



Knowledge Enhanced Graph Neural Networks for Graph Completion

Luisa Werner, Nabil Layaïda, Pierre Genevès, Sarah Chlyah

► To cite this version:

Luisa Werner, Nabil Layaïda, Pierre Genevès, Sarah Chlyah. Knowledge Enhanced Graph Neural Networks for Graph Completion. DSAA 2023 - The 10th IEEE International Conference on Data Science and Advanced Analytics, Oct 2023, Thessaloniki, Greece. hal-04041691v3

HAL Id: hal-04041691

<https://inria.hal.science/hal-04041691v3>

Submitted on 30 Aug 2023

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License

Knowledge Enhanced Graph Neural Networks

^{1st} Luisa Werner
Université Grenoble Alpes, INRIA
Grenoble, France
luisa.werner@inria.fr

^{2nd} Nabil Layaïda
INRIA
Grenoble, France
nabil.layaïda@inria.fr

^{3rd} Pierre Genevès
CNRS, INRIA
Grenoble, France
pierre.geneves@inria.fr

^{4th} Sarah Chlyah
INRIA
Grenoble, France
sarah.chlyah@inria.fr

Abstract—Graph data is omnipresent and has a wide variety of applications, such as in natural science, social networks, or the semantic web. However, while being rich in information, graphs are often noisy and incomplete. As a result, graph completion tasks, such as node classification or link prediction, have gained attention. On one hand, neural methods, such as graph neural networks, have proven to be robust tools for learning rich representations of noisy graphs. On the other hand, symbolic methods enable exact reasoning on graphs. We propose Knowledge Enhanced Graph Neural Networks (KeGNN), a neuro-symbolic framework for graph completion that combines both paradigms as it allows for the integration of prior knowledge into a graph neural network model. Essentially, KeGNN consists of a graph neural network as a base upon which knowledge enhancement layers are stacked with the goal of refining predictions with respect to prior knowledge. We instantiate KeGNN in conjunction with two well-known graph neural networks, Graph Convolutional Networks and Graph Attention Networks, and evaluate KeGNN on multiple benchmark datasets for node classification.

Index Terms—neuro-symbolic integration, graph neural networks, relational learning, knowledge graphs, fuzzy logic

I. INTRODUCTION

Graphs are ubiquitous across diverse real-world applications such as e-commerce [1], natural science [2] or social networks [3]. Graphs connect nodes by edges and allow to enrich them with features. This makes them a versatile and powerful data structure that encodes relational information. As graphs are often derived from noisy data, incompleteness and errors are common issues. Consequently, graph completion tasks such as node classification or link prediction have become increasingly important. These tasks are approached from different directions. In the field of deep learning, research on graph neural networks (GNNs) has gained momentum. Numerous models have been proposed for various graph topologies and applications [4]–[7]. The key strength of GNNs is to find meaningful representations of noisy graph data, that can be used to improve prediction tasks [8]. Despite this advantage, as a subcategory of deep learning methods, GNNs are criticized for their limited interpretability and large data consumption [9]. Alongside, the research field of symbolic AI addresses the above-mentioned tasks. In symbolic AI, solutions are found by performing logic-like reasoning steps that are exact, interpretable and data-efficient [10]. For large graphs, however, symbolic methods are often computationally expensive or even infeasible. Since techniques from deep learning and from symbolic AI have complementary pros and cons, the field of

neuro-symbolic AI aims to combine both paradigms. Neuro-symbolic AI not only paves the way towards the application of AI to learning with limited data, but also allows for jointly using symbolic information (in the form of logical rules) and sub-symbolic information (in the form of real-valued data). This helps to overcome the black-box nature of deep learning methods and to improve interpretability through symbolic representations [9], [11], [12].

In this work, we present the neuro-symbolic approach Knowledge enhanced Graph Neural Networks (KeGNN) to conduct node classification given graph data and a set of prior knowledge. In KeGNN, knowledge enhancement layers [13] are stacked on top of a GNN and adjust its predictions in order to increase the satisfaction of some prior knowledge. In addition to the parameters of the GNN, the knowledge enhancement layers contain learnable clause weights that reflect the impact of the prior knowledge on the predictions. Both components form an end-to-end differentiable model. KeGNN can be seen as a variant of knowledge enhanced neural networks (KENN), which stack knowledge enhancement layers onto a multi-layer perceptron (MLP) and have been proven successful in semantic point cloud segmentation, image segmentation and multi-label classification [14]–[16]. However, relational information in sparse graphs can only be introduced through the logical clauses with binary predicates in the knowledge enhancement layer and not at base neural network level. In contrast, KeGNN is based on GNNs that process the graph structure, which makes both the neural and symbolic components sufficiently powerful to exploit the graph structure. In this work, we instantiate KeGNN in conjunction with two well-known GNNs: Graph Attention Networks [17] and Graph Convolutional Networks [18]. We apply KeGNN to the benchmark datasets for node classification Cora, Citeseer, PubMed [19] and Flickr [20].

II. METHOD: KEGNN

KeGNN is a neuro-symbolic approach that can be applied to node classification tasks with the capacity of handling graph structure at the base neural network level. The model takes two types of input: (1) real-valued graph data and (2) prior knowledge expressed in first-order logic.

A. Graph-structured Data

A Graph $G = (N, E)$ consists of a set of n nodes N and a set of k edges E where each edge of the form (v_i, v_j)

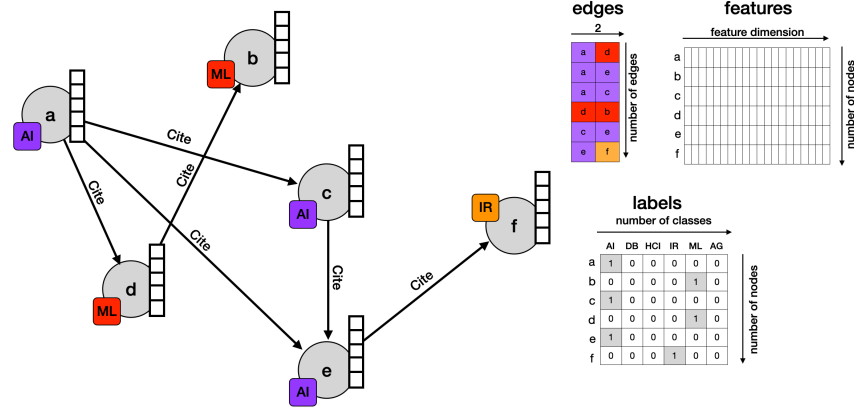


Fig. (1) Example extract of the Citeseer citation graph.

connects two nodes $v_i \in \mathcal{N}$ and $v_j \in \mathcal{N}$. The neighborhood $\mathcal{N}(v_i)$ describes the set of first-order neighbors of v_i . For an *attributed* and *labelled* graph, nodes are enriched with features and labels. Each node has a feature vector $\mathbf{x} \in \mathbb{R}^d$ of dimension d and a label vector $\mathbf{y} \in \mathbb{R}^m$. The label vector $\mathbf{y}$ contains one-hot encoded ground truth labels for m classes. In matrix notation, the features and labels of the entire graph are described as $\mathbf{X} \in \mathbb{R}^{n \times d}$ and $\mathbf{Y} \in \mathbb{R}^{n \times m}$. A graph is *typed* if type functions f_E and f_N assign edge types and node types to the edges and nodes, respectively. A graph with constant type functions (that assign the same edge and node type to all edges and nodes) is called *homogeneous*, whereas for *heterogeneous* graphs, nodes and edges may have different types [5].

Example 2.1: A Citation Graph $\mathbf{G}_{\text{Cit}}$ consists of documents and citations. Figure 1 shows an extract of the Citeseer citation graph that is used as example to guide through this paper. The documents are represented by a set of nodes $\mathcal{N}_{\text{Cit}}$ and citations by a set of edges $\mathcal{E}_{\text{Cit}}$. Documents can be attributed with features $\mathbf{X}_{\text{Cit}}$ that describe their content as Word2Vec [21] vectors. Each node is labelled with one of the six topic categories $\{\text{AI}, \text{DB}, \text{HCI}, \text{IR}, \text{ML}, \text{AG}\}^1$ that are encoded in $\mathbf{Y}_{\text{Cit}}$. Since all nodes (documents) and edges (citations) have the same type, $\mathbf{G}_{\text{Cit}}$ is homogeneous.

