

Analysis of the Performance of Representation Learning Methods for Entity Alignment: Benchmark vs. Real-world Data

Ensiyeh Raoufi ^{a,*}, Bill Gates Happi Happi ^b, Pierre Larmande ^b, François Scharffe ^a and Konstantin Todorov ^a

^a *LIRMM, University of Montpellier, CNRS, France*

E-mails: ensiyeh.raoufi@umontpellier.fr, francois.scharffe@lirmm.fr, konstantin.todorov@lirmm.fr

^b *DIADE, IRD, CIRAD, University of Montpellier, France*

E-mails: bill.happi@ird.fr, pierre.larmande@ird.fr

Abstract.

Representation learning for Entity Alignment (EA) aims to map, across two Knowledge Graphs (KG), distinct entities that correspond to the same real-world object using an embedding space. Hence, the similarity of the learned entity embeddings serves as a proxy for that of the actual entities. Although many embedding-based models show very good performance on established synthetic benchmark datasets, in this paper we demonstrate that benchmark overfitting limits the applicability of these methods in real-world scenarios, where we deal with highly heterogeneous, incomplete, and domain-specific data. While there have been efforts to employ sampling algorithms to generate benchmark datasets reflecting as much as possible real-world scenarios, there is still a lack of comprehensive analysis and comparison between the performance of methods on synthetic benchmark and original real-world heterogeneous datasets. In addition, most existing models report their performance by excluding from the alignment candidate search space entities that are not part of the validation data. This under-represents the knowledge and the data contained in the KGs, limiting the ability of these models to find new alignments in large-scale KGs. We analyze models with competitive performance on widely used synthetic benchmark datasets, such as the cross-lingual DBP15K. We compare the performance of the selected models on real-world heterogeneous datasets beyond DBP15K and we show that most of the current approaches are not effectively capable of discovering mappings between entities in the real-world, due to the above-mentioned drawbacks. We compare the utilized methods from different aspects and measure joint semantic similarity and profiling properties of the KGs to explain the models' performance drop on real-world datasets. Furthermore, we show how tuning the EA models by restricting the search space only to validation data affects the models' performance and causes them to face generalization issues. By addressing practical challenges in applying EA models to heterogeneous datasets and providing valuable insights for future research, we signal the need for more robust solutions in real-world applications.

Keywords: Entity alignment, Knowledge graphs, Representation learning, Knowledge graph heterogeneity, EA Benchmarks

1. Introduction

Knowledge Graphs (KG) have emerged as powerful tools for a range of applications, including information retrieval, question answering, and data federation [1]. An entity in a knowledge graph refers to a distinct and identifiable concept, which can be, for example, a concrete object, an abstract idea, or an event. Entities are represented

*Corresponding author. E-mail: ensiyeh.raoufi@umontpellier.fr.

as nodes, forming the building blocks of the graph. The relationships between entities are modeled as edges, serving as connections that establish associations or interactions between the nodes. These edges convey various meanings, representing **attribute properties** or **relation properties** of the entities. Examples of attribute properties include "birth date", "genre" or "description" and examples of relation properties include "located in", "established by" or "worked with". Based on the given definition, (Person1, birth date, "1989-09-30") indicates an attribute triple, while (City1, located in, Country1) indicates a relation triple.¹ Essentially, the combination of entities and their interconnecting relationships forms a structured representation of knowledge. Hence, knowledge graphs are designed in a way that facilitates storage, access, semantic understanding of data and reasoning over it and are widely used in a variety of domains, including the Semantic Web in general [2–4], cultural heritage [5–8], biomedicine [9–12], sociology [13–15], and data-driven industries [16, 17].

As data comes from different sources, it is often scattered across multiple knowledge graphs (KGs), even if it conveys the same information, leading to various challenges. One such challenge is identifying and matching entities from a source KG to their equivalents in a target KG that represent the same real-world object [18]—a task known as **Entity Alignment (EA)**. EA in turn facilitates data integration, information retrieval, and entity disambiguation across diverse knowledge sources [19–22].

We start by providing some *terminology*. We use the term **KG heterogeneity** in the sense of [22], common in the fields of ontology and entity alignment. In a nutshell, it concerns any difference in the expression of a given piece of knowledge across two KGs (be it structural, syntactical, terminological, or other). Sticking to the data heterogeneity concept provided in [22] and algebraic properties of the graphs, in this paper, we are interested in differences in value and structural levels of the KGs in each dataset. The pair of KGs in the datasets might differ in the graph's structural properties such as size, degree distribution, etc. Also, each pair of aligned entities in two KGs may have descriptive and data quality heterogeneities, following [22].

By **dataset**, we mean an EA dataset, which consists of a pair of KGs: a source and a target KG to be interlinked, together with a reference alignment that helps evaluate or train the models. A **reference (or seed) alignment** is a manually curated set of correspondences or alignments (often together with a confidence score) between entities across the two different KGs. By **unmatchable entities**, we mean pairs of entities from the source and target KGs that are not to be aligned (i.e. they refer to different real-world entities).

We distinguish two main types of datasets. A **synthetic benchmark dataset** refers to a dataset that consists mostly of KGs sampled from larger ones following a motivation of having smaller KGs that mimic certain characteristics of the real large KGs. Additionally, in machine learning applications, benchmark datasets could potentially be fully synthetic, i.e. entirely generated from scratch, not using a real KG as a start point. Often in benchmark datasets the source entities are matched with their corresponding counterparts in the target KG under the **1-to-1 assumption** (meaning that each source entity has exactly one match in the target graph).² In this paper, we use the terms "benchmark dataset" and "synthetic benchmark dataset" interchangeably. A **real-world dataset**, on the other hand, is one that is issued from a real-world scenario and contains the unchanged KGs that are not sub-sampled from larger KGs under some conditions like being sparse or dense, or retaining a similar degree distribution as the KGs they are sampled from.

For the purposes of this study, we categorize entity alignment techniques into two main groups: **embedding** and **non-embedding**-based methods. Non-embedding-based approaches apply user-crafted representations of entities and relations and align the entities across the KGs based on similarity measures or logic axioms. This group of approaches prioritizes symbolic reasoning, logical inferences and linking specifications defined by domain experts to guide the alignment process [23]. Embedding-based approaches, in contrast, automatically represent entities in a feature space and predict alignments based on similarity metrics over the learned embeddings. **Embedding** refers to representing an object as a vector in a continuous space based on a given number of constraints (e.g. close in meaning entities should have vectors that live close to one another in the embedding space) [23]. Following this paradigm, embedding-based EA models commonly use an embedding and an alignment module. While the embedding module

¹While this work was conducted with knowledge graphs represented in the RDF format, the research is not restricted to this paradigm and applies to property graphs as well.

²Note that the 1:1 assumption is sometimes "imposed" by the models, while in many cases it is inherent to the benchmark datasets, as is the case in the OpenEA datasets, discussed in detail below.

represents each KG entity as a vector in a low-dimensional embedding space, the alignment module ensures that aligned entities are close together in a unique embedding space or learns a mapping between KGs with respect to the reference alignment. During the model's training phase, through iterative interactions between the embedding and alignment modules, all entities from both KGs are embedded, and predictions are made regarding which entities are most likely to align.

Relying on non-embedding-based methods may be more suitable in scenarios dealing with sparse or small datasets, or in situations where there is not a high variety of heterogeneities such as predicate, class, or graph linking problems (as defined by [24]). Real-world datasets often do not meet these ideal conditions and therefore, embedding-based EA methods promise more efficiency, as they are flexible in cross-lingual scenarios, scalable to large KGs, and globally consistent in representations across KGs.

This paper's focus is on analyzing specifically the embedding-based EA models, with respect precisely to both synthetic benchmark datasets and real-world datasets, as well as considering the different training and evaluation strategies and model types. Hence, we add to several recent surveys and studies that investigate that question from a critical viewpoint [23, 25, 26]. For example, [27–29] analyze the performance of embeddings-based entity alignment models and compare them regarding their performance on benchmark datasets [23, 30–35]. Previous research has extracted more realistic EA benchmark datasets [30] from large knowledge bases like DBpedia [36]. We enhance the work of the studies cited above by expanding the benchmarks with datasets containing a low percentage of matchable entities to better reflect real-world scenarios, and evaluating the performance of two of the leading embedding models (RDGCN and BERT-INT) when we include all entities in the target KG as alignment candidates during the model evaluation. Furthermore, we include an in-depth discussion on the evaluation metrics that are commonly used for the EA task, building on preliminary remarks found in [37]. As an overarching question, we consider several recent EA embedding-based models having state-of-the-art performance on synthetic benchmark datasets and analyze their capacities when they deal with heterogeneous real-world data. We show a considerable drop in performance in the latter scenarios. To help understand this observation, on the one hand we analyze and compare the real-world and synthetic benchmark datasets with respect to a set of dataset profiling features [38] studied and applied for the entity alignment task. On the other hand, we pair these observations with a look into the underlying nature of the embedding-based models. [39] observes that cutting-edge entity alignment models did not address the particular properties of data well because they prioritized genericity and automation. Indeed, the results of our study demonstrate that while embedding-based models perform well on certain synthetic benchmark datasets, they struggle in real-world scenarios due to insufficient consideration of the inherent characteristics and nature of the data. Finally in order to be able to compare the embedding-based models to methods from the non-embeddings group, we include in our analyses the DLinker system [40], for reasons explained in Section 3.

The main contributions of this analysis paper are:

- A novel look into and comparison of the frameworks of established EA methods having different embedding bases: we propose a novel categorization of embedding-based EA methods based on their embedding approaches.
- A comparison of the features of synthetic benchmark and real-world datasets from aspects related to entity alignment: although it appears difficult to isolate a structure-related meta-feature which explains the performance of all methods on the different datasets (because each method embeds the structure from a different aspect), we find that the semantic similarity is the dataset meta-feature that correlates at the strongest with the performance of embedding-based EA methods.
- A discussion of the commonly used evaluation metrics for the EA task: we explain how Hit@1 is equivalent to precision and recall when under the 1-to-1 assumption in the validation set and when and why each evaluation metrics should be applied.
- An analysis of the performance drop of EA methods on real-world datasets in comparison to their performance on established synthetic benchmarks: we present shreds of evidence and probable reasons to explain the observed drop in performance; we go beyond the 1-to-1 assumption during the model evaluation and investigate the performance drop of the EA models using both Hit@1 and F_1 -score.
- Analysis of the different categories of embeddings models with respect to synthetic and real-world datasets: we find out interaction training models to be the best-performing category of EA methods on real-world large-scale data.

The paper is organized as follows: Section 2 summarizes related surveys and empirical studies on EA methods, positioning our work within that context. Section 3 outlines the key features of the embedding-based EA methods selected for performance analysis. In Section 4, we compare several benchmark and real-world datasets, highlighting the greater heterogeneity and complexity of the latter. Finally, Section 5 examines the performance of EA models on heterogeneous real-world data and explores the reasons for their performance decline on these datasets.

2. Related Work

This section provides an overview of surveys and related analytical studies that critically examine existing EA approaches, with a particular focus on embeddings-based methods. In line with the scope of the paper, specific alignment methods are not discussed here.

Several studies have contributed to the understanding and advancement of KG embeddings and their applications such as link prediction, knowledge graph completion and reasoning, and entity alignment [1, 41–44]. While [42] reveals sharp differences in the geometry of embeddings produced by various KG embedding methods, [45] introduces a multi-embedding interaction mechanism for analyzing KG embedding models like DistMult [46] and ComplEx [47]. The latter study unifies and generalizes these models, offering an intuitive perspective for their effective use. The authors of [48] introduce a scalable and open-source Python library for multi-source knowledge graph embeddings. Supporting joint representation learning, it implements 26 KG embedding models and 16 benchmark datasets. Moreover, [49] categorizes the existing KG Embedding (KGE) models based on representation spaces and discusses whether they have algebraic, geometric, or analytical structures.

Several surveys and experimental studies have been conducted on methods for entity alignment across knowledge graphs [32, 50, 51]. Broadly, these studies categorize EA techniques into two main groups: embedding-based methods and traditional approaches [23, 30, 31]. Traditional EA methods rely on user-defined rules, OWL reasoning, and/or similarity computations based on symbolic features of entities. We refer to these as non-embedding-based methods.

Moving now the focus towards embedding-based methods, Sun et al. [30] created an open-source toolkit, named OpenEA. The authors discuss the characteristics and functionalities of embedding-based methods, highlighting how they predict matching entities through nearest-neighbor searches among target entity embeddings. Two combination paradigms are outlined: one encoding KGs in independent spaces and learning a mapping using seed alignment, and another representing KGs in a unified space considering highly similar embeddings for aligned entities. The study underscores the incorporation of entity relation and attribute properties into embedding modules to enhance accuracy, categorizing relation embeddings into triple-based, path-based, and neighborhood-based groups. Attribute embedding, achieved through correlation or literal methods, is also explored for improving entity similarity assessment. In [32], Fanourakis et al. present the meta-features of the OpenEA datasets, which adhere to the 1-to-1 assumption, and explain the technical details of several embedding-based EA models. However, they do not include details regarding the generation of the dataset’s meta-features, such as description similarity. In this work, we provide formulas to compute the extracted meta-features, we analyze the performance of EA models on both benchmark and real-world datasets that do not follow the 1-to-1 assumption.

