909.03193*, 2019.
- [Zhang *et al.*, 2020] Shuailiang Zhang, Hai Zhao, Yuwei Wu, Zhuosheng Zhang, Xi Zhou, and Xiang Zhou. Dcmn+: Dual co-matching network for multi-choice reading comprehension. In *AAAI*, 2020.
- [Zhang *et al.*, 2021] Zhanqiu Zhang, Jie Wang, Jiajun Chen, Shuiwang Ji, and Feng Wu. Cone: Cone embeddings for multi-hop reasoning over knowledge graphs. In *NeurIPS*, 2021.
- [Zhao *et al.*, 2020] Yu Zhao, Anxiang Zhang, Ruobing Xie, Kang Liu, and Xiaojie Wang. Connecting embeddings for knowledge graph entity typing. In *ACL*, 2020.