B. Prior Knowledge

Some prior knowledge $\mathcal{K}$ is provided to KeGNN. It can be described as a set of ℓ logical clauses expressed in the logical language $\mathcal{L}$ that is defined as sets of constants $\mathcal{C}$, variables $\mathcal{X}$ and predicates $\mathcal{P}$. Predicates have an arity r of one (unary) or two (binary): $\mathcal{P} = \mathcal{P}_U \cup \mathcal{P}_B$. Predicates of arity $r > 2$ are not considered in this work. Unary predicates express properties, whereas binary predicates express relations. $\mathcal{L}$ supports negation ($\neg$) and disjunction ($\vee$). Each clause $\varphi \in \mathcal{K} = \{\varphi_1, \dots, \varphi_\ell\}$ can be formulated as a disjunction of (possibly negated) atoms $\bigvee_{j=1}^q o_j$ with q atoms $\{o_1, \dots, o_q\}$.

¹The classes are abbreviations for the categories *Artificial Intelligence*, *Databases*, *Human-Computer Interaction*, *Information Retrieval*, *Machine Learning* and *Agents*.

Since the prior knowledge is general, all clauses are assumed to be universally quantified. Clauses can be *grounded* by assigning constants to the free variables. A grounded clause is denoted as $\varphi[x_1, x_2, \dots | c_1, c_2, \dots]$ with variables $x_i \in \mathcal{X}$ and constants $c_i \in \mathcal{C}$. The set of all grounded clauses in a graph is $\mathcal{G}(\mathcal{K}, \mathcal{C})$.

Example 2.2: The graph $\mathbf{G}_{\text{Cit}}$ in Figure 1 can be expressed in $\mathcal{L}$. Nodes are represented by a set of constants $\mathcal{C} = \{a, b, \dots, f\}$. Node labels are expressed as a set of unary predicates $\mathcal{P}_U = \{\text{AI}, \text{DB}, \dots, \text{AG}\}$ and edges as a set of binary predicates $\mathcal{P}_B = \{\text{Cite}\}$. $\mathcal{L}$ has a set of variables $\mathcal{X} = \{x, y\}$. The atom $\text{AI}(x)$, for example, expresses the membership of x to the class AI and $\text{Cite}(x, y)$ expresses the existence of a citation between x and y . Some prior knowledge $\mathcal{K}$ can be written as a set of $\ell = 6$ disjunctive clauses in $\mathcal{L}$. Here, the assumption is denoted that two papers that cite each other have the same document class:

$$\begin{aligned} \varphi_{\text{AI}} : \forall xy \neg \text{AI}(x) \vee \neg \text{Cite}(x, y) \vee \text{AI}(y) \\ \varphi_{\text{DB}} : \forall xy \neg \text{DB}(x) \vee \neg \text{Cite}(x, y) \vee \text{DB}(y) \\ \dots \end{aligned}$$

The atoms are grounded by replacing the variables x and y with the constants $\{a, b, \dots, f\}$ to obtain the sets of unary groundings $\{\text{AI}(a), \text{ML}(b), \dots, \text{IR}(f)\}$ and binary groundings $\{\text{Cite}(a, d), \text{Cite}(a, e), \dots, \text{Cite}(a, f)\}$. Assuming a closed world and exclusive classes, other facts could be derived, such as $\{\neg \text{DB}(a), \neg \text{IR}(a), \dots, \neg \text{Cite}(a, b)\}$. For the sake of simplicity, these are omitted here.

C. Node Classification

Node classification is a subtask of knowledge graph completion on a graph $\mathbf{G}$ with the objective to assign classes to nodes where they are unknown. This task is accomplished given node features $\mathbf{X}$, edges $\mathbf{E}$ and some prior knowledge $\mathcal{K}$ encoded as a set of clauses in $\mathcal{L}$. A predictive model is trained on a subset of the graph $\mathbf{G}_{\text{train}}$ with ground truth labels $\mathbf{Y}_{\text{train}}$ and validated on a test set $\mathbf{G}_{\text{test}}$ for which the ground truth

labels are compared to the predictions in order to assess the predictive performance. Node classification can be studied in a *transductive* or *inductive* setting. In a transductive setting, the entire graph is available for training, but the true labels of the test nodes are masked. In an inductive setting, only the nodes in the training set and the edges connecting them are available, making it more challenging to classify unseen nodes.

D. Fuzzy Semantics

Let us consider an attributed and labelled graph $\mathbf{G}$ and the prior knowledge $\mathcal{K}$. While $\mathcal{K}$ can be defined in the logic language $\mathcal{L}$, the neural component in KeGNN relies on continuous and differentiable representations. To interpret Boolean logic in the real-valued domain, KeGNN uses fuzzy logic [22], which maps Boolean truth values to the continuous interval $[0, 1] \subset \mathbb{R}$. A constant in $\mathcal{C}$ is interpreted as a real-valued feature vector $\mathbf{x} \in \mathbb{R}^d$. A predicate $P \in \mathcal{P}$ with arity r is interpreted as a function $f_P : \mathbb{R}^{r \times d} \mapsto [0, 1]$ that takes r feature vectors as input and returns a truth value.

Example 2.3: In the example, a unary predicate $P_U \in \mathcal{P}_U = \{\text{AI}, \text{DB}, \dots\}$ is interpreted as a function $f_{P_U} : \mathbb{R}^d \mapsto [0, 1]$ that takes a feature vector $\mathbf{x}$ and returns a truth value indicating whether the node belongs to the class encoded as P_U . The binary predicate $\text{Cite} \in \mathcal{P}_B$ is interpreted as the function

$$f_{\text{Cite}}(v_i, v_j) = \begin{cases} 1, & \text{if } (v_i, v_j) \in \mathbf{E}_{\text{Cit}} \\ 0, & \text{else.} \end{cases}$$

f_{Cite} returns the truth value 1 if there is an edge between two nodes v_i and v_j in $\mathbf{G}_{\text{Cit}}$ and 0 otherwise.

T-conorm functions $\perp : [0, 1] \times [0, 1] \mapsto [0, 1]$ [23] take real-valued truth values of two literals² and define the truth value of their disjunction. The Gödel t-conorm function for two truth values $\mathbf{t}_i, \mathbf{t}_j$ is defined as

$$\perp(\mathbf{t}_i, \mathbf{t}_j) \mapsto \max(\mathbf{t}_i, \mathbf{t}_j).$$

To obtain the truth value of a clause $\varphi : o_1 \vee \dots \vee o_q$, the function $\perp$ is extended to a vector $\mathbf{t}$ of q truth values: $\perp(\mathbf{t}_1, \mathbf{t}_2, \dots, \mathbf{t}_q) = \perp(\mathbf{t}_1, \perp(\mathbf{t}_2, \dots, \perp(\mathbf{t}_{q-1}, \mathbf{t}_q)))$. Fuzzy negation over truth values is defined as $\mathbf{t} \mapsto 1 - \mathbf{t}$ [22].

Example 2.4: Given the clause $\varphi_{\text{AI}} : \forall xy \neg \text{AI}(x) \vee \neg \text{Cite}(x, y) \vee \text{AI}(y)$ and its grounding $\varphi_{\text{AI}}[x, y|a, b] : \text{AI}(a) \vee \neg \text{Cite}(a, b) \vee \text{AI}(b)$ to the constants a and b and truth values for the grounded predicates $\text{AI}(a) = \mathbf{t}_1$, $\text{AI}(b) = \mathbf{t}_2$ and $\text{Cite}(a, b) = \mathbf{t}_3$, the truth value of $\varphi_{\text{AI}}[x, y|a, b]$ is $\max\{\max\{(1 - \mathbf{t}_1), (1 - \mathbf{t}_3)\}, \mathbf{t}_2\}$.