In [31], Zhang et al. analyze the performance of Translational and GNN-based EA methods with respect to the seed alignment and dataset sizes, the use (or not) of attribute triples, the presence of multilingual data, and the embedding size. They propose new benchmark datasets sampled from large-scale knowledge graphs like Wikidata [52] and Freebase [53] that do not fulfill the 1-to-1 assumption (40% and 75% of the entities in every pair of KGs combined in the datasets do not have matches). The authors then tested several EA models on the newly sampled dataset. However, one of the issues with this approach is the fact that the 1-to-1 assumption is not a condition that only holds for the datasets but, also many EA models consider that constraint during the model evaluation. Hence, even though Zhang et al. generate new data including not only 1:1 mappings, none of the non-matchable entities are considered during the evaluation for some of the models they applied (such as MultiKE [46] and TransEdge [54]) due to the nature of these models that only consider the alignment candidates from the reference alignment. Similarly, Jiang et al. [35] evaluate the performance of EA methods on newly generated, highly heterogeneous KGs that differ in scale and structure, with fewer overlapping entities than benchmark datasets (sharing 60% and

25% of the entities). They introduced two more realistic datasets, ICEWS-WIKI and ICEWS-YAGO, which were derived from knowledge bases with significantly different degree distributions. Although these new datasets deviate from the typical 1-to-1 assumption, they tested EA methods like GCN-Align [55], RDGCN, and FuAlign [56] on subsampled heterogeneous datasets, overlooking the fact that these methods do not account for non-matchable entities as candidates in the evaluation phase. In contrast to these studies and building on these observations, we discuss in detail this search-space-related issue in our experiments in Section 5.3.2.

In [23], Zeng et al. provide a brief overview of research in entity alignment, covering traditional methods, knowledge representation learning, and alignment based on representation learning in knowledge graphs. They conducted their research only on one single dataset – DBP15K, which is a synthetic benchmark datasets holding the 1:1 assumption. They tested only two categories of models: Translational and GNN-based in the context of multilingualism. In contrast, our work covers a wider variety of both datasets and models that differ in nature.

In [57], Fanourakis et al. explores indirect biases of EA methods due to structural diversity in the KGs and introduces a sampling algorithm to generate challenging benchmark datasets by changing the properties of the KGs. In that way, the authors assess EA methods robustness against such diversity. Modifications include changing connectivity metrics such as "average node degree", "max component size" (maxCS), and "ratio of weakly connected components" (wccR) to control the level of structural heterogeneity of the generated datasets. In our work, we do not use a sampling algorithm, instead, we experiment with EA methods having different design bases on widely used benchmark datasets and real-world datasets.

In [37], Leone et al. provide a discussion on the evaluation metrics for EA for datasets that do not follow the 1-to-1 assumption. To go beyond this assumption, the authors generate sub-sampled datasets whose KG sizes are different, where each dataset variant includes about 30% unmatchable entities. In comparison to Leone's study, in addition to the fact that our real-world datasets are original and not obtained by sampling, the proportion of unmatchable entities for each of our real-world datasets is more than 80% (KGs in Doremus and AgroLD on average have more than 87% and 82% unmatchable entities, respectively). Furthermore, we report the Hit@k measures for the case that a 1-to-1 assumption does not hold on our datasets to compare the models performance with the one reported in previous studies.

To sum up, our work builds on previous research by extending and refining key aspects of EA evaluation. In particular, we study the performance of embedding-based EA methods with distinct representation learning principles on real-world and benchmark datasets. In that, our study stands out for its attention to data quality considerations [38]. While prior studies provide valuable insights into meta-features and sampling methods, we advance this by offering explicit formulas for meta-feature extraction and testing EA models on both benchmark and real-world datasets that do not follow the 1-to-1 assumption. Instead of generating sampled datasets, we focus on real-world original data with over 80% unmatchable entities, providing a more rigorous evaluation. In this way, our work continues and enhances existing research, bringing new perspectives to real-world EA challenges.

3. Methods for Entity Alignment via Representation Learning

Certain studies categorize embedding-based methods according to their use of semantic information to represent the KGs [56], while others categorize them according to whether they use attribute or relation predicates for embedding learning, their alignment modules (i.e., whether they embed both KGs in the same space or separately), or their learning strategy (supervised, semi-supervised, or unsupervised) [30, 32]. Based on recent studies [23, 30–32, 35, 56] and our analysis, we propose to classify the embedding-based EA models into four groups: (1) Translational, (2) GNN-based, (3) Graph Transformers-based, and (4) Interaction training models.

Several entity alignment models such as MTransE [58] and IPTransE [59] have been designed by using **translational techniques** like TransE [60] for KG embedding and entity alignment across KGs [61]. A knowledge graph is usually represented as a directed graph, in which nodes refer to entities and edges refer to relations between entities, or simply by a set of triples of the type $\langle head\ entity, relation, tail\ entity \rangle$. The translational model's framework embeds a relation predicate as a translation vector from a head to a tail entity. **GNN-based methods**, such as GCN-Align [62], RREA [63], and GMNN [64] employ Graph Neural Networks (GNNs) [65–67] to represent the graphs and link them. These models rely on the signature GNN message-passing system to integrate the information of each

entity's neighbors. Inspired by the successful application of the Transformer [68] model in representing sequences for the automatic translation task [69], several works developed and applied the **Graph Transformers** (GT) model for representing graphs and their entities [70–74]. Recently, models built on top of transformers have been developed specifically for the entity alignment task [75, 76], adopting the self-attention mechanism from Transformers to represent entities. As a point of comparison between GNNs and Transformers, we underline that Transformers use multi-head attention, treating the entire sequence as a local neighborhood, whereas standard GNNs aggregate features from local nodes [77]. Applying Transformer architecture to GNNs, as in the Graph Transformers approaches, is motivated by the need to overcome the issue of information dispersion between distant elements in structural data [78]. Graph Transformers address the limitations of traditional GNNs by leveraging Transformers' ability to capture long-term dependencies. By integrating GNNs and Transformers, GTs expand the receptive field of GNNs, effectively utilize graph structure information, and establish a collaborative framework where each module reinforces the other's strengths [79].

The final group of models we introduce (and that prior works categorize as "others" [35]) have a common important characteristic: learning the embeddings of the two KGs simultaneously [56, 80–82]. We refer to this group as the **interaction training group**. Unlike other methods that embed entire KGs independently and then align entities, interaction training models do not need to embed entire KGs which makes their inference more adaptable to unseen data. Instead, these models embed pairs of entities from both source and target KGs simultaneously capturing interactions between the entities. This is done by comparing the entities features—using techniques like aggregation or averaging—to generate interaction vectors, which are then embedded through Neural Networks or similar techniques. The final predictions are based on a distance margin or threshold. If the interaction embedding of a pair of entities belonging to the source and target KGs is measured to be above the threshold (measured using the vector norms), then the entities are aligned together. The aim is to keep a marginal distance of aligned entities with non-aligned ones. Interaction training methods might use translational or GNN or any other basic model to initially embed the entities but in contrast to the three other groups, these models can provide insights on the correlation between the features of entities belonging to two KGs, whether they are aligned or not.

After analyzing many comparative studies on benchmark datasets, we decided to focus in this study on the following recently-proposed embedding-based EA methods, representative of each of the four groups outlined above: MultiKE [46] (translational model), RDGCN [83] (GNN-based model), i-Align [76] (GT-based model), and BERT-INT (interaction training model). We choose these established models because they are scalable to run on real-world large KGs and have state-of-the-art performance on well-known benchmark datasets [35]. We give more detail on each of them in the following.

MultiKE considers the two distinct KGs to be aligned as one large KG. To connect these two KGs and augment the number of relation triples, the method connects each entity in the source KG to the neighbors of its counterpart entity in the target KG and vice versa by replacing the head and tail entities of each relation triple with their counterparts in the reference alignment. To further enhance the relation triples, the method identifies matching relation and attribute predicates by comparing their literal or relation embeddings and selecting those that exceed a similarity threshold. Once the predicates are matched across the KGs, each relation is replaced by its counterpart, augmenting the relation triples accordingly. Then, it represents each entity and relation using a variant of TransE. To generate the final entity alignment predictions, the model combines these representations with encoded local names of entities and predicates, which are then fed into a Convolutional Neural Network (CNN) [46].

BERT-INT begins by generating initial entity embeddings using a pre-trained BERT-based model, leveraging the entities' descriptions or names/labels. It then constructs a similarity matrix based on the initial embeddings for each pair of training entities. Next, the method creates a neighborhood similarity matrix to co-train each entity pair in the candidate set. For training the interactions of the KGs' structural embeddings, BERT-INT relies exclusively on the direct neighbors of the entities. It is worth noticing that BERT-INT computes interactions between the attributes of entities being compared, rather than simply aggregating attribute information. This approach can reduce the impact of noisy or irrelevant attribute matches. The attribute-view interactions are processed in a unified way along with name/description and neighbor interactions, contributing to the overall alignment decision. It then aggregates all the vectors obtained by the similarity matrices to represent each pair of entities and finalizes the entity pair representations using a Multi-Layer Perceptron (MLP).

Table 1
Comparison of embedding-based EA methods

Method	KG embedding approach	Best-evaluated benchmark dataset	Hit@1 on benchmark dataset	Input features		
				Relation predicate	Attribute predicate	Entity name
MultiKE	Translational	DBP_WD_100K	0.918	Relation name	Attr name/value	Entity name
RDGCN	GNN	DBP15K	0.886 (on FR-EN)	-	-	Entity name
i-Align	Graph Transformer	DBP_YG_15K	0.912	-	Attr name/value	Entity name
BERT-INT	Interaction training	DBP15K	0.992 (on FR-EN)	Relation name	Attr name/value	Entity name/ description

RDGCN leverages the information of relations into entity representations employing a two-step process. First, a dual relation graph is constructed based on the input KG (the context graph), which is nothing but the line graph of the context graph.³ In the dual graph, each node represents a type of relation and two nodes are connected together if they have a common head or tail in the main KG. Then, a graph attention mechanism (GAT) is applied to arouse reciprocal actions between the two graphs. The resulting vertex representations in the context graph are fed to Graph Convolutional Networks (GCN) [84] layers to capture the graph's structural information through a message-passing system. In the last step, the obtained entity representations are used for aligning pairs of entities.

i-Align uses two transformer-based architectures to represent the entities based on their graph structures and textual attribute values. The model uses a graph encoder to aggregate the entities' structural information that can also effectively handle large KGs. The model's other transformer obtains the interconnection between the entity attributes using the embeddings of attribute keys and values as inputs. i-Align provides explanations of the alignment results in the form of a set of the most influential attribute predicates and entity neighbors based on the attention weights of its two transformers.

We summarize the main properties of these four methods in Table 1, including their respective results in terms of Hit@1 measure as reported in the original papers introducing these methods. As we can see, all methods report a Hit@1 of more than 88% exceeding competing methods in the respective studies. MultiKE and i-Align have been evaluated on the DBP_WD_100K [61] and DBP_YG_15K [31] benchmark datasets, respectively, while the BERT-INT and RDGCN – on the DBP15K [27] dataset. All four methods use entity names for embedding as an extra, i-Align utilizes the attribute predicate's names and values in its embedding procedure. To use the maximum descriptive information of entities, BERT-INT employs the entity's descriptions instead of their names when such descriptions exist.

Finally, to be able to compare between non-embedding-based and embedding-based methods, we include in our analyses DLinker [40] as a representative method of the **non-embedding-based group**. The method applies an average aggregation between the similarity measures derived from the instance objects calculated by the longest common sub-sequence algorithm. DLinker has a performance that is close to that of the best-performing system LogMap [85] on several OAEI⁴ entity linking tracks. Furthermore, because it is developed in this paper's authors team, having full control of the tool facilitates further experiments.

In the sequel, we move on to describing and comparing the datasets that we consider in this study, specifically in terms of their different heterogeneity aspects.

4. Datasets and Their Heterogeneity Aspects

In this section, after giving a summary of the datasets we consider and motivating their choice, we study the degrees of their heterogeneities using specific metrics, introduced below. The study showed these datasets to be diverse and highly heterogeneous. Hence, we believe the analysis of the performance of the four chosen models

³The line graph of an undirected graph G is another graph that represents the adjacencies between edges of G . The line graph of the given graph G is constructed by making a node (vertex) instead of each edge in G . Then, for every two edges in G that have a vertex in common, we make an edge between their corresponding vertices in the line graph of G . See https://en.wikipedia.org/wiki/Line_graph

⁴<https://oei.ontologymatching.org/>

(described above) on this particular collection of datasets would give us adequate insights beyond the specific choice of datasets and models and in particular for better understanding the challenges for the EA task when dealing with real-world highly diverse datasets.

4.1. Datasets

We proceed to present and analyze the chosen datasets, coming from the two groups identified in the introduction: synthetic benchmark datasets and real-world datasets.

Benchmark Datasets. We consider DBP15K [80], which is a benchmark dataset that a significant number of state-of-the-art methods report their results on [23, 86]. DBP15K consists of three pairs of KGs that differ in the used language (French, Japanese, and Chinese). We pick the French-English dataset (DBP15K_{FR-EN}) as a sample of the whole. Furthermore, we consider SPIMBENCH,⁵ the OAEI reference instance matching dataset, which is much smaller than DBP15K but similar to it in several heterogeneity aspects that we will discuss below. Additionally, we consider the ICEWS-WIKI and ICEWS-YAGO proposed in [35]. They have the particularity to have been generated from KGs having highly different degree distributions, hence mimicking graphs from real-world scenarios. Hence, these two datasets are highly heterogeneous in structure as compared to DBP15K_{FR-EN} (for full comparison we refer to [35]), where the source and target KGs in each of these two datasets have very different scales.

Real-world Datasets. Because heterogeneity of KGs has a broader meaning than linguistic differences [22] and also, benchmarks often present idealized scenarios with a limited set of relationships, controlled noise, and specific characteristics [87], we added to our investigation two real-world datasets: DOREMUS [5], and AgroLD [88] that differ from benchmarks in terms of the types of their heterogeneity. DOREMUS is a real-world music-related dataset consisting of three interconnected datasets that describe classical music works and the related events and entities. The data is multilingual with a majority of French text and comes from catalogs and archives of three major French cultural institutions (Radio France, La Philharmonie de Paris, and the French National Library) [89]. AgroLD consolidates data relevant to the plant science community, including crops like rice, wheat, and arabidopsis [90]. With approximately 900 million triples, AgroLD is the result of annotating and integrating over 100 datasets from 15 diverse sources [88].

To get an idea of the extent to which the KGs in each dataset differ in scale, we show in Table 2 the sizes of the source and target KG for each dataset (denoted by #S and #T, respectively). The remaining columns of the table will be introduced and explained as we proceed in this section.

4.2. Evaluating the Heterogeneity of the Datasets

We start by taking a bird's view on the datasets and showing the degree distribution of the underlying KGs i.e. the undirected graphs in these datasets in Figure 1. We count the number of entities relative to their degrees and visualize this for degrees up to the point where 90% of the nodes have a degree below that threshold. We also considered visualizing up to the median or median unique degree. However, we believe it is not appropriate to plot up to these values, as the median only represents the point at which 50% of the entities are below or equal, and fewer than 5 entities have the median unique degree.

The figure shows that the number of the nodes having the same degrees are similar in the pair of KGs in both DBP15K and SPIMBENCH datasets, indicating similar degree distributions in these two datasets. However, this is not the case for DOREMUS, AgroLD, ICEWS-WIKI, and ICEWS-YAGO, where we can see that the number of nodes having the same degree is different across the KGs in each dataset.

To get a more in-depth understanding of the the dataset heterogeneities, we proceed to compute three statistical and two qualitative metrics for each pair of KGs in our datasets.

⁵https://hobbit-project.github.io/OAEI_2022.html

Table 2

Comparing each two KGs for each dataset (all numbers indicate percentages except for KG Sizes which indicates the number of entities).

Dataset	JS divergence	Max difference in percentage of nodes	KG Sizes	Size similarity	Reference alignment	
					Levenshtein normalized similarity	EA semantic similarity
DBP15K _{FR-EN}	5.55	1.87	#S 19661 #T 19993	98.3	60.1	90.5
SPIMBENCH	4.41	2.45	#S 2966 #T 3082	96.2	36.6	66.2
ICEWS-WIKI	36.1	6.21	#S 11047 #T 15831	69.78	-	-
ICEWS-YAGO	43.0	10.37	#S 26863 #T 22555	83.9	-	-
DOREMUS	16.8	14.0	#S 2057 #T 1889	92.8	30.3	46.6
AgroLD	6.84	6.22	#S 96117 #T 51488	53.6	19.6	56.6

4.2.1. Jensen–Shannon Divergence

To statistically analyze the underlying distribution of degree sequences in each pair of KGs, we opt for applying the Jensen–Shannon divergence test. We found Jensen–Shannon (JS) divergence or JS distance [91] a suitable statistical metric that captures the amount of overlap between two distributions by using a bi-directional Kullback–Leibler (KL) divergence [92]. KL divergence, defined in equation 1, measures how one reference probability distribution P is different from a second probability distribution Q .

$$D_{KL}(P||Q) = \sum_x P(x) \log\left(\frac{P(x)}{Q(x)}\right). \quad (1)$$

The JS divergence is a symmetrized and smoothed version of KL divergence $D_{KL}(P||Q)$, defined as follows:

$$D_{JS}(P||Q) = \frac{1}{2}D_{KL}(P||M) + \frac{1}{2}D_{KL}(Q||M), \quad (2)$$

where $M = \frac{1}{2}(P+Q)$ is a mixture distribution of P and Q . JS divergence gives us a number in the range $[0, 1]$, where the higher the number, the more divergent the distributions. The results of our study on JS divergence are given in Table 2.⁶ We notice that the degree distributions of KGs in DBP15K_{FR-EN} and SPIMBENCH are less divergent than their counterparts in DOREMUS and AgroLD, as well as that there is a much higher level of heterogeneity in degree distributions of the ICEWS-WIKI and ICEWS-YAGO datasets than the other datasets. This is not surprising, since these two datasets have been generated to mimic the JS ratios of ICEWS, YAGO, and WIKI KGs, which have very different degree distributions.

4.2.2. Maximum Difference in Percentage of Nodes w.r.t Node Degrees

To better recognize the differences in the KGs' degree distributions, we calculate our second statistical metric which measures the maximum difference in the percentage of the entities with respect to the degrees in each pair of KGs. By looking at the second column of Table 2, we can see that the maximum difference in the percentage of the nodes across the KGs (w.r.t. the node degrees) in DOREMUS is much higher than in all other datasets. This confirms the observation in Figure 1, showing that the percentage of entities having the same degree in the two KGs varies less in the benchmark datasets (DBP15K_{FR-EN} and SPIMBENCH), as compared to the other datasets.

⁶Note that the JS measurements have been multiplied by 100 to show the percentage.

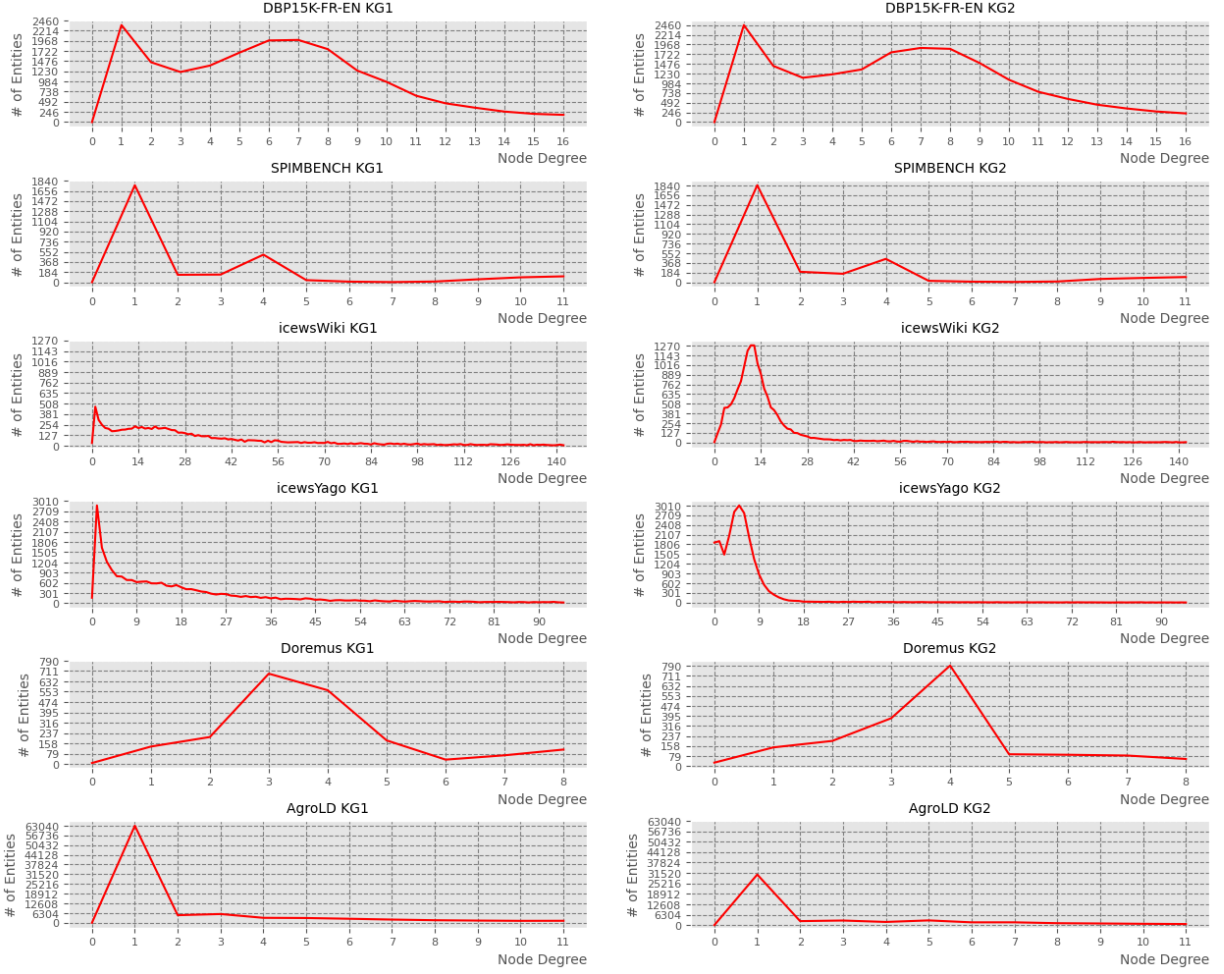


Fig. 1. Degree distribution of each two KGs for each dataset.

4.2.3. Size Similarity

As a third statistical metric, we calculate the normalized difference between the number of entities in every pair of KGs, so we can compute the similarity in the size of KGs using Equation 3:

$$Size_similarity = 1 - \frac{|s(KG_1) - s(KG_2)|}{\max(s(KG_1), s(KG_2))}, \quad (3)$$

where $s(KG)$ returns the size (i.e. number of nodes/entities) of a given KG. We divide the absolute value of the difference in sizes of the KGs by the size of the larger KG. Table 2 shows that the size of the KGs in the benchmark datasets is almost the same, while in real-world cases the two KGs might include very different numbers of entities, as is the case for the AgroLD dataset. The size similarity of 53.6% for AgroLD indicates that the size of one of its KGs is almost twice as bigger as that of other KG. The two KGs in the AgroLD real-world dataset differ more strongly in scale even than their counterparts in the two highly heterogeneous synthetically-generated datasets of ICEWS-WIKI and ICEWS-YAGO.

Analyzing the results of the three statistical features (Table 2), in combination with observing the degree distributions (Figure 1), reveals the higher level of structural heterogeneity in KGs of DOREMUS and AgroLD as compared to the two synthetic benchmark datasets. Moreover, there is less JS distance between the degree distribution of KGs

in benchmark datasets in comparison with the other datasets. Furthermore, in all non-benchmark datasets, ICEWS-YAGO dataset contains KGs having the least overlaps in their degree distributions, and AgroLD is the dataset that includes KGs having the most difference in scales.

Nevertheless, none of these three statistical metrics are indicators of the textual / lingual properties of the entities. We therefore turn our attention to string-level features.

4.2.4. Normalized Levenshtein Similarity

In order to get a more enhanced understanding of how dataset heterogeneities affect each model's performance, we need a qualitative heterogeneity metric, especially, for approaches like BERT-INT, MultiKE, and i-Align that use the textual attribute values of the entities. Even RDGCN, instead of random initialization, uses a representation vector of entity name as the entity's initial embedding. In addition, since all of the four analyzed methods have been trained in a supervised manner, they all use some part (30%) of the reference alignment as their training data. Hence, the quality of data of the reference alignment affects the model's performance directly. Hence, we explore two qualitative metrics - one based on the Levenshtein similarity (in this subsection) and one based on the embeddings' semantic similarity (in the following subsection).

As a first qualitative metric, we get an average over the Levenshtein normalized similarity of the attribute values for all pairs of aligned entities in the reference alignment. Levenshtein, or edit distance, is originally a measure of the closeness of two strings. It quantifies the minimum number of single-character edits (insertions, deletions, or substitutions) required to transform one string into the other [93]. The resulting distance is always in the interval $[0, 1]$. The similarity is computed by subtracting the normalized distance by 1. To analyze the quality of the data, we focus on the reference alignment, and for the first step, we compute the average over the Levenshtein normalized similarity of the attribute values of each pair of aligned entities (cf. results in Table 2).