E. Model Architecture

The way KeGNN computes the final predictions can be divided in two stages. First, a GNN predicts the node classes given the features and the edges. Subsequently, the knowledge enhancement layers use the predictions as truth values for the grounded unary predicates and update them with respect to the knowledge. An overview of KeGNN is given in Figure 2.

²A literal is a (possibly negated) grounded atom, e.g. $\text{AI}(a)$

1) Neural Component: The role of the GNN in the neural component is to exploit feature information in the graph structure. The key strength of a GNN is to enrich node representations with graph structure by nesting k message passing layers [8]. Per layer, the representations of neighboring nodes are aggregated and combined to obtain updated representations. The node representation v_i^{k+1} in the k -th message passing layer is

$$v_i^{k+1} = \text{combine}(v_i^k, \text{aggregate}(\{v_j^k | v_j^k \in \mathcal{N}(v_i)\})).$$

The layers contain learnable parameters that are optimized with backpropagation. In this work, we consider two well-known GNNs as components for KeGNN: Graph Convolutional Networks (GCN) [18] and Graph Attention Networks (GAT) [17]. While GCN considers the graph structure as given, GAT allows for assessing the importance of the neighbors with attention weights α_{ij} between node v_i and node v_j . In case of multi-head attention, the attention weights are calculated multiple times and concatenated which allows for capturing different aspects of the input data. In KeGNN, the GNN implements the functions f_{P_U} (see Section II-D). In other words, the predictions are used as truth values for the grounded unary predicates in the symbolic component.

2) Symbolic Component: To refine the predictions of the GNN, one or more knowledge enhancement layers are stacked onto the GNN to update its predictions $\mathbf{Y}$ to $\mathbf{Y}'$. The goal is to increase the satisfaction of the prior knowledge. The predictions $\mathbf{Y}$ of the GNN serve as input to the symbolic component where they are interpreted as fuzzy truth values for the unary grounded predicates $\mathbf{U} := \mathbf{Y}$ with $\mathbf{U} \in \mathbb{R}^{n \times m}$. Fuzzy truth values for the groundings of binary predicates are encoded as a matrix $\mathbf{B}$ where each row represents an edge (v_i, v_j) and each column represents an edge type e . In the context of node classification, the GNN returns only predictions for the node classes, while the edges are assumed to be given. A binary grounded predicate is therefore set to truth value 1 (true) if an edge between two nodes v_i and v_j exists:

$$\mathbf{B}_{[(v_i, v_j), e]} = \begin{cases} 1, & \text{if } (v_i, v_j) \text{ of type } e \in \mathbf{E} \\ 0, & \text{else.} \end{cases}$$

Example 2.5: In case of the beforementioned citation graph of Figure 1, $\mathbf{U}$ and $\mathbf{B}$ are defined as:

$$\mathbf{U} := \begin{bmatrix} \text{AI}(a) & \dots & \text{AG}(a) \\ \text{AI}(b) & \dots & \text{AG}(b) \\ \vdots & & \vdots \\ \text{AI}(f) & \dots & \text{AG}(f) \end{bmatrix} \quad \mathbf{B} := \begin{bmatrix} \text{Cite}(a, d) \\ \text{Cite}(a, e) \\ \text{Cite}(a, c) \\ \vdots \\ \text{Cite}(c, e) \\ \text{Cite}(e, f) \end{bmatrix}$$

To enhance the satisfaction of clauses that contain both unary and binary predicates, their groundings are joined into one matrix $\mathbf{M} \in \mathbb{R}^{k \times p}$ with $p = 2 \cdot |\mathcal{P}_U| + |\mathcal{P}_B|$. $\mathbf{M}$ is computed by joining $\mathbf{U}$ and $\mathbf{B}$ so that each row of $\mathbf{M}$ represents an edge

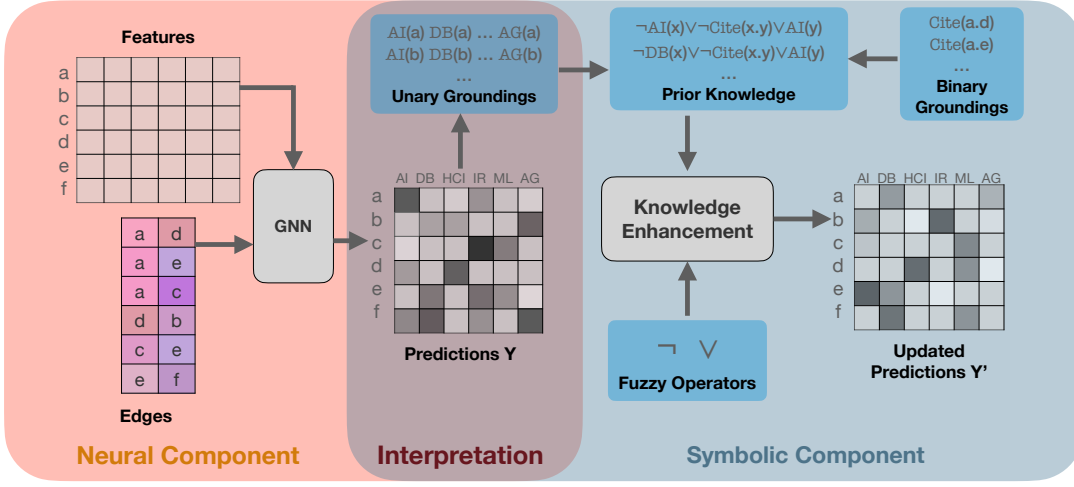


Fig. (2) Overview of KeGNN.

(v_i, v_j) . As a result, $\mathbf{M}$ contains all required grounded unary predicates for v_i and v_j .

Example 2.6: For the example citation graph, we obtain $\mathbf{M}$ as follows:

$$\mathbf{M} = \begin{array}{c|ccc|ccc|ccc} \text{groundings unary predicates for } x & & & & \text{groundings unary predicates for } y & & & \text{groundings Cite}(x,y) & & & \\ \hline \text{AI}^x(a) & \text{DB}^x(a) & \dots & \text{AG}^x(a) & \text{AI}^y(d) & \text{DB}^y(d) & \dots & \text{AG}^y(d) & \text{Cit}(a,d) & & \\ \text{AI}^x(a) & \text{DB}^x(a) & \dots & \text{AG}^x(a) & \text{AI}^y(e) & \text{DB}^y(e) & \dots & \text{AG}^y(e) & \text{Cit}(a,e) & & \\ \text{AI}^x(a) & \text{DB}^x(a) & \dots & \text{AG}^x(a) & \text{AI}^y(c) & \text{DB}^y(c) & \dots & \text{AG}^y(c) & \text{Cit}(a,c) & & \\ \vdots & \vdots & & \vdots & \vdots & \vdots & & \vdots & \vdots & & \\ \text{AI}^x(c) & \text{DB}^x(c) & \dots & \text{AG}^x(c) & \text{AI}^y(e) & \text{DB}^y(e) & \dots & \text{AG}^y(e) & \text{Cit}(c,e) & & \\ \text{AI}^x(e) & \text{DB}^x(e) & \dots & \text{AG}^x(e) & \text{AI}^y(f) & \text{DB}^y(f) & \dots & \text{AG}^y(f) & \text{Cit}(e,f) & & \end{array}$$

A knowledge enhancement layer consists of multiple *clause enhancers*. A clause enhancer is instantiated for each clause $\varphi \in \mathcal{K}$. Its aim is to compute updates $\delta\mathbf{M}_\varphi$ for the groundings in $\mathbf{M}$ that increase the satisfaction of φ .