Because minor variations in the input data do not affect the performance of the language models in embedding the texts [94, 95] and regarding the fact that in EA methods that utilize the attribute values of entities (including three of our employed methods), a language model is used for initial embeddings of the entities [96, 97], we first lemmatized and stemmed all the words in each attribute value. Then, we compared all attribute values for each pair of entities in the reference alignment and we computed the maximum Levenshtein similarity between each pair of attribute values and averaged all. Due to the multilingual nature of DBP15K_{FR-EN}, we used Google Translate to get the English version of the French KG.

The results of the Levenshtein measurements reported in Table 2 indicate that, despite possible translation errors in the DBP15K French KG [83], the normalized Levenshtein similarity of aligned entities in the benchmark datasets is higher than in the real-world datasets. This suggests that EA methods, particularly those relying on textual descriptions of entities, may perform better on benchmark datasets. Since attribute triples are not included in the ICEWS-WIKI and ICEWS-YAGO datasets, calculating the Levenshtein measure based on attribute values is not possible for these datasets.

4.2.5. Semantic Similarity

While the normalized Levenshtein similarity offers insights into the textual closeness of aligned entities in each dataset, it primarily focuses on character-level differences and does not capture the semantic or contextual similarities between entity pairs. Therefore, we further investigate the semantic similarity between the aligned entities in the knowledge graphs.

As mentioned, approaches using the entity or predicate features as input, usually utilize a language model to embed the entities. The studies show that the performance of these methods is relevant to the quality of the initial embeddings [30, 98]. Hence, we want to measure the similarity of two aligned KGs [99, 100] based on initial embeddings of entities in the reference alignment. Because language models capture the semantic similarity of words, we rely on the entities embeddings. We apply the well-known normalized Euclidean relative distance over the pairs of entities of the reference alignment, which is a common choice [32], given as follows:

$$Semantic_similarity(KG_1, KG_2) = 1 - \frac{1}{s} \sum_{i=0}^s f(d_i), \quad (4)$$

where

$$d_i = \frac{\|v(e_i) - v(e'_i)\|}{\|\sum_{j=0}^s v(e_i) - v(e'_j)\|} \quad (5)$$

with s being the size of the set of seed alignments $S = \{(e_0, e'_0), \dots, (e_s, e'_s)\}$, and $v(e_i)$ – the initial representation vector of entity e_i that has been obtained by a pre-trained multilingual BERT model over entity’s descriptions.⁷ The relative distances d_i have been normalized using the MinMax scaler function f :

$$f(d_i) = \frac{d_i - \text{Min}(d_i)}{\text{Max}(d_i) - \text{Min}(d_i)}, \quad i = 1, \dots, s \quad (6)$$

Notice that by definition, having a lower value of the EA semantic similarity for a dataset indicates a higher level of semantic heterogeneity. The corresponding semantic similarities on our pairs of KGs in each dataset are reported in the last column of Table 2. The similarities are calculated based on the initial embeddings of attribute values of the aligned entities using the pre-trained BERT model. Note that this measurement is important when we analyze the performance of the embedding-based EA models which use the entities’ attribute values in their frameworks. As reflected by the semantic similarity results, the real-world datasets have a higher degree of semantic heterogeneity in comparison to the benchmark datasets. This confirms what we already observed with the help of the Levenshtein similarity, showing that from both character-level and conceptual perspectives, aligned entities from the benchmark datasets have more similar textual descriptions than those in real-world datasets.

To visualize how the semantic similarity in different synthetic benchmark and real-world datasets differs, we applied t-SNE [101]. t-SNE, or t-distributed Stochastic Neighbor Embedding is a dimensionality reduction technique commonly used for visualizing high-dimensional data in lower-dimensional space, typically 2D or 3D. In Figure 2, we visualize the entity embedding spaces of the SPIMBENCH and DOREMUS datasets that, according to Table 2, are datasets with a low and a high level of EA semantic similarity, respectively.⁸

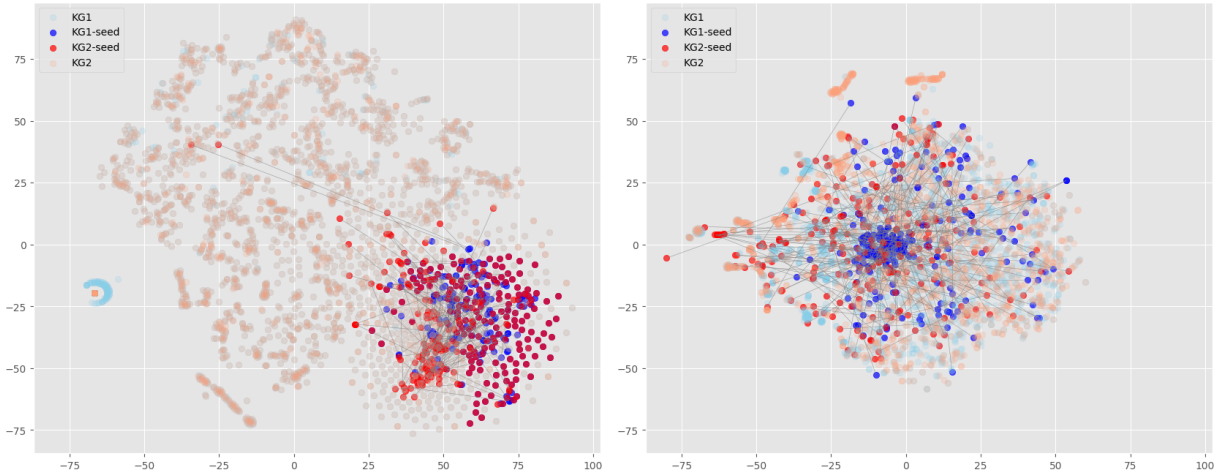


Fig. 2. Reduced-dimension BERT-based Initial entity embeddings of SPIMBENCH (to the left) and DOREMUS (to the right).

The dark blue and red points represent the seed alignment of the KGs, while each entity in the seed alignment is connected to its counterpart using grey lines. Looking at the grey lines that show the distance between the initial

⁷<https://huggingface.co/google-bert/bert-base-multilingual-cased>

⁸Although the amount of semantic similarity for the DBP15K dataset is much higher than that of SPIMBENCH, we visualized SPIMBENCH because it has much fewer entities and this facilitates the visualization.

embeddings of the entities in the reference alignment, one can easily recognize how far the entities are located in the DOREMUS real-world dataset. In SPIMBENCH, only two aligned entities have a long distance, and for the other aligned samples, the distance is much shorter than what we can see for the DOREMUS dataset. It is important to note that in Figure 2, for the SPIMBENCH dataset, several red dots appear in the bottom-right corner, seemingly unlinked. In reality, they are connected to nearly overlapping dark blue dots. The line between them is not visible because the initial embeddings of the source and target entities are so similar that the connection line becomes imperceptible. Additionally, there are some red dots located in the middle-left of the figure, surrounded by light-blue dots. These light-blue dots represent entities in the source KG that have very similar initial embeddings but are not part of the seed alignment, meaning they do not have a corresponding match in the target KG. The plots confirm the higher level of semantic heterogeneity in the real-world DOREMUS dataset as compared to the benchmark dataset SPIMBENCH.

5. Comparative Analysis of the Embedding-based Methods

In this section, we present the results of implementing and applying the selected EA models on the chosen datasets. We first explain the challenges of using the models on real-world and less well-known benchmark datasets and how we overcome these issues, this being part of the lessons learnt in this work. Next, we discuss the evaluation metrics employed by the models and present the results of our experiments. Further on, we provide an overview of how the models perform on both benchmark and real-world datasets. We also investigate how these performances relate to the dataset features in light of the discussion in Section 4. Finally, we look into the inference capacities of the models when facing the full-scale graphs (instead of their corresponding validation sets).

5.1. Datasets Preparation for Applying the EA Models

For each dataset, we have a file of the source KG, a file of the target KG, and a file containing the reference alignments in XML, turtle, or Ntriples format. To feed the data to each model, we prepare a series of files following the naming convention and formats required by each model, e.g. json, pickle, txt, or other. In this process we confronted issues either related to the dataset design itself, e.g. using blank nodes, or related to the model input's design, e.g. when there is no instruction about the proper model input. Even with correctly formatted inputs, runtime errors can still occur unexpectedly due to minor changes in the input data. This necessitates a thorough process of data validation to ensure the models function correctly. Data validation in an ML pipeline ensures that training data is error-free and accurate, preventing issues that could degrade model performance during deployment and safeguarding against errors introduced during data processing [102]. Hence, we need to handle the data lifecycle of inputs to each embedding-based EA model [103], and address as many problems as we face to prepare the suitable data. After writing the codes to prepare the proper input files to the four models that we use and validating them on the different benchmark and real-world datasets, we share the codes on a GitHub page⁹ to pave the way for researchers to benefit from employing these methods on their specific datasets. We included all the links to the original models on our GitHub repository.

Note that, despite the high heterogeneity aspects of ICEWS-WIKI and ICEWS-YAGO, which make them more similar to the real-world KGs, the only model that we employ for them is RDGCN. The reason is that these two datasets lack the attribute triples which are the essential features utilized by DLinker and the other three methods and they only contain the relation triples. We opted not to use additional datasets due to significant structural differences and limited accessibility. The OpenEA datasets for instance were generated under the 1-to-1 assumption, omitting unmatched entities. In response to this limitation, Leone et al. [37] introduced new, more realistic datasets that do not follow this assumption, which we could not access on their repository. While the sampling algorithm code is provided to regenerate the datasets, doing so would result in datasets that differ from DOREMUS and AgroLD, as Leone's datasets are derived from larger knowledge graphs, whereas DOREMUS and AgroLD are not. Therefore, We chose not to use additional datasets to maintain consistency in the real-world data analysis.

⁹https://github.com/dace-dl-anr/Create_Input_Data_to_EA_Models

5.2. Evaluation Metrics

There are two commonly used families of evaluation metrics: (1) Precision and Recall (and the resulting F_1 -score) and (2) Hit@k or, often, Hit@1. Traditionally EA methods rely on metrics in (1), but the advent of embeddings-based methods have moved the focus to metrics of type (2), as they follow the common tradition in link prediction settings. We provide definitions of the two types of metrics and discuss their commonalities and differences in what follows.

We define precision, recall, and F_1 -score in the context of a classification problem: precision is the ratio of true positive predictions to the total number of positive predictions made by the model. It measures the accuracy of the positive predictions made by the model. Recall is the ratio of true positive predictions to the total number of actual positive instances in the dataset. It measures the model's ability to identify all relevant instances. The F_1 -score is the harmonic mean of precision and recall. It balances both precision and recall in one metric. The respective formulas are given as follows:

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}; \quad \text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}; \quad F_1\text{-score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}, \quad (7)$$

where TP stands for True Positives i.e. instances for which the model correctly predicted a positive label, FP stands for False Positives, and FN stands for False Negatives i.e. instances for which the model incorrectly predicted a negative label. Hit@1, on the other side, is defined as:

$$\text{Hit@1} = \frac{\text{Number of times the top-ranked prediction is correct}}{\text{Total number of predictions}}. \quad (8)$$

A Hit@K, analogically, is recorded when the correct entity matching the ground truth appears within the top K positions in the ranked list of predictions.

In this work, we make use of both types of metrics: Hit@k allows us to compare our results with the state-of-the-art EA embeddings-based methods which mainly use this metric to report their performance. However, Leone et al. [37] argue that precision, recall, and F_1 -score better represent the performance of EA models. For that reason, we also report these metrics regarding the models we employed on our datasets. We consider it worthwhile going into more depth regarding the similarities and differences of these two evaluation paradigms. The following discussion contributes to shedding light on the cases when they are considered equivalent, and on those when one of the two types is to be used by preference.

During the validation phase of embedding-based EA methods, since the model generates a ranked list of candidates for each entity, and given that the validation set adheres to the 1-to-1 assumption, the precision, recall, and F_1 -score are equivalent to Hit@1 [30]. In [37], Leone et al. mention that using the Hit@k metric for evaluating the model's performance is based on the unfair 1-to-1 assumption. However, this assumption does not hold in real-world EA datasets. Therefore, for entity alignment, evaluation metrics should be used that can take into account the case where no aligned entity is predicted for the given input entity. Based on the similarity measures (or the predicted alignment probabilities) that the EA methods provide at their last evaluation step, Leone et al. define an assignment module that matches a given entity in the source KG to at most one entity in the target KG (do note that there could be no match for the given source entity). Then, precision and recall are calculated based on their definitions in the field of information retrieval.

A question that might come to mind is why under the 1-to-1 assumption Hit@1 is equivalent to precision, recall, and F_1 -score. While a few hints have been proposed in [30] by Sun et al. and in the appendix C of the technical report of [37] by Leone et al., we believe that detailing on this claim contributes further to the discussion. Current EA methods evaluate model performance on validation sets that include only positive samples, meaning pairs of corresponding entities, which, as a direct consequence, satisfy the 1-to-1 assumption, i.e. for every given entity $e_l \in \text{Source_KG}$ included in the validation set, there exists exactly one entity $e'_l \in \text{Target_KG}$ where $e_l \equiv e'_l$ (e_l is the same as e'_l). Suppose n is the size of the validation set (by size, we mean the number of samples included in the validation set), and suppose $\forall i, i = 1, \dots, n$, in the validation set we have $e_i \equiv e'_i$. Hence, because the 1-to-1

assumption holds in the validation set, for entity $e_i \in Source_KG$ we have: $\forall i | i \neq l, e_i \neq e'_i, e'_i \in Target_KG$. To show that the precision and recall (and the F_1 -score) are equal, based on the definitions, it is enough to show $FP = FN$. Afterwards, if we show that the total number of the predictions (the denominator in the Hit@1 formula) equals $TP + FP$, we have shown that Hit@1 equals both precision and recall. Suppose that the class of aligned and non-aligned entities are the positive and negative classes, respectively. In Table 3, we formally define TP, FP, FN, and TN (True Negative) in the context of EA models that predict a ranked list for evaluating the model's performance under the 1-to-1 assumption in their validation set.

Table 3

Defining TP, FP, TN, and FN for EA models based on Hit@1 predictions, when the validation set of size n is free of the negative samples and meets the 1-to-1 assumption, where $i, j \in 1, \dots, n$.

Actual \ Predicted	Aligned (P)	Non-Aligned (N)	Implicitly Non-Aligned (IN)
Aligned (P)	TP: $e_i \equiv e'_i$	FN: $e_i \neq e'_i$	$e_i \equiv e'_j,$ $ i \neq j$
Non-Aligned (N)	FP: $e_i \equiv e'_j,$ $ i \neq j$	TN: $e_i \neq e'_j,$ $ i \neq j$	$e_i \equiv e'_i$

Considering Hit@1 as the final prediction by EA models, for each entity $e_i \in Source_KG$ included in the validation set, the model prediction is either its correct match (TP) which is e'_i , or a wrong match (FP) like $e'_j, j \neq i$ in $Target_KG$, and $i, j \in 1, \dots, n$. Hence, it is trivial that $TP + FP$ equals the total number of predictions. Furthermore, while we don't have any trivial non-aligned predictions by the model, each Hit@1 implicitly gives us some non-aligned predictions. For example, for a false negative, based on the definition, the model should predict that $e_i \neq e'_i$ for some $i \in 1, \dots, n$. While the model never predicts a false negative explicitly, but each time that the model predicts that $e_i \equiv e'_j$ where j indicates any index other than i , the 1-to-1 assumption which has been met in the validation set implies that e_i in $Source_KG$ is not aligned with any other entity than e'_j in the $Target_KG$ including the e'_i . Hence, whenever the model predicts $e_i \equiv e'_j | i \neq j$ which is a false positive, following the 1-to-1 assumption, it implicitly has predicted that $e_i \neq e'_i$ which is a false negative. Accordingly, $FP = FN$, and we have shown that precision equals the recall (and the F_1 -score). Furthermore, because the number of times the top-ranked prediction is correct equals the TP, and as we have earlier shown the denominator of the precision and Hit@1 are equal, we can see that precision equals the Hit@1. As a result, all the aforementioned metrics are equivalent whenever the 1-to-1 assumption has been met during the evaluation.

5.3. Evaluation Tasks and Analyses

5.3.1. Real-world vs. Benchmark datasets

The comparative performance of the selected models on the collection of chosen datasets is given in Table 4. Following the discussion presented in Section 5.2, we compare the Hit@1 results of embedding-based models with DLinker, considering cases where the 1-to-1 assumption was met during the evaluation of embedding-based EA models on each dataset. The best-performing method for each dataset is highlighted in bold. We can observe an overall drop of the performance of embedding-based models when tested on real-world datasets, as compared to benchmark ones. In what follows, we discuss the results of each of the models of interest in light of that global observation, while also considering the internal mechanisms of each of the models that differentiate them from one another and could provide insights into this observations.

The performance of the BERT-INT model is strong on datasets like DBP15K and SPIMBENCH, achieving high Hit@1 rates (99.3% and 82.4%, respectively). However, its performance drops significantly on our real-world datasets DOREMUS and AgroLD (Hit@1 rates of 47.9% and 21.1%, respectively). The reason for this drop is that BERT-INT relies heavily on the quality and amount of textual information (entity descriptions), as observed in Table 1. Hence, datasets with less textual and semantic similarity and fewer descriptive features, like DOREMUS and AgroLD (Table 2), lead to a decrease in its performance. This emphasizes the importance of high-quality data descriptions for BERT-INT's success.

Observing the results of RDGCN (Table 4), we can see that its performance is more than 88% and 77% on DBP15K and SPIMBENCH, respectively, and it also drops significantly in the real-world scenarios (close to 0% Hit@1). RDGCN relies solely on graph structure and a GloVe word embedding¹⁰ on entity names (see Table 1). This, along with the statistical metrics from Table 2, which highlight the greater structural heterogeneity of the real-world dataset, helps explain this outcome. Since RDGCN is adaptable to different initial embeddings, we modified the initial embedding to the model to use a multilingual pre-trained BERT model to generate the initial embeddings for entity names. This adjustment improved the Hit@1 and Hit@10 scores on the dataset by only 0.4% and 6%, respectively. This modest improvement is likely due to the fact that most entity names in are represented by IDs rather than meaningful text, which limits the impact of the embeddings. We have uploaded the code for generating the initial embeddings with BERT using GPUs to our GitHub repository. We also experimented with using BERT-based embeddings of the entities' descriptions (used in BERT-int) as the initial embeddings in RDGCN. To facilitate this, we implemented a method to create a dictionary of entities' descriptions for any pair of given RDF graph, which is available in our repository. This approach led to improvements in the percentage of the Hit@1 scores across several datasets: 93.70 on DBP15k, 99.53 on Spimbench, 22.49 on DOREMUS, and 7.21 on AgroLD. These represent enhancement percentages of 5.1%, 21.83%, 21.16%, and 7.19%, respectively. This demonstrates that more informative initial embeddings significantly boost RDGCN's performance.

Looking at Figure 1, we can see the long-tailed issue of AgroLD's KGs. The long-tailed problem in graphs [104, 105] is described as an issue where a small number of nodes have a substantial number of neighbors, while the majority (referred to as tail nodes) have only a few neighbors [106]. GNNs used in RDGCN under-represent tail nodes during the training of the model and lead to a low-quality KG embedding [107] and this can explain the drop of performance on this dataset for RDGCN. However, as we observe in Figure 1, SPIMBENCH also has a long tail problem, which does not appear to be an issue for RDGCN. We found two main differences between SPIMBENCH and AgroLD that could explain the drop in performance from one to another of this method. (1) The number of common neighbors: In the SPIMBENCH dataset, many entities in the reference alignment share common neighbors. On average, 48% of the entities in the reference alignment, across its two knowledge graphs, have at least one common neighbor with their linked entities. According to the results presented in Wang et al. [108], a higher number of common neighbors improves the efficiency of embeddings generated by GCNs for these entities. However, in the AgroLD KGs, none of the linked entity pairs share a common neighbor. As Wang et al. [108] demonstrated, the performance of GNN-based graph embedding models, including GCNs, correlates more strongly with the number of common neighbors than with node degrees. This suggests that the lack of common neighbors in AgroLD could negatively impact the performance of the RDGCN model. (2) The KGs in AgroLD are bipartite: Giamphy et al. [109] discuss how GNN-based graph embedding of a large bipartite graph is difficult due to the challenge of merging heterogeneous node and graph-level information while ensuring scalability to handle the graph's increasing size. They also propose a list of available resources that perform better on bipartite graph embedding (but unfortunately, we found none of them working on the multi-relational graph embedding which is our case). Moreover, RDGCN uses word embedding models to produce the initial embeddings of the entities using entity names. Because the names (which are the last part of the entity URIs by the model's default) of the musical works and the proteins/genes in DOREMUS and AgroLD have been defined by IDs in their respective ontologies, the initial entity embeddings would not be able to guide the embedding module to a better result. All the aforementioned observations can explain the low performance of RDGCN on DOREMUS (1.2% of Hit@1) and AgroLD (less than 1% of Hit@1). Furthermore, the results of RDGCN on the ICEWS-WIKI and ICEWS-YAGO datasets suggest again the contribution of the high-quality entity names to improving the model's performance. As a result, these two datasets are structurally less complex for the model compared to real-world datasets.

Although MultiKE outperforms several EA Translational-based methods [46] using multi-view KG embedding technique, this model overall performs the weakest among the employed embedding- and non-embedding-based models on the selected datasets. Similarly to its predecessors (Table 4), MultiKE's performance also drops for DOREMUS and AgroLD. Recall that we observe a higher level of structural and qualitative heterogeneities in these two real-world datasets than in the benchmark datasets (Table 2). Hence, the fact that MultiKE relies on both the

¹⁰<https://nlp.stanford.edu/projects/glove/>

Table 4

Measured evaluation metrics of analyzed EA models on the datasets (All numbers indicate percentages). Following the 1-to-1 assumption during the models' evaluation, Hit@1 equals the precision, recall, and F_1 -score for each model.

		Methods								
		BERT-INT		RDGCN		MultiKE		i-Align		DLinker
		Hit@1	Hit@10	Hit@1	Hit@10	Hit@1	Hit@10	Hit@1	Hit@10	
Datasets	DBP15K _{FR-EN}	99.3	99.8	88.6*	95.7*	37.5	43.6	26.6	43.2	-
	SPIMBENCH	82.4	82.4	77.7	94.7	57.1	57.1	75.0	86.5	70.2
	ICEWS-WIKI	-	-	75.1	84.2	-	-	-	-	-
	ICEWS-YAGO	-	-	68.3	80.8	-	-	-	-	-
	DOREMUS	47.9	64.1	1.33	5.92	2.70	8.70	53.1	68.0	95.6
	AgroLD	21.1	33.2	0.02	0.3	2.30	5.7	4.4	12.1	59.0

* The reported numbers are derived from the original study.

graph structure and textual information of entities and their attributes (Table 1) can explain the gap in the performance of this model. Furthermore, while DBP15K is less heterogeneous than the SPIMBENCH, we suspect the reason for the worse performance of MultiKE on this dataset being the multilinguality that could not be handled using a pre-trained English word2vec model¹¹ that MultiKE is employing for the entities' local name embeddings. To embed the French language in Doremus and DBP15K using MultiKE, we were unable to use a pre-trained multilingual BERT model due to the method's strong reliance on word2vec. Instead, we added a French word2vec dictionary to the existing English one, which led to significant improvements in MultiKE's performance on the DOREMUS dataset. Specifically, Hit@1 increased to 30.7% and Hit@10 rose to 34.4%, which seems reasonable given the predominance of French text in DOREMUS. However, the inclusion of this multilingual collection significantly reduces MultiKE's performance on DBP15K, with Hit@1 and Hit@10 dropping to 0.53% and 3.18%, respectively. This represents a decrease of 37% in Hit@1 and 40.4% in Hit@10. We believe this decline is due to a lack of mappings between the French and English word vectors, which causes conflicts in the embeddings of the two languages.

Finally, i-Align uses two transformer encoders for text and graph embeddings. As Table 4 shows, it performs better on SPIMBENCH and DOREMUS (75.0% and 53.1% of Hit@1, respectively) as compared to DBP15K (26.6% of Hit@1) and its performance drops significantly when it comes to AgroLD (4.4% of Hit@1). We suspect that the reason for the model performing worse on DBP15K as compared to DOREMUS is the fact that only the first ten characters of the attribute values were considered, while the rest of the sequence was ignored by the textual transformer-based encoder. Due to the curse of multilinguality issue by transformers [110, 111] and inter-language competition for the model parameters, it seems this limited amount of data may not suffice to train the same text transformer's parameters. Additionally, during our experiments, we discovered that reducing the length of textual properties of the entities in the BERT-INT model can result in a significant reduction in performance by as much as 19%. This again illustrates the importance of retaining the informative attribute descriptions included in the values.

As a baseline of non-embedding-based approaches, we used the DLinker method. Because this model fundamentally finds the longest common subsequence in the descriptions of a pair of entities belonging to two different KGs, it does not support entity alignment on the multilingual dataset of DBP15K_{FR-EN}. Moreover, by comparing the Hit@1 of the embedding-based EA models (Table 4), DLinker is not the best-performing method on the SPIMBENCH, but it shows the top performance on the real-world DOREMUS and AgroLD datasets.