First, fuzzy negation is applied to the columns of $\mathbf{M}$ that correspond to negated atoms in φ . Then $\delta\mathbf{M}_\varphi$ is computed by a *t-conorm boost function* ϕ [13]. This function $\phi : [0, 1]^q \mapsto [0, 1]^q$ takes q truth values and returns changes to those truth values such that the satisfaction is increased: $\perp(\mathbf{t}) \leq \perp(\mathbf{t} + \phi(\mathbf{t}))$. [13] propose the following differentiable t-conorm boost function

$$\phi_{w_\varphi}(\mathbf{t})_i = w_\varphi \cdot \frac{e^{\mathbf{t}_i}}{\sum_{j=1}^q e^{\mathbf{t}_j}}.$$

The boost function ϕ_{w_φ} employs a clause weight w_φ that is initialized in the beginning of the training and optimized during training as a learnable parameter. The updates for the groundings calculated by ϕ_{w_φ} are proportional to w_φ . Therefore, w_φ determines the magnitude of the update and thus reflects the impact of a clause. The changes to the atoms that do not appear in a clause are set to zero. The boost function is applied row-wise to $\mathbf{M}$ as illustrated in the following example.

Example 2.7: Given the clause $\varphi_{AI} : \forall xy \neg \text{AI}(x) \vee \neg \text{Cite}(x,y) \vee \text{AI}(y)$ and the clause weight w_{AI} , the changes

for this clause are $\delta\mathbf{M}_{\varphi_{AI}} =$

$$w_{AI} \cdot \begin{bmatrix} \delta_{\neg \text{AI}^x(a)} & 0 & \dots & \delta_{\text{AI}^y(c)} & 0 & \dots & \delta_{\neg \text{Cit}(a,c)} \\ \delta_{\neg \text{AI}^x(a)} & 0 & \dots & \delta_{\text{AI}^y(e)} & 0 & \dots & \delta_{\neg \text{Cit}(a,e)} \\ \delta_{\neg \text{AI}^x(a)} & 0 & \dots & \delta_{\text{AI}^y(d)} & 0 & \dots & \delta_{\neg \text{Cit}(a,d)} \\ \vdots & \vdots & & \vdots & \vdots & & \vdots \\ \delta_{\neg \text{AI}^x(e)} & 0 & \dots & \delta_{\text{AI}^y(f)} & 0 & \dots & \delta_{\neg \text{Cit}(e,f)} \end{bmatrix}$$

The values of $\delta\mathbf{M}_{\varphi_{AI}}$ are calculated by $\phi_{w_{AI}}$, for example:

$$\delta_{\neg \text{AI}^x(a)} = \phi_{w_{AI}}(\mathbf{z})_a = -\frac{e^{-\mathbf{z}_{AI(a)}}}{e^{-\mathbf{z}_{AI(a)}} + e^{-\mathbf{z}_{\text{Cit}(a,c)}} + e^{\mathbf{z}_{AI(c)}}}$$

Each clause enhancer computes updates $\delta\mathbf{M}_\varphi$ to increase the satisfaction of a clause independently. The updates of all clause enhancers are finally added, resulting in a matrix $\delta\mathbf{M} = \sum_{\varphi \in \mathcal{K}} \delta\mathbf{M}_\varphi$. To apply the updates to the initial predictions, $\delta\mathbf{M}$ has to be added to $\mathbf{Y}$. The updates in $\delta\mathbf{M}$ can not directly be applied to the predictions $\mathbf{Y}$ of the GNN. Since the unary groundings $\mathbf{U}$ were joined with the binary groundings $\mathbf{B}$, multiple changes may be proposed for the same grounded unary atom. For example, for the grounded atom $\text{AI}(c)$ the changes $\delta_{\neg \text{AI}^y(c)}$ and $\delta_{\neg \text{AI}^x(c)}$ are proposed, since c appears in the grounded clauses $\varphi_{AI}[x, y|a, c]$ and $\varphi_{AI}[x, y|c, e]$. In $\mathbf{G}_{\text{Cit}}$ the node c appears in first place of edge (a, c) and in second place of edge (c, e) . Therefore, all updates for the same grounded atom are summed, which reduces the size of $\mathbf{M}$ to the size of $\mathbf{U}$.

To ensure that the updated predictions remain truth values in the range of $[0, 1]$, the knowledge enhancement layer updates at first the preactivations $\mathbf{Z}$ of the GNN and then applies the activation function σ to the updated preactivations $\mathbf{Z}'$ in order to obtain the final predictions: $\mathbf{Y}' = \sigma(\mathbf{Z}')$. Therefore, a knowledge enhancement layer transforms $\mathbf{Z}$ to $\mathbf{Z}'$ (with $\mathbf{Z}, \mathbf{Z}' \in \mathbb{R}^{n \times m}$). In the last step, the updates by the knowledge enhancer are added to the preactivations $\mathbf{Z}$ of the GNN and passed to σ to obtain the updated predictions

$$\mathbf{Y}' = \sigma\left(\mathbf{Z} + \sum_{\varphi \in \mathcal{K}} \delta\mathbf{U}_\varphi\right)$$

where δU_φ is the matrix obtained by extracting the changes to the unary predicates from δM_φ . Regarding the binary groundings, the values in $\mathbf{B}$ are set to a high positive value that results in one when σ is applied.

III. RELATED WORK

The field of knowledge graph completion is addressed from several research directions. Symbolic methods exist that conduct link prediction given a set of prior knowledge [24] [25]. Embedding-based methods [26] are mostly sub-symbolic methods to obtain node embeddings that are used for knowledge graph completion tasks. Usually, their common objective is to find similar embeddings for nodes that are located closely in the graph. The majority of these methods only encodes the graph structure, but does not consider node-specific feature information [27]. However, KeGNN is based on GNNs that are suited for learning representations of graphs attributed with node features. It stacks additional layers that interpret the outputs of the GNN in fuzzy logic and modify them to increase the satisfiability. Therefore, it is considered a neuro-symbolic method. In the multifaceted neuro-symbolic field, KeGNN can be placed in the category of knowledge-guided learning [13], where the focus lies on learning in the presence of additional supervision introduced as prior knowledge. Within this category, KeGNN belongs to the model-based approaches, where prior knowledge in the form of knowledge enhancement layers is an integral part of the model [14]. Beyond, loss-based methods such as logic tensor networks [28] exist that encode the satisfiability of prior knowledge as an optimization objective.

Further, in [29] neuro-symbolic approaches dealing with graph structures are classified into three categories. First, logically informed embedding approaches [30], [31] use predefined logical rules that provide knowledge to a neural system, while both components are mostly distinct. Second, approaches for knowledge graph embedding with logical constraints [32], [33] use prior knowledge as constraints on the neural knowledge graph embedding method in order to modify predictions or embeddings. Thirdly, neuro-symbolic methods are used for learning rules for graph reasoning tasks [34], [35]. This allows for rule generation or confidence scores for prior knowledge and makes the models robust to exceptions or soft knowledge. KeGNN best falls into the second category, since the prior knowledge is interpreted in fuzzy logic to be integrated with the neural model and update the GNN's predictions. The idea of confidence values in category three shares the common property of weighting knowledge as with KeGNN's clause weights. However, even though KeGNN's clause weights introduce a notion of impact of a clause when predictions are made, they cannot directly be interpreted as the confidence in a rule.

In the well-known *Kautz Taxonomy* [36] that classifies neuro-symbolic approaches according to the integration of neural and symbolic modules, KeGNN falls best into the category *Neuro[Symbolic]* (Type 6) of fully-integrated

neuro-symbolic systems that embed symbolic reasoning in a neural architecture.

IV. EXPERIMENTS

To evaluate the performance of KeGNN, we apply it to the datasets Citeseer, Cora, PubMed and Flickr that are common benchmarks for node classification in a transductive setting. In the following, KeGNN is called KeGCN and KeGAT when instantiated to a GCN or a GAT, respectively. As additional baseline, we consider KeMLP, that stacks knowledge enhancement layers onto an MLP, as proposed in [14]. Further, the standalone neural models MLP, GCN and GAT are used as baselines. While Citeseer, Cora and PubMed are citation graphs that encode citations between scientific papers (as in Example 2.2), Flickr contains images and shared properties between them. All datasets can be modelled as homogeneous, labelled and attributed graphs as defined in Section II-A. Table I gives an overview of the named datasets in this work. The datasets are publicly available on the dataset collection³ of PyTorch Geometric [37]. For the split into train, valid and test set, we take the predefined splits in [38] for the citation graphs and in [20] for Flickr. Word2Vec vectors [21] are used as node features for the citation graphs and image data for Flickr. Figure 1 visualizes the graph structure of the underlying datasets in this work as a homogeneous, attributed and labelled graph on the example of Citeseer.