Furthermore, since DLinker finds the alignments based on the greedy strategy of finding the longest common subsequence and ignoring the rest of structural or literal information, and since DLinker's performance is significantly better than the embedding-based methods, we can conclude that in cases where we have real-world data, taking the extra volume of information into account does not help the quality of the embeddings but enforces more noise to them. Although this conclusion holds for all categories of embedding-based EA methods, the Translational and GNN-based methods, which rely primarily on graph structure, introduce more noise into the embeddings. However, i-Align also embeds the graph structures using a graph transformer to embed the local subgraphs containing the

¹¹<https://fasttext.cc/docs/en/english-vectors.html>

nodes that have high interconnectivity with the given entities in two KGs. By comparing the performance of i-Align with RDGCN and MultiKE, we recognize that using the transformer attention mechanism for local subgraph embedding propagates noise less than the GNN message passing and Translational systems in the more heterogeneous real-world cases (See results of these three methods on and AgroLD datasets in Table 4). The other reason that i-Align has a better performance than the RDGCN and MultiKE seems to be the fact that it relies more than those on the literals and textual properties of the entities, as we see that both i-Align and BERT-INT perform better than the other methods for the real-world cases. Furthermore, by comparing the performance of i-Align and BERT-INT on AgroLD, we can conclude that an interaction training method which uses almost all textual properties of the entities is the best choice for the case where we have structurally and semantically heterogeneous large-scale KGs. Because the interaction training methods mostly rely on the comparison between the properties of pairs of entities belonging to two KGs rather than comparing them as particles of the large KGs they belong to, we can conclude that for doing an embedding-based EA task on real-world datasets, a local comparison of the given entities in two KGs will guide the model to predict higher quality alignments.

Generalizability Assessment. Inspired by prior work on domain generalization [112, 113], we examine the robustness of existing EA models by evaluating their performance on datasets that differ from synthetic benchmarks in structure and semantic similarity levels. However, unlike these studies that focus on training on one domain and testing on a different one, we investigate how well-established EA methods perform when applied to more heterogeneous and less curated datasets, without altering their training or optimization procedures. To quantify this, we compute the average Hit@1 on the real-world datasets DOREMUS and AgroLD. As shown in Table 4, BERT-INT reaches an average Hit@1 of 34.5%, while i-Align, MultiKE, and RDGCN score 28.75%, 2.5%, and 0.675%, respectively. This superior performance suggests that interaction-based models like BERT-INT generalize more effectively to heterogeneous real-world scenarios compared to structure-dependent embedding models, although further investigation would be needed to confirm their generalizability across other domains. Overall, our results show that even though embedding-based models perform very well on some benchmark datasets (e.g. 99.3% of Hit@1 for BERT-INT on DBP15K_{FR-EN}), it seems that they overfit on the benchmark data.¹² Consequently, using these models could lead to errors in producing alignments in heterogeneous real-world datasets. Hence, we observed how dataset features presented in Table 2 together with the results of our experiments in this section can justify the gap between the performance of the selected methods on benchmark and real-world datasets.

5.3.2. Analyzing the Models' Effectiveness of Inference

In this section, we focus on the capabilities of models in the inference phase. As a common practice, under the 1-to-1 assumption, models like RDGCN and BERT-INT consider only a subset of the reference alignment as a validation set, during the model evaluation, ignoring the rest of the space. That corresponds to the subspace of dark-colored points in Figure 2 (reference alignments). Such under-representation of the search space undermines the reliability of the reported results, as well as the efficiency of these methods in predicting correct alignments beyond the validation set. Indeed, some real-world studies have removed the 1-to-1 assumption in dataset generation, allowing for more complex scenarios with non-matchable entities. However, many EA models still only focus on data that involve ground truth entities, sometimes even using only ground truth for training. As a result, these models fail to consider the non-matchable entities added to the dataset as candidates. This means the model's performance remains the same as if it were working under the 1-to-1 assumption, because it doesn't effectively handle the additional non-matchable data. In Figure 3, we illustrate the comparison space of EA models that impose the 1-to-1 assumption during evaluation. In this context, the similarity matrix used to identify the top-ranked predictions is a square matrix, as depicted by the green one in Figure 3. This matrix excludes comparisons between entities in the validation set and other entities in the knowledge graphs, such as non-matchable entities and those utilized during training. Later in this section, when we extend the comparison space—first to compare source-to-target and then to compare target-to-source, the results show a decrease in Hit@1. This suggests that the best-predicted match in the extended comparison is not always the same as in the more restricted case. Moreover, for certain entities in the

¹²Benchmark overfitting, meant as models being too good on benchmark datasets and less so on unseen real-world ones [39], is not to be confused with data overfitting of models while training.

validation set, their most likely alignment may actually be outside the validation set, highlighting a lack of efficient embedding for these entities.

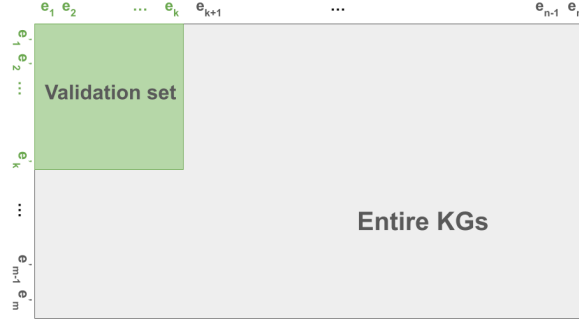


Fig. 3. Illustration of the search (comparison) spaces of the EA models in the limited case (imposing the 1-to-1 assumption) and in the extended case (entire graphs). n and m are the sizes of source and target KGs, respectively.

Therefore, in what follows we assess the models' performance on two versions of the datasets: (1) a limited validation set scenario, and (2) an extended scenario where each source entity's candidate search space includes a broader set of entities from the target KG, rather than being restricted to only those in the validation set. In Figure 4, we visualized the results of our experiments on the performance of BERT-INT and RDGCN. We utilized the repository provided by Leone et al. in [37] to measure F_1 -score in the extended case. However, we encountered difficulties in replicating its original performance when applied to the RDGCN method. As a result, we re-implemented its assignment module to ensure accurate computation of these metrics. It is important to note that the module presented in [37], excludes the portion of the ground truth that had been seen during training as alignment candidates for the entities in the validation set. In contrast, we applied a stricter condition: we considered all entities from the target KG (including the ground truth entities used during training) as candidates for aligning with a given entity from the source KG in the validation set. We make two main observations: although the Hit@1 measure decreases from the limited to the extended search space case, this decrease is not significant. However, if we look at the F_1 -score, the situation is drastically different, where an important drop of performance of the two models can be observed from the limited to the extended search space scenario.

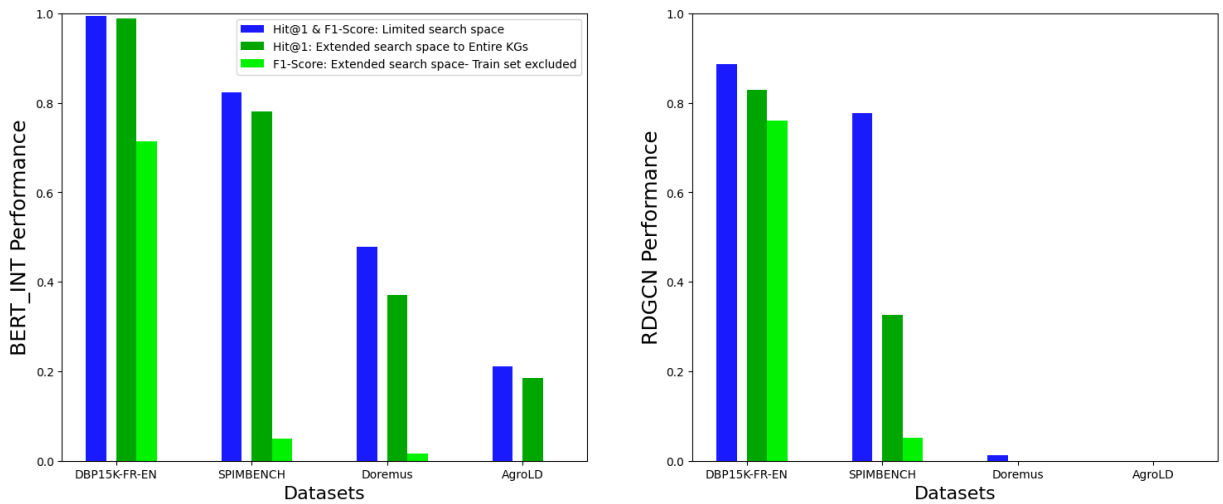


Fig. 4. Hit@1 and F_1 -score of BERT-INT and RDGCN models on the validation data considering the candidate search space limited to the reference alignment (shown in blue) or extended to the whole KG space (shown in green).

As illustrated by the comparison between the blue and dark-green bars in Figure 4, there is a performance decline of 5.66% and 45.15%, as measured by Hit@1, for RDGCN on the DBP15K_{FR-EN} and SPIMBENCH datasets, respectively, when the search space is extended. This drop is of 0.5% and 4.3% for the same datasets, respectively, when it comes to using BERT-INT. Due to RDGCN’s very low performance in the limited scenario, this is not surprising that extending the candidates’ search space of DOREMUS and AgroLD causes Hit@1 of RDGCN to drop to zero. Furthermore, extending this space on DOREMUS and AgroLD causes Hit@1 of BERT-INT to drop by 10.8% and 2.6%, respectively. However, as Leone et al. [37] pointed out, precision, recall, and F_1 -score are fairer metrics for evaluating EA tasks. To assess this, we used their repository¹³ to test BERT-INT and we re-implement the assignment module to calculate these metrics for RDGCN, by expanding the search space and removing the 1-to-1 assumption. Note that in [37], the candidates are chosen from both the validation set and non-matchable entities, not the entire KGs. The results of this experiment are visualized in light-green bars of Figure 4.

Recall that Hit@1, as shown in Table 4, is equivalent to F_1 -score under the 1-to-1 assumption, i.e. in the limited search space case, while this equivalence does not hold anymore in the extended case. Comparing the blue and light-green bars in Figure 4 shows that extending the candidates’ search space to include also the non-matchable entities in the KGs causes the F_1 -score of RDGCN to drop by 12.6%, 72.6%, 1.33%, and 0.02% on DBP15K_{FR-EN}, SPIMBENCH, DOREMUS, and AgroLD, respectively. Similarly, BERT-INT’s performance is decreased by 27.9%, 77.4%, 46.2%, and 21.1% for the same datasets, respectively. Although we applied a more strict extension condition, looking at Figure 4, we can see that for the vast majority of cases, Hit@1 is still higher than the F_1 -score of the models in the extended case. Since the reduction in F_1 -score is more notable than the Hit@1, and since the assignment module in [37] predicts a pair of entities as an alignment only in the case that those entities are predicted as counterparts in source-to-target and target-to-source predictions, our results suggest that these models have difficulty making symmetric predictions. These results also show that embedding-based EA models still encounter generalizability issues and need improvements in order to be able to find alignments in the realistic search space of the knowledge graphs.

6. Conclusion and Future Work

The objective of this work is to build upon and complement recent empirical studies in the field of embedding-based entity alignment (EA) [30–32, 37, 57], offering a critical perspective on the different models and their limitations, particularly in relation to the challenges posed by various types of datasets and the evaluation process. Therefore, we aim for this study to open new methodological avenues, without focusing on proposing a new model.

We conducted an in-depth analysis of the features of several real-world datasets compared to popular benchmark datasets. Also, we presented an empirical study analyzing the performance of embedding-based EA models beyond test data and on real-world heterogeneous data. We observed that a number of entity alignment embedding-based models like BERT-INT and RDGCN with very strong performance for the task of entity alignment on the well-known DBP15K dataset, suffer a drop in performance on real-world data with heterogeneous textual properties. Hence, the results of our study shed light on the benchmark overfitting issue of EA methods discussed in [39, 114] i.e. the scenario where the model is tuned excessively to perform well on specific benchmark datasets or evaluation metrics, at the expense of its generalization ability to new, unseen real-world data.

It appears challenging to identify a single structure-related meta-feature that accounts for the performance drops of all methods across different datasets, as each method captures the structure from a different perspective. Since, however, heterogeneity is not just limited to diversity in size and degree distribution, we observed the semantic similarity over the reference alignments to be well-correlated with the performance of EA models that employ a language model, helping explain the performance issues. Then, by investigating the reasons for performance fluctuations of EA models regarding the heterogeneities of real-world datasets, we found interaction training models better fit for driving the EA task on real-world especially large-scale scenarios. Although interaction training models showed promise on real-world data, we could not conduct a deeper analysis of how they handle noise and ambiguity, due to time and resource constraints. We leave this as an important direction for future work.

¹³<https://github.com/epfl-dlab/entity-matchers>

Most of the existing embedding-based EA methods simplify the inference process by considering the 1-to-1 assumption [23] and use just a limited portion of the embedding space for the evaluation. This seems to be neither fair nor practical when it comes to using them to discover unseen alignments. As a result, there is a need to go toward an inductive learning EA approach, in which models are trained on pairs of entities from two aligned KGs to predict alignments between unseen entities belonging to the same KGs, as well as matches between entities in other unseen KGs. By addressing this challenge, we believe that EA models will be able to uncover a significantly larger number of alignments across different pairs of knowledge graphs.