The set of prior logic for the knowledge enhancement layers is manually defined. In this work, we encode the assumption that the existence of an edge for a node pair points to their membership to the same class and hence provides added value to the node classification task. In the context of citation graphs, this implies that two documents that cite each other refer to the same topic, while for Flickr, linked images share the same properties. Following this pattern for all datasets, a clause $\varphi: \forall xy : \neg \text{Cls}_i(x) \vee \neg \text{Link}(x, y) \vee \text{Cls}_i(y)$ is instantiated for each node class $\text{Cls}_i, i \in \{1, \dots, m\}$. More details on the experiments are given in Section IV-B. The source code of the experiments are publicly available⁴.

A. Results

To compare the performance of all models, we examine the average test accuracy over 50 runs (10 for Flickr) for the knowledge enhanced models KeMLP, KeGCN, KeGAT and the standalone base models MLP, GCN, GAT on the named datasets. The results are given in Table II. For Cora and Citeseer, KeMLP leads to a significant improvement over MLP (p-value of one-sided t-test $\ll 0.05$). In contrast, no significant advantage of KeGCN or KeGAT in comparison to the standalone base model is observed. Nevertheless, all GNN-based models are significantly superior to KeMLP for Cora. This includes not only KeGCN and KeGAT, but also the GNN baselines. For Citeseer, KeGAT and GAT both outperform KeMLP. In the case of PubMed, only a significant improvement of KeMLP over MLP can be observed, while the

³<https://pytorch-geometric.readthedocs.io/en/latest/modules/datasets.html>

⁴<https://gitlab.inria.fr/tyrex/kegenn>

Name	#nodes	#edges	#features	#Classes	train/valid/test split
Citeseer	3,327	9,104	3,703	6	1817/500/1000
Cora	2,708	10,556	1,433	7	1208/500/1000
PubMed	19,717	88,648	500	3	18217/500/1000
Flickr	89,250	899,756	500	7	44624/22312/22312

TABLE (I) Overview of the datasets Citeseer, Cora, PubMed and Flickr

	MLP	KeMLP	GCN	KeGCN	GAT	KeGAT
Cora	0.7098 (0.0080)	0.8072 (0.0193)	0.8538 (0.0057)	0.8587 (0.0057)	0.8517 (0.0068)	0.8498 (0.0066)
CiteSeer	0.7278 (0.0081)	0.7529 (0.0067)	0.748 (0.0102)	0.7506 (0.0096)	0.7718 (0.0072)	0.7734 (0.0073)
PubMed	0.8844 (0.0057)	0.8931 (0.0048)	0.8855 (0.0062)	0.8840 (0.0087)	0.8769 (0.0040)	0.8686 (0.0081)
Flickr	0.4656 (0.0018)	0.4659 (0.0012)	0.5007 (0.0063)	0.4974 (0.0180)	0.4970 (0.0124)	0.4920 (0.0189)

TABLE (II) Average test accuracy of 50 runs (10 for Flickr). The standard deviations are reported in brackets.

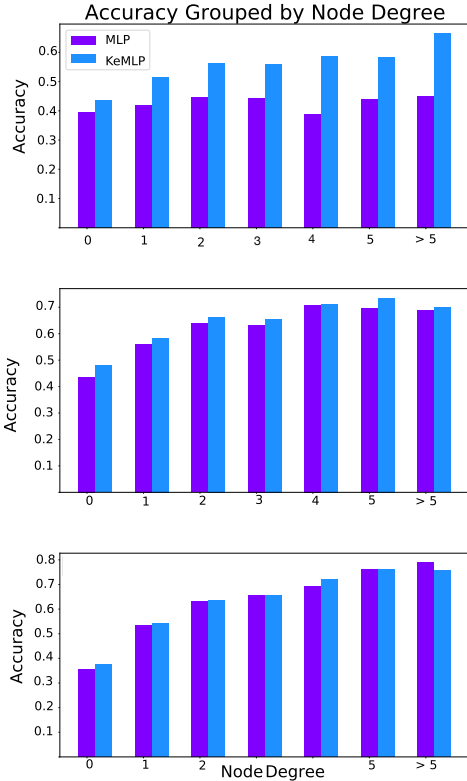


Fig. (3) The accuracy grouped by the node degree for MLP vs. KeMLP (above) and GCN vs. KeGCN (middle) and GAT vs. KeGAT (below) on Citeseer.

GNN-based models and their enhanced versions do not provide any positive effect. For Flickr, no significant improvement between the base model and the respective knowledge enhanced model can be observed. Nevertheless, all GNN-based models outperform KeMLP, reporting significantly higher mean test accuracies for KeGAT, GAT, GCN and KeGCN.

1) *Exploitation of the Graph Structure*: It turns out that the performance gap between MLP and KeMLP is larger than for KeGNN in comparison to the standalone GNN. To explain this observation, we examine how the graph structure affects

the prediction performance. Therefore, in Figure 3 we analyze the accuracy grouped by the node degree for the entire graph for MLP vs. KeMLP and GCN vs. KeGCN. The findings for KeGAT are in line with those for KeGCN. It is observed that KeMLP performs better compared to MLP as the node degree increases. By contrast, when comparing GCN and KeGCN, for both models, the accuracy is observed superior for nodes with a higher degree.

This shows that rich graph structure is helpful for the node classification in general. Indeed, the MLP is a simple model that misses information on the graph structure and thus benefits from graph structure contributed by KeMLP in the form of binary predicates. On the contrary, standalone GNNs can process graph structure by using message passing techniques to transmit learned node representations between neighbors. The prior knowledge introduced in the knowledge enhancer is simple. It encodes that two neighbors are likely to be of the same class. An explanation for the small difference in performance is that GNNs may be able to capture and propagate this simple knowledge across neighbors implicitly, using its message passing technique. In other words we observe that, in this particular case, the introduced knowledge happens to be redundant for GNNs. However, the introduced knowledge significantly improves the accuracy of MLPs. In this context, we discuss perspectives for future work in Section V.

2) *Robustness to wrong knowledge*: Furthermore, a question of interest is how the knowledge enhanced model find a balance between knowledge and graph data in case of knowledge that is not consistent with the graph data. In other words, can the KeGNN successfully deal with nodes having mainly neighbors that belong to a different ground truth class and thus contribute misleading information to the node classification?

To analyze this question, we categorize the accuracy by the proportion of misleading nodes in the neighborhood, see Figure 4. Misleading nodes are node neighbors that have a different ground truth class than the node to be classified. It turns out that KeMLP is particularly helpful over MLP when the neighborhood provides the right information. However, if the neighborhood is misleading (if most or even all of the

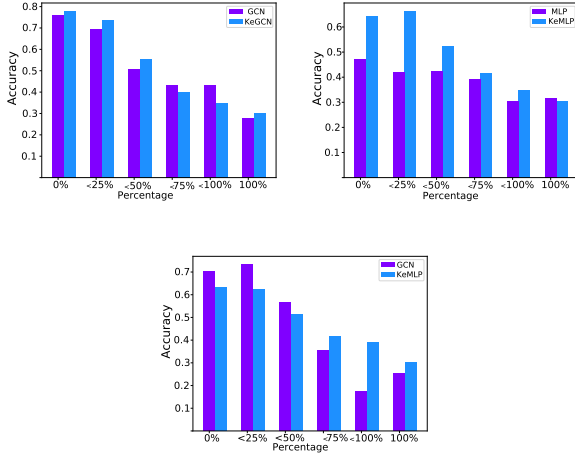


Fig. (4) The accuracy grouped by the ratio of misleading first-order neighbors for GCN vs. KeGCN (left), MLP vs. KeMLP (right), GCN vs. KeMLP (below) on Citeseer.

neighbors belong to a different class), an MLP that ignores the graph structure can lead to even better results. When comparing KeGCN and GCN, there is no clear difference. This is expected, since both models are equally affected by misleading nodes as they utilize the graph structure. Just as a GCN, the KeGCN is not necessarily robust to wrong prior knowledge since the GCN component uses the entire neighborhood, including the misleading nodes.

When comparing GCN to KeMLP, see plot below in Figure 4, KeMLP is more robust to misleading neighbors. While GCN takes the graph structure as given and includes all neighbors equally in the embeddings by graph convolution, the clause weights in the knowledge enhancement layers provide a way to devalue knowledge. If the data frequently contradicts a clause, the model has the capacity to reduce the respective clause weight in the learning process and reduce its impact.

3) *Clause Weight Learning*: Further, we want to examine whether the clause weights learned during training are aligned with the knowledge in the ground truth data. The clause weights provide insights on the magnitude of the updates made by a clause. The *clause compliance* [13] measures how well

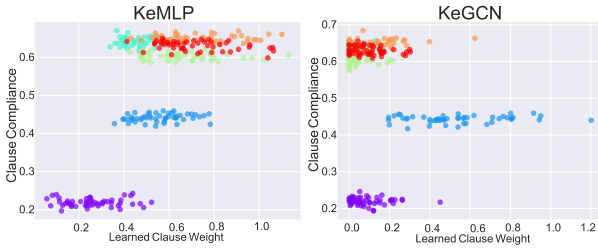


Fig. (5) Learned clause weights vs. clause compliance for KeMLP (left) and KeGCN (right) on Citeseer.

the prior knowledge is satisfied in a graph. Given a clause φ ,

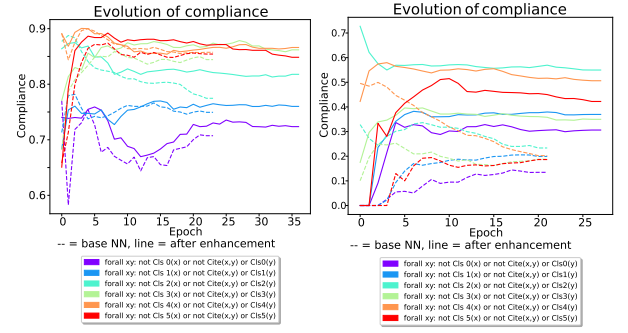


Fig. (6) Clause compliance during training for GCN vs. KeGCN (left) and MLP vs. KeMLP (right) on Citeseer.

a class Cls_i , a set of nodes $\mathbf{V}$, a set of nodes of the class Cls_i : $\mathbf{V}_i = \{v_i | v_i \in \mathbf{V} \wedge \text{Cls}(v_i) == i\}$, and the neighborhood $\mathcal{N}(v_i)$ of v_i , the clause compliance of clause φ on graph $\mathbf{G}$ is defined as follows:

$$\text{Compliance}(\mathbf{G}, \varphi) = \frac{\sum_{v_i \in \mathbf{V}_i} \sum_{v_j \in \mathcal{N}(v)} \mathbf{1}[\text{if } v_j \in \mathbf{V}_i]}{\sum_{v_i \in \mathbf{V}_i} |\mathcal{N}(v_i)|} \quad (1)$$

In other words, the clause compliance counts how often among nodes of a class Cls_i the neighboring nodes have the same class [13]. The clause compliance can be calculated on the ground truth classes of the training set or the predicted classes. As a reference, we measure the clause compliance based on the ground truth labels in the training set. Figure 5 displays the learned clause weights for KeGCN and KeMLP versus the clause compliance on the ground truth labels of the training set. For KeMLP, a positive correlation between the learned clause weights and the clause compliance on the training set is observed. This indicates that higher clause weights are learned for clauses that are satisfied in the training set. Consequently, these clauses have a higher impact on the updates of the predictions. In addition, the clause weights corresponding to clauses with low compliance values make smaller updates to the initial predictions. Accordingly, clauses that are rarely satisfied learn lower clause weights during the training process. In the case of KeGCN, the clause weights are predominantly set to values close to zero. This is in accordance with the absence of a significant performance gap between GCN and KeGCN. Since the GCN itself already leads to valid classifications, smaller updates are required by the clause enhancers.

Furthermore, we analyze how the compliance evolves during training to investigate whether the models learn predictions that increase the satisfaction of the prior knowledge. Figure 6 plots the evolution of the clause compliance for the six clauses for GCN vs. KeGCN and MLP vs. KeMLP. It is observed that GCN and KeGCN yield similar results as the evolution of the compliance during training for both models is mostly aligned. For MLP vs. KeMLP the clause compliance of the prediction of the MLP converges to lower values for all classes than the clause compliance obtained with the KeMLP. This

gives evidence that the knowledge enhancement layer actually improves the satisfiability of the prior knowledge. As already observed, this gives evidence that the standalone GCN is able to implicitly satisfy the prior knowledge even though it is not explicitly defined.

B. Additional Experiment Details

1) *Implementation:* The code⁴ is based on PyTorch [39] and the graph learning library PyTorch Geometric [37]. The Weights & Biases tracking tool [40] is used to monitor the experiments. All experiments are conducted on a machine running an Ubuntu 20.4 equipped with an Intel(R) Xeon(R) Silver 4114 CPU 2.20GHz processor, 192G of RAM and one GPU Nvidia Quadro P5000.

2) *Model Parameters and Hyperparameter Tuning:* KeGNN contains a set of hyperparameters. Batch normalization [41] is applied after each hidden layer of the GNN. The Adam optimizer [42] is used as optimizer for all models. Concerning the hyperparameters specific to the knowledge enhancement layers, the initialization of the preactivations of the binary predicates (which are assumed to be known) is taken as a hyperparameter. They are set to a high positive value for edges that are known to exist and correspond to the grounding of the binary predicate. Furthermore, different initializations of clause weights and constraints on them are tested. Moreover, the number of stacked knowledge enhancement layers is a hyperparameter. We further allow the model to randomly neglect a proportion of edges by setting an edges drop rate parameter. Further, we test whether the normalization of the edges with the diagonal matrix $\tilde{\mathbf{D}} = \sum_j \tilde{\mathbf{A}}_{i,j}$ (with $\tilde{\mathbf{A}} = \mathbf{A} + \mathbf{I}$) is helpful.

To find a suitable set hyperparameters for each dataset and model, we perform a random search with up to 800 runs and 48h time limit and choose the parameter combination which leads to the highest accuracy on the validation set. The hyperparameter tuning is executed in Weights and Biases [40]. The following hyperparameter values are tested:

- Adam optimizer parameters: β_1 : 0.9, β_2 : 0.99, ϵ : 1e-07
- Attention heads: {1, 2, 3, 4, 6, 8, 10}
- Batch size: {128, 512, 1024, 2048, full batch}
- Binary preactivation: {0.5, 1.0, 10.0, 100.0, 500.0}
- Clause weights initialization: {0.001, 0.1, 0.25, 0.5, random uniform distribution on [0,1]}
- Dropout rate: 0.5
- Edges drop rate: random uniform distribution [0.0, 0.9]
- Edge normalization: {true, false}
- Early stopping: δ_{min} : 0.001, patience: {1, 10, 100}
- Hidden layer dimension: {32, 64, 128, 256}
- Learning rate: random uniform distribution [0.0001, 0.1]
- Clause weight clipping: w_{min} : 0.0, w_{max} : random uniform distribution: [0.8, 500.0]
- Number of knowledge enhancement layers: {1, 2, 3, 4, 5, 6}
- Number of hidden layers: {2, 3, 4, 5, 6}
- Number of epochs 200 (unless training stopped early)

The obtained parameter combinations for the models KeMLP, KeGCN and KeGAT for Cora, Citeseer, PubMed and Flickr

are displayed in Table III. We set the random seed for all experiments to 1234.

The reference models MLP, GCN and GAT are trained with the same parameter set as the respective knowledge enhanced models.