Acknowledgements

This work is (partially) supported by the French National Research Agency under the ANR DACE-DL project, grant number ANR-21-CE23-0019. Our team was granted access to the HPC resources of IDRIS under the allocation 2024-AD010315119R1 made by GENCI.

The authors acknowledge the ISO 9001 certified IRD i-Trop HPC (South Green Platform) at IRD montpellier for providing HPC resources that have contributed to the research results reported within this paper.

References

- [1] S. Ji, S. Pan, E. Cambria, P. Marttinen and S.Y. Philip, A survey on knowledge graphs: Representation, acquisition, and applications, *IEEE transactions on neural networks and learning systems* **33**(2) (2021), 494–514.
- [2] P.A. Bonatti, S. Decker, A. Polleres and V. Presutti, Knowledge graphs: New directions for knowledge representation on the semantic web (dagstuhl seminar 18371), in: *Dagstuhl reports*, Vol. 8, Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik, 2019.
- [3] V. Ryen, A. Soylu and D. Roman, Building semantic knowledge graphs from (semi-) structured data: a review, *Future Internet* **14**(5) (2022), 129.
- [4] B. Villazón-Terrazas, F. Ortiz-Rodríguez, S.M. Tiwari and S.K. Shandilya, Knowledge graphs and semantic web, *Communications in Computer and Information Science* **1232** (2020), 1–225.
- [5] M. Achichi, P. Lisena, K. Todorov, R. Troncy and J. Delahousse, DOREMUS: A graph of linked musical works, in: *The Semantic Web–ISWC 2018: 17th International Semantic Web Conference, Monterey, CA, USA, October 8–12, 2018, Proceedings, Part II* 17, Springer, 2018, pp. 3–19.
- [6] V.A. Carriero, A. Gangemi, M.L. Mancinelli, L. Marinucci, A.G. Nuzzolese, V. Presutti and C. Veninata, ArCo: The Italian cultural heritage knowledge graph, in: *The Semantic Web–ISWC 2019: 18th International Semantic Web Conference, Auckland, New Zealand, October 26–30, 2019, Proceedings, Part II* 18, Springer, 2019, pp. 36–52.
- [7] J. Dou, J. Qin, Z. Jin and Z. Li, Knowledge graph based on domain ontology and natural language processing technology for Chinese intangible cultural heritage, *Journal of Visual Languages & Computing* **48** (2018), 19–28.
- [8] E. Marchand, M. Gagnon and A. Zouaq, Extraction of a knowledge graph from French cultural heritage documents, in: *ADBIS, TPD and EDA 2020 Common Workshops and Doctoral Consortium: International Workshops: DOING, MADEISD, SKG, BBIGAP, SIMPDA, AIMinScience 2020 and Doctoral Consortium, Lyon, France, August 25–27, 2020, Proceedings* 24, Springer, 2020, pp. 23–35.
- [9] G. Sanou, V. Giudicelli, N. Abdollahi, S. Kossida, K. Todorov and P. Duroux, IMGT-KG: A Knowledge Graph for Immunogenetics, in: *International Semantic Web Conference*, Springer, 2022, pp. 628–642.
- [10] D.N. Nicholson and C.S. Greene, Constructing knowledge graphs and their biomedical applications, *Computational and structural biotechnology journal* **18** (2020), 1414–1428.
- [11] P. Ernst, A. Siu and G. Weikum, Knowlife: a versatile approach for constructing a large knowledge graph for biomedical sciences, *BMC bioinformatics* **16** (2015), 1–13.
- [12] D.R. Unni, S.A. Moxon, M. Bada, M. Brush, R. Bruskiewich, J.H. Caufield, P.A. Clemons, V. Dancik, M. Dumontier, K. Fecho et al., Biolink Model: A universal schema for knowledge graphs in clinical, biomedical, and translational science, *Clinical and translational science* **15**(8) (2022), 1848–1855.
- [13] A. Tchechmedjiev, P. Fafalios, K. Boland, M. Gasquet, M. Zloch, B. Zapilko, S. Dietze and K. Todorov, ClaimsKG: A knowledge graph of fact-checked claims, in: *The Semantic Web–ISWC 2019: 18th International Semantic Web Conference, Auckland, New Zealand, October 26–30, 2019, Proceedings, Part II* 18, Springer, 2019, pp. 309–324.
- [14] Z. Wang, T. Chen, J. Ren, W. Yu, H. Cheng and L. Lin, Deep reasoning with knowledge graph for social relationship understanding, *arXiv preprint arXiv:1807.00504* (2018).
- [15] L. Cao, H. Zhang and L. Feng, Building and using personal knowledge graph to improve suicidal ideation detection on social media, *IEEE Transactions on Multimedia* **24** (2020), 87–102.
- [16] M. Kejriwal, J.F. Sequeda and V. Lopez, Knowledge graphs: Construction, management and querying., *Semantic Web* **10**(6) (2019), 961–962.

- [17] S.R. Bader, I. Grangel-Gonzalez, P. Nanjappa, M.-E. Vidal and M. Maleshkova, A knowledge graph for industry 4.0, in: *The Semantic Web: 17th International Conference, ESWC 2020, Heraklion, Crete, Greece, May 31–June 4, 2020, Proceedings 17*, Springer, 2020, pp. 465–480.
- [18] A. Ferrara, A. Nikolov and F. Scharffe, Data linking for the semantic web, *International Journal on Semantic Web and Information Systems (IJSWIS)* 7(3) (2011), 46–76.
- [19] V. Beretta, K. Boland, L.L. Seen, S. Harispe, K. Todorov and A. Tchechmedjiev, Can Knowledge Graph Embeddings Tell Us What Fact-checked Claims Are About?, in: *Workshop on Insights from Negative Results in NLP*, Association for Computational Linguistics, 2020.
- [20] X. Zou, A survey on application of knowledge graph, in: *Journal of Physics: Conference Series*, Vol. 1487, IOP Publishing, 2020, p. 012016.
- [21] A. Saha, V. Pahuja, M. Khapra, K. Sankaranarayanan and S. Chandar, Complex sequential question answering: Towards learning to converse over linked question answer pairs with a knowledge graph, in: *Proceedings of the AAAI conference on artificial intelligence*, Vol. 32, 2018.
- [22] M. Achichi, Z. Bellahsene, M.B. Ellefi and K. Todorov, Linking and disambiguating entities across heterogeneous RDF graphs, *Journal of Web Semantics* 55 (2019), 108–121.
- [23] K. Zeng, C. Li, L. Hou, J. Li and L. Feng, A comprehensive survey of entity alignment for knowledge graphs, *AI Open* 2 (2021), 1–13.
- [24] R.C. Salazar, C. Jonquet and D. Symeonidou, Classification of Linking Problem Types for linking semantic data, in: *SEMANTICS*, 2023.
- [25] F. Ardjani, D. Bouchiha and M. Malki, Ontology-Alignment Techniques: Survey and Analysis., *International Journal of Modern Education & Computer Science* 7(11) (2015).
- [26] P. Shvaiko and J. Euzénat, Ontology matching: state of the art and future challenges, *IEEE Transactions on knowledge and data engineering* 25(1) (2011), 158–176.
- [27] Z. Sun, W. Hu and C. Li, Cross-lingual entity alignment via joint attribute-preserving embedding, in: *The Semantic Web–ISWC 2017: 16th International Semantic Web Conference, Vienna, Austria, October 21–25, 2017, Proceedings, Part I 16*, Springer, 2017, pp. 628–644.
- [28] Y. Bengio, A. Courville and P. Vincent, Representation learning: A review and new perspectives, *IEEE transactions on pattern analysis and machine intelligence* 35(8) (2013), 1798–1828.
- [29] W.L. Hamilton, R. Ying and J. Leskovec, Representation learning on graphs: Methods and applications, *IEEE Data Eng. Bull.* 40 (2017), 52–74.
- [30] Z. Sun, Q. Zhang, W. Hu, C. Wang, M. Chen, F. Akrami and C. Li, A benchmarking study of embedding-based entity alignment for knowledge graphs, *Proc. VLDB Endow.* 13(12) (2020), 2326–2340.
- [31] R. Zhang, B.D. Trisedya, M. Li, Y. Jiang and J. Qi, A benchmark and comprehensive survey on knowledge graph entity alignment via representation learning, *The VLDB Journal* 31(5) (2022), 1143–1168.
- [32] N. Fanourakis, V. Efthymiou, D. Kotzinos and V. Christophides, Knowledge Graph Embedding Methods for Entity Alignment: An Experimental Review, *Data Mining and Knowledge Discovery* 37(5) (2023), 2070–2137.
- [33] J. Huang, J. Wang, Y. Li and W. Zhao, A Survey of Entity Alignment of Knowledge Graph Based on Embedded Representation, in: *Journal of Physics: Conference Series*, Vol. 2171, IOP Publishing, 2022, p. 012050.
- [34] J. Euzénat, M.-E. Roşoiu and C. Trojahn, Ontology matching benchmarks: generation, stability, and discriminability, *Journal of web semantics* 21 (2013), 30–48.
- [35] X. Jiang, C. Xu, Y. Shen, F. Su, Y. Wang, F. Sun, Z. Li and H. Shen, Rethinking GNN-based Entity Alignment on Heterogeneous Knowledge Graphs: New Datasets and A New Method, *arXiv preprint arXiv:2304.03468* (2023).
- [36] J. Lehmann, R. Isele, M. Jakob, A. Jentzsch, D. Kontokostas, P.N. Mendes, S. Hellmann, M. Morsey, P. Van Kleef, S. Auer et al., Dbpedia—a large-scale, multilingual knowledge base extracted from wikipedia, *Semantic web* 6(2) (2015), 167–195.
- [37] M. Leone, S. Huber, A. Arora, A. García-Durán and R. West, A critical re-evaluation of neural methods for entity alignment, *Proc. VLDB Endow.* 15 (2022), 1712–1725.
- [38] M. Ben Ellefi, Z. Bellahsene, J.G. Breslin, E. Demidova, S. Dietze, J. Szymański and K. Todorov, RDF dataset profiling—a survey of features, methods, vocabularies and applications, *Semantic Web* 9(5) (2018), 677–705.
- [39] K. Todorov, Datasets First! A Bottom-up Data Linking Paradigm., in: *ISWC (Satellites)*, 2019, pp. 338–342.
- [40] B.G. Happi Happi, G. Fokou Pelap, D. Symeonidou and P. Larmande, DLinker Results for OAEI 2022, in: *17th International Workshop on Ontology Matching, OM 2022, CEUR Workshop Proceedings*, Vol. 3324, CEUR-WS, 2022, pp. 166–173.
- [41] Q. Wang, Z. Mao, B. Wang and L. Guo, Knowledge graph embedding: A survey of approaches and applications, *IEEE Transactions on Knowledge and Data Engineering* 29(12) (2017), 2724–2743.
- [42] A. Sharma, P. Talukdar et al., Towards understanding the geometry of knowledge graph embeddings, in: *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 2018, pp. 122–131.
- [43] S. Choudhary, T. Luthra, A. Mittal and R. Singh, A survey of knowledge graph embedding and their applications, *arXiv preprint arXiv:2107.07842* (2021).
- [44] F. Lu, P. Cong and X. Huang, Utilizing textual information in knowledge graph embedding: A survey of methods and applications, *IEEE Access* 8 (2020), 92072–92088.
- [45] H.N. Tran and A. Takasu, Analyzing knowledge graph embedding methods from a multi-embedding interaction perspective, *arXiv preprint arXiv:1903.11406* (2019).
- [46] Q. Zhang, Z. Sun, W. Hu, M. Chen, L. Guo and Y. Qu, Multi-view Knowledge Graph Embedding for Entity Alignment, in: *International Joint Conference on Artificial Intelligence*, 2019.