V. LIMITATIONS AND PERSPECTIVES

The method of KeGNN is limited in some aspects, which we present in this section. In this work, we focus on homogeneous graphs. In reality, however, graphs are often heterogeneous with multiple node and edge types [4]. Adaptations are necessary on both the neural and the symbolic side to apply KeGNN to heterogeneous graphs. The restriction to homogeneous graphs also limits the scope of formulating complex prior knowledge. Eventually, the datasets used in this work and the set of prior knowledge are too simple for KeGNN to exploit its potential and lead to a significant improvement over the GNN. The experimental results show that the knowledge encoded in the symbolic component leads to significant improvement over an MLP that is not capable to capture and learn that knowledge. This indicates that for more complex knowledge that is harder for a GNN to learn, KeGNN has the potential to bring higher improvements. A perspective for further work is the extension of KeGNN to more generic data structures such as incomplete and heterogeneous knowledge graphs in conjunction with more complex prior knowledge.

Another limitation of KeGNN is scalability. With an increasing number of stacked knowledge enhancement layers, the affected node neighborhood grows exponentially, which can lead to significant memory overhead. This problem is referred as neighborhood explosion [7] and is particularly problematic in the context of training on memory-constrained GPUs. This affects both the GNN and the knowledge enhancement layers that encode binary knowledge. Methods from scalable graph learning [20], [43], [44] represent potential solutions for the neighborhood explosion problem in KeGNN.

Furthermore, limitations appear in the context of link prediction with KeGNN. For link prediction, a neural component is required that predicts fuzzy truth values for binary predicates. At present, KeGNN can handle clauses containing binary predicates, but their truth values are initialized with artificial predictions, where a high value encodes the presence of an edge. This limits the application of KeGNN to datasets for which the graph structure is complete and known a priori.

VI. CONCLUSION

In this work, we introduced KeGNN, a neuro-symbolic model that integrates GNNs with symbolic knowledge enhancement layers to create an end-to-end differentiable model. This allows the use of prior knowledge to improve node classification while exploiting the strength of a GNN to learn expressive representations. Experimental studies show that the inclusion of prior knowledge has the potential to improve simple neural models (as observed in the case of MLP). However, the knowledge enhancement of GNNs is harder to achieve on the underlying and limited benchmarks for

	mode	atten- tion heads	batch size	binary preacti- vation	initial clause weight	edges drop rate	es pati- ence	hidden chan- nels	learn- ing rate	norma- lize edges	KE layers	hidden layers
PubMed	KeMLP	-	1024	10.0	0.001	0.22	100	256	0.057	false	2	4
	KeGCN	-	full batch	1.0	random	0.66	10	256	0.043	false	1	2
	KeGAT	8	1024	10.0	0.5	0.07	10	256	0.016	true	5	2
Flickr	KeMLP	-	128	10.0	0.001	0.2	10	32	0.001	true	1	2
	KeGCN	-	1024	500.0	0.001	0.24	10	128	0.016	true	4	4
	KeGAT	8	2048	500.0	0.1	0.12	100	64	0.0039	false	1	3
Cora	KeMLP	-	512	10.0	0.5	0.47	1	32	0.026	true	4	2
	KeGCN	-	512	100.0	random	0.17	1	256	0.032	false	2	2
	KeGAT	1	full batch	1.0	0.5	0.27	10	64	0.033	true	1	2
Citeseer	KeMLP	-	128	10.0	0.5	0.01	10	256	0.028	true	1	2
	KeGCN	-	full batch	0.5	0.25	0.35	10	128	0.037	false	3	5
	KeGAT	3	1024	0.5	0.1	0.88	10	32	0.006	true	2	2

TABLE (III) Hyperparameters and experiment configuration for PubMed and Flickr

which the injection of simple knowledge concerning local neighborhood is redundant with the representations that GNNs are able to learn. Nevertheless, KeGNN has not only the potential to improve graph completion tasks from a performance perspective, but also to increase interpretability through clause weights. This work is a step towards a holistic neuro-symbolic method on incomplete and noisy semantic data, such as knowledge graphs.

ACKNOWLEDGMENTS

This work has been partially supported by the MIAI Knowledge communication and evolution chair (ANR-19-P3IA-0003).

REFERENCES

- [1] W. Liu, Y. Zhang, J. Wang, Y. He, J. Caverlee, P. Chan, D. Yeung, and P.-A. Heng, "Item relationship graph neural networks for e-commerce," vol. PP, 03 2021, pp. 1–15.
- [2] A. Sanchez-Gonzalez, N. Heess, J. T. Springenberg, J. Merel, M. Riedmiller, R. Hadsell, and P. Battaglia, "Graph networks as learnable physics engines for inference and control," in *Proceedings of the 35th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, J. Dy and A. Krause, Eds., vol. 80. PMLR, 10–15 Jul 2018, pp. 4470–4479. [Online]. Available: <https://proceedings.mlr.press/v80/sanchez-gonzalez18a.html>
- [3] Y. Wu, D. Lian, Y. Xu, L. Wu, and E. Chen, "Graph convolutional networks with markov random field reasoning for social spammer detection," vol. 34, no. 01, Apr. 2020, pp. 1054–1061. [Online]. Available: <https://ojs.aaai.org/index.php/AAAI/article/view/5455>
- [4] X. Yang, M. Yan, S. Pan, X. Ye, and D. Fan, "Simple and efficient heterogeneous graph neural network," 2022. [Online]. Available: <https://arxiv.org/abs/2207.02547>
- [5] Y. Ma and J. Tang, *Deep Learning on Graphs*. Cambridge University Press, 2021.
- [6] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, and P. S. Yu, "A comprehensive survey on graph neural networks," vol. 32, no. 1. Institute of Electrical and Electronics Engineers (IEEE), jan 2021, pp. 4–24. [Online]. Available: <https://doi.org/10.1109%2Ftnnls.2020.2978386>
- [7] K. Duan, Z. Liu, P. Wang, W. Zheng, K. Zhou, T. Chen, X. Hu, and Z. Wang, "A comprehensive study on large-scale graph training: Benchmarking and rethinking," in *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022.
- [8] L. Wu, P. Cui, J. Pei, and L. Zhao, *Graph Neural Networks: Foundations, Frontiers, and Applications*. Singapore: Springer Singapore, 2022.
- [9] Z. Susskind, B. Arden, L. K. John, P. Stockton, and E. B. John, "Neuro-symbolic AI: an emerging class of AI workloads and their characterization," vol. abs/2109.06133, 2021.
- [10] L. De Raedt, S. Dumančić, R. Manhaeve, and G. Marra, "From statistical relational to neural-symbolic artificial intelligence," in *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence*, ser. IJCAI'20, 2021.
- [11] M. Garlato and M. Shanahan, "Reconciling deep learning with symbolic artificial intelligence: representing objects and relations," vol. 29, 2019, pp. 17–23, artificial Intelligence. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2352154618301943>
- [12] G. Marra, M. Diligenti, F. Giannini, M. Gori, and M. Maggini, "Relational neural machines," in *ECAI 2020 - 24th European Conference on Artificial Intelligence, 29 August-8 September 2020, Santiago de Compostela, Spain, August 29 - September 8, 2020 - Including 10th Conference on Prestigious Applications of Artificial Intelligence (PAIS 2020)*, ser. Frontiers in Artificial Intelligence and Applications, G. D. Giacomo, A. Catalá, B. Dilkina, M. Milano, S. Barro, A. Bugarín, and J. Lang, Eds., vol. 325. IOS Press, 2020, pp. 1340–1347. [Online]. Available: <https://doi.org/10.3233/FAIA200237>
- [13] A. Daniele and L. Serafini, "Neural networks enhancement with logical knowledge," <https://arxiv.org/abs/2009.06087>, 2020, unpublished.
- [14] —, "Knowledge enhanced neural networks for relational domains," in *AIXIA 2022 – Advances in Artificial Intelligence*, A. Dovier, A. Montanari, and A. Orlandini, Eds. Cham: Springer International Publishing, 2023, pp. 91–109.
- [15] E. Grilli, A. Daniele, M. Bassier, F. Remondino, and L. Serafini, "Knowledge enhanced neural networks for point cloud semantic segmentation," vol. 15, no. 10, 2023. [Online]. Available: <https://www.mdpi.com/2072-4292/15/10/2590>
- [16] A. Daniele and L. Serafini, "Knowledge enhanced neural networks," in *PRICAI 2019: Trends in Artificial Intelligence*, A. C. Nayak and A. Sharma, Eds. Cham: Springer International Publishing, 2019, pp. 542–554.
- [17] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, and Y. Bengio, "Graph attention networks," in *International Conference on Learning Representations*, 2018.
- [18] T. N. Kipf and M. Welling, "Semi-Supervised Classification with Graph Convolutional Networks," in *Proceedings of the 5th International Conference on Learning Representations*, ser. ICLR '17, 2017. [Online]. Available: <https://openreview.net/forum?id=SJU4ayYgl>
- [19] Z. Yang, W. W. Cohen, and R. Salakhutdinov, "Revisiting semi-supervised learning with graph embeddings," in *Proceedings of the 33rd International Conference on International Conference on Machine Learning - Volume 48*, ser. ICML'16. JMLR.org, 2016, p. 40–48.
- [20] H. Zeng, H. Zhou, A. Srivastava, R. Kannan, and V. Prasanna, "Graph-saint: Graph sampling based inductive learning method," in *International Conference on Learning Representations*, 2020.

- [21] H. Caselles-Dupré, F. Lesaint, and J. Royo-Letelier, “Word2vec applied to recommendation: Hyperparameters matter,” in *Proceedings of the 12th ACM Conference on Recommender Systems*, ser. RecSys ’18. New York, NY, USA: Association for Computing Machinery, 2018, p. 352–356. [Online]. Available: <https://doi.org/10.1145/3240323.3240377>
- [22] L. Zadeh, “Fuzzy logic,” *Computer*, vol. 21, no. 4, pp. 83–93, 1988.
- [23] E. Klement, R. Mesiar, and E. Pap, *Triangular Norms*, ser. Trends in Logic. Springer Netherlands, 2013. [Online]. Available: <https://books.google.fr/books?id=HXzvCAAAQBAJ>
- [24] D. Dou, H. Wang, and H. Liu, “Semantic data mining: A survey of ontology-based approaches,” in *Proceedings of the 2015 IEEE 9th International Conference on Semantic Computing (IEEE ICSC 2015)*, 2015, pp. 244–251.
- [25] C. Meilicke, M. W. Chekol, D. Ruffinelli, and H. Stuckenschmidt, “Anytime bottom-up rule learning for knowledge graph completion,” in *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, ser. IJCAI’19. AAAI Press, 2019, p. 3137–3143.
- [26] Y. Dai, S. Wang, N. N. Xiong, and W. Guo, “A survey on knowledge graph embedding: Approaches, applications and benchmarks,” *Electronics*, vol. 9, no. 5, 2020. [Online]. Available: <https://www.mdpi.com/2079-9292/9/5/750>
- [27] R. Abboud and I. I. Ceylan, “Node classification meets link prediction on knowledge graphs,” 2021. [Online]. Available: <https://arxiv.org/abs/2106.07297>
- [28] S. Badreddine, A. d’Avila Garcez, L. Serafini, and M. Spranger, “Logic tensor networks,” vol. 303. Elsevier BV, feb 2022, p. 103649. [Online]. Available: <https://doi.org/10.1016%2Fj.artint.2021.103649>
- [29] L. N. DeLong, R. F. Mir, M. Whyte, Z. Ji, and J. D. Fleuriot, “Neurosymbolic ai for reasoning on graph structures: A survey,” 2023. [Online]. Available: <https://arxiv.org/abs/2302.07200>
- [30] W. Li, R. Peng, and Z. Li, “Knowledge graph completion by jointly learning structural features and soft logical rules,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 3, pp. 2724–2735, 2023.
- [31] N. Jain, T.-K. Tran, M. H. Gad-Elrab, and D. Stepanova, “Improving knowledge graph embeddings with ontological reasoning,” in *The Semantic Web – ISWC 2021: 20th International Semantic Web Conference, ISWC 2021, Virtual Event, October 24–28, 2021, Proceedings*. Berlin, Heidelberg: Springer-Verlag, 2021, p. 410–426. [Online]. Available: https://doi.org/10.1007/978-3-030-88361-4_24
- [32] B. Fatemi, S. Ravanbakhsh, and D. Poole, “Improved knowledge graph embedding using background taxonomic information,” vol. 33, 07 2019, pp. 3526–3533.
- [33] S. Guo, Q. Wang, L. Wang, B. Wang, and L. Guo, “Jointly embedding knowledge graphs and logical rules,” in *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*. Austin, Texas: Association for Computational Linguistics, Nov. 2016, pp. 192–202. [Online]. Available: <https://aclanthology.org/D16-1019>
- [34] Y. Hu, Z. Ye, M. Wang, J. Yu, D. Zheng, M. Li, Z. Zhang, Z. Zhang, and Y. Wang, “Featgraph: A flexible and efficient backend for graph neural network systems,” in *SC20: International Conference for High Performance Computing, Networking, Storage and Analysis*, 2020, pp. 1–13.
- [35] M. Qu, J. Chen, L.-P. Xhonneux, Y. Bengio, and J. Tang, “{RNNL}ogic: Learning logic rules for reasoning on knowledge graphs,” in *International Conference on Learning Representations*, 2021.
- [36] H. A. Kautz, “The third ai summer: Aaai robert s. engelmore memorial lecture,” <https://onlinelibrary.wiley.com/doi/10.1002/aaai.12036>, 2022.
- [37] M. Fey and J. E. Lenssen, “Fast graph representation learning with PyTorch Geometric,” 2019. [Online]. Available: <https://arxiv.org/abs/1903.02428>
- [38] J. Chen, T. Ma, and C. Xiao, “Fastgcn: Fast learning with graph convolutional networks via importance sampling,” in *ICLR (Poster)*. OpenReview.net, 2018. [Online]. Available: <http://dblp.uni-trier.de/db/conf/iclr/iclr2018.html#ChenMX18>
- [39] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, A. Desmaison, A. Köpf, E. Yang, Z. DeVito, M. Raison, A. Tejani, S. Chilamkurthy, B. Steiner, L. Fang, J. Bai, and S. Chintala, “Pytorch: An imperative style, high-performance deep learning library,” 2019. [Online]. Available: <https://arxiv.org/abs/1912.01703>
- [40] L. Biewald, “Experiment tracking with weights and biases,” 2020, software available from wandb.com. [Online]. Available: <https://www.wandb.com/>
- [41] S. Ioffe and C. Szegedy, “Batch normalization: Accelerating deep network training by reducing internal covariate shift,” in *Proceedings of the 32nd International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, F. Bach and D. Blei, Eds., vol. 37. Lille, France: PMLR, 07–09 Jul 2015, pp. 448–456. [Online]. Available: <https://proceedings.mlr.press/v37/ioffe15.html>
- [42] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” in *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7–9, 2015, Conference Track Proceedings*, Y. Bengio and Y. LeCun, Eds., 2015. [Online]. Available: <http://arxiv.org/abs/1412.6980>
- [43] M. Fey, J. E. Lenssen, F. Weichert, and J. Leskovec, “Gnnautoscale: Scalable and expressive graph neural networks via historical embeddings,” in *Proceedings of the 38th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, M. Meila and T. Zhang, Eds., vol. 139. PMLR, 18–24 Jul 2021, pp. 3294–3304. [Online]. Available: <https://proceedings.mlr.press/v139/fey21a.html>
- [44] W. L. Hamilton, R. Ying, and J. Leskovec, “Inductive representation learning on large graphs,” in *Proceedings of the 31st International Conference on Neural Information Processing Systems*, ser. NIPS’17. Red Hook, NY, USA: Curran Associates Inc., 2017, p. 1025–1035.