- [47] T. Trouillon, J. Welbl, S. Riedel, É. Gaussier and G. Bouchard, Complex embeddings for simple link prediction, in: *International conference on machine learning*, PMLR, 2016, pp. 2071–2080.
- [48] X. Luo, Z. Sun and W. Hu, μ KG: A Library for Multi-source Knowledge Graph Embeddings and Applications, in: *International Semantic Web Conference*, Springer, 2022, pp. 610–627.
- [49] J. Cao, J. Fang, Z. Meng and S. Liang, Knowledge Graph Embedding: A Survey from the Perspective of Representation Spaces, *ACM Comput. Surv.* **56**(6) (2024), 1–42.
- [50] X. Zhao, W. Zeng, J. Tang, W. Wang and F.M. Suchanek, An experimental study of state-of-the-art entity alignment approaches, *IEEE Transactions on Knowledge and Data Engineering* **34**(6) (2020), 2610–2625.
- [51] X. Zhao, Y. Jia, A. Li, R. Jiang and Y. Song, Multi-source knowledge fusion: a survey, *World Wide Web* **23** (2020), 2567–2592.
- [52] D. Vrandečić and M. Krötzsch, Wikidata: a free collaborative knowledgebase, *Communications of the ACM* **57**(10) (2014), 78–85.
- [53] K. Bollacker, C. Evans, P. Paritosh, T. Sturge and 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.
- [54] Z. Sun, J. Huang, M. Chen, L. Guo and Y. Qu, TransEdge: Translating Relation-Contextualized Embeddings for Knowledge Graphs, in: *The Semantic Web – ISWC 2019*, C. Ghidini, O. Hartig, M. Maleshkova, V. Svátek, I. Cruz, A. Hogan, J. Song, M. Lefrançois and F. Gandon, eds, Springer International Publishing, Cham, 2019, pp. 612–629.
- [55] Z. Wang, Q. Lv, X. Lan and Y. Zhang, Cross-lingual Knowledge Graph Alignment via Graph Convolutional Networks, in: *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, E. Riloff, D. Chiang, J. Hockenmaier and J. Tsujii, eds, Association for Computational Linguistics, Brussels, Belgium, 2018, pp. 349–357.
- [56] C. Wang, Z. Huang, Y. Wan, J. Wei, J. Zhao and P. Wang, FuAlign: Cross-lingual entity alignment via multi-view representation learning of fused knowledge graphs, *Information Fusion* **89** (2023), 41–52.
- [57] N. Fanourakis, V. Efthymiou, V. Christophides, D. Kotzinos, E. Pitoura and K. Stefanidis, Structural Bias in Knowledge Graphs for the Entity Alignment Task, in: *The Semantic Web: 20th International Conference, ESWC 2023, Hersonissos, Crete, Greece, May 28–June 1, 2023, Proceedings*, Springer-Verlag, Berlin, Heidelberg, 2023, pp. 72–90.
- [58] M. Chen, Y. Tian, M. Yang and C. Zaniolo, Multilingual knowledge graph embeddings for cross-lingual knowledge alignment, in: *Proceedings of the 26th International Joint Conference on Artificial Intelligence, IJCAI’17*, AAAI Press, 2017, pp. 1511–1517.
- [59] H. Zhu, R. Xie, Z. Liu and M. Sun, Iterative Entity Alignment via Joint Knowledge Embeddings., in: *IJCAI*, Vol. 17, 2017, pp. 4258–4264.
- [60] A. Bordes, N. Usunier, A. Garcia-Duran, J. Weston and O. Yakhnenko, Translating embeddings for modeling multi-relational data, *Advances in neural information processing systems* **26** (2013).
- [61] Z. Sun, W. Hu, Q. Zhang and Y. Qu, Bootstrapping entity alignment with knowledge graph embedding., in: *IJCAI*, Vol. 18, 2018.
- [62] Z. Wang, Q. Lv, X. Lan and Y. Zhang, Cross-lingual knowledge graph alignment via graph convolutional networks, in: *Proceedings of the 2018 conference on empirical methods in natural language processing*, 2018, pp. 349–357.
- [63] X. Mao, W. Wang, H. Xu, Y. Wu and M. Lan, Relational reflection entity alignment, in: *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*, 2020, pp. 1095–1104.
- [64] K. Xu, L. Wang, M. Yu, Y. Feng, Y. Song, Z. Wang and D. Yu, Cross-lingual Knowledge Graph Alignment via Graph Matching Neural Network, in: *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, A. Korhonen, D. Traum and L. Màrquez, eds, Association for Computational Linguistics, 2019, pp. 3156–3161.
- [65] J. Zhou, G. Cui, S. Hu, Z. Zhang, C. Yang, Z. Liu, L. Wang, C. Li and M. Sun, Graph neural networks: A review of methods and applications, *AI open* **1** (2020), 57–81.
- [66] F. Scarselli, M. Gori, A.C. Tsoi, M. Hagenbuchner and G. Monfardini, The graph neural network model, *IEEE transactions on neural networks* **20**(1) (2008), 61–80.
- [67] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang and P.S. Yu, A Comprehensive Survey on Graph Neural Networks, *IEEE Transactions on Neural Networks and Learning Systems* **32**(1) (2021), 4–24.
- [68] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A.N. Gomez, L.u. Kaiser and I. Polosukhin, Attention is All you Need, in: *Advances in Neural Information Processing Systems*, I. Guyon, U.V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan and R. Garnett, eds, NIPS’17, Vol. 30, Curran Associates, Inc., 2017, pp. 6000–6010.
- [69] T. Lin, Y. Wang, X. Liu and X. Qiu, A survey of transformers, *AI Open* (2022).
- [70] Y. Li, X. Liang, Z. Hu, Y. Chen and E.P. Xing, Graph transformer (2018).
- [71] D.Q. Nguyen, T.D. Nguyen and D. Phung, Universal graph transformer self-attention networks, in: *Companion Proceedings of the Web Conference 2022*, 2022, pp. 193–196.
- [72] J. Zhang, H. Zhang, C. Xia and L. Sun, Graph-bert: Only attention is needed for learning graph representations, *arXiv preprint arXiv:2001.05140* (2020).
- [73] L. Müller, M. Galkin, C. Morris and L. Rampásek, Attending to graph transformers, *arXiv:2302.04181* (2023).
- [74] V.P. Dwivedi and X. Bresson, A Generalization of Transformer Networks to Graphs, *ArXiv abs/2012.09699* (2020).
- [75] W. Cai, W. Ma, J. Zhan and Y. Jiang, Entity Alignment with Reliable Path Reasoning and Relation-aware Heterogeneous Graph Transformer, in: *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, 2022, pp. 1930–1937.
- [76] B.D. Trisedya, F.D. Salim, J. Chan, D. Spina, F. Scholer and M. Sanderson, i-Align: an interpretable knowledge graph alignment model, *Data Mining and Knowledge Discovery* (2023), 1–23.
- [77] M.S. Hussain, M.J. Zaki and D. Subramanian, Edge-augmented graph transformers: Global self-attention is enough for graphs, *arXiv preprint arXiv:2108.03348* (2021).
- [78] D. Chen, L. O’Bray and K. Borgwardt, Structure-aware transformer for graph representation learning, in: *International Conference on Machine Learning*, PMLR, 2022, pp. 3469–3489.

- [79] E. Min, R. Chen, Y. Bian, T. Xu, K. Zhao, W. Huang, P. Zhao, J. Huang, S. Ananiadou and Y. Rong, Transformer for graphs: An overview from architecture perspective, *arXiv preprint arXiv:2202.08455* (2022).
- [80] X. Tang, J. Zhang, B. Chen, Y. Yang, H. Chen and C. Li, BERT-INT: A BERT-based Interaction Model For Knowledge Graph Alignment, in: *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI-20*, 2020, pp. 3174–3180.
- [81] K. Yang, S. Liu, J. Zhao, Y. Wang and B. Xie, COTSAE: co-training of structure and attribute embeddings for entity alignment, in: *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 34, 2020, pp. 3025–3032.
- [82] W. Zeng, X. Zhao, W. Wang, J. Tang and Z. Tan, Degree-aware alignment for entities in tail, in: *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2020, pp. 811–820.
- [83] Y. Wu, X. Liu, Y. Feng, Z. Wang, R. Yan and D. Zhao, Relation-Aware Entity Alignment for Heterogeneous Knowledge Graphs, in: *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19*, 2019, pp. 5278–5284.
- [84] T.N. Kipf and M. Welling, Semi-Supervised Classification with Graph Convolutional Networks, in: *International Conference on Learning Representations*, 2017. <https://openreview.net/forum?id=SJU4ayYgl>.
- [85] E. Jiménez-Ruiz and B. Cuenca Grau, Logmap: Logic-based and scalable ontology matching, in: *The Semantic Web–ISWC 2011: 10th International Semantic Web Conference, Bonn, Germany, October 23-27, 2011, Proceedings, Part I 10*, Springer, 2011, pp. 273–288.
- [86] M. Berrendorf, E. Faerman, V. Melnychuk, V. Tresp and T. Seidl, Knowledge graph entity alignment with graph convolutional networks: lessons learned, in: *Advances in Information Retrieval: 42nd European Conference on IR Research, ECIR 2020, Lisbon, Portugal, April 14–17, 2020, Proceedings, Part II 42*, Springer, 2020, pp. 3–11.
- [87] I. Fundulaki and A.-C.N. Ngomo, Instance matching benchmark for spatial data: a challenge proposal to OAEI., in: *OM@ ISWC*, 2016, pp. 233–234.
- [88] P. Larmande and K. Todorov, AgroLD: A knowledge graph for the plant sciences, in: *The Semantic Web–ISWC 2021: 20th International Semantic Web Conference, ISWC 2021, Virtual Event, October 24–28, 2021, Proceedings 20*, Springer, 2021, pp. 496–510.
- [89] P. Lisena, M. Achichi, P. Choffé, C. Cecconi, K. Todorov, B. Jacquemin and R. Troncy, Improving (re-) usability of musical datasets: An overview of the doremus project, *Bibliothek Forschung und Praxis* **42**(2) (2018), 194–205.
- [90] A. Venkatesan, G. Tagny Ngompé, N.E. Hassouni, I. Chentli, V. Guignon, C. Jonquet, M. Ruiz and P. Larmande, Agronomic Linked Data (AgroLD): A knowledge-based system to enable integrative biology in agronomy, *PLoS One* **13**(11) (2018), e0198270.
- [91] D.M. Endres and J.E. Schindelin, A new metric for probability distributions, *IEEE Transactions on Information theory* **49**(7) (2003), 1858–1860.
- [92] S. Kullback and R.A. Leibler, On information and sufficiency, *The annals of mathematical statistics* **22**(1) (1951), 79–86.
- [93] L. Yujian and L. Bo, A Normalized Levenshtein Distance Metric, *IEEE Transactions on Pattern Analysis and Machine Intelligence* **29**(6) (2007), 1091–1095.
- [94] K. Heylen, Y. Peirsman, D. Geeraerts and D. Speelman, Modelling Word Similarity: an Evaluation of Automatic Synonymy Extraction Algorithms., in: *LREC*, Vol. 8, 2008, pp. 3243–3249.
- [95] A. Elekes, M. Schaefer and K. Boehm, On the Various Semantics of Similarity in Word Embedding Models. In 2017 ACM/IEEE Joint Conference on Digital Libraries (JCDL), 1–10, 2017.
- [96] J. Shen, C. Wang, L. Gong and D. Song, Joint Language Semantic and Structure Embedding for Knowledge Graph Completion, in: *Proceedings of the 29th International Conference on Computational Linguistics*, International Committee on Computational Linguistics, Gyeongju, Republic of Korea, 2022, pp. 1965–1978.
- [97] Z. Zhang, X. Liu, Y. Zhang, Q. Su, X. Sun and B. He, Pretrain-KGE: learning knowledge representation from pretrained language models, in: *Findings of the Association for Computational Linguistics: EMNLP 2020*, 2020, pp. 259–266.
- [98] B. Wang, A. Wang, F. Chen, Y. Wang and C.-C.J. Kuo, Evaluating word embedding models: Methods and experimental results, *APSIPA transactions on signal and information processing* **8** (2019), e19.
- [99] G. Zhu and C.A. Iglesias, Computing semantic similarity of concepts in knowledge graphs, *IEEE Transactions on Knowledge and Data Engineering* **29**(1) (2016), 72–85.
- [100] I. Traverso, M.-E. Vidal, B. Kämpgen and Y. Sure-Vetter, GADES: a graph-based semantic similarity measure, in: *Proceedings of the 12th International Conference on Semantic Systems*, 2016, pp. 101–104.
- [101] L. Van der Maaten and G. Hinton, Visualizing data using t-SNE., *Journal of machine learning research* **9**(11) (2008).
- [102] N. Polyzotis, S. Roy, S.E. Whang and M. Zinkevich, Data lifecycle challenges in production machine learning: a survey, *ACM SIGMOD Record* **47**(2) (2018), 17–28.
- [103] V. Gudivada, A. Apon and J. Ding, Data quality considerations for big data and machine learning: Going beyond data cleaning and transformations, *International Journal on Advances in Software* **10**(1) (2017), 1–20.
- [104] E.M. Hamedani and M. Kaedi, Recommending the long tail items through personalized diversification, *Knowledge-Based Systems* **164** (2019), 348–357.
- [105] L. Shi, Trading-off among accuracy, similarity, diversity, and long-tail: a graph-based recommendation approach, in: *Proceedings of the 7th ACM Conference on Recommender Systems*, 2013, pp. 57–64.
- [106] Z. Liu, T.-K. Nguyen and Y. Fang, Tail-gnn: Tail-node graph neural networks, in: *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery & Data Mining*, 2021, pp. 1109–1119.
- [107] L. Liang, Z. Xu, Z. Song, I. King, Y. Qi and J. Ye, Tackling Long-tailed Distribution Issue in Graph Neural Networks via Normalization, *IEEE Transactions on Knowledge and Data Engineering* (2023).
- [108] Y. Wang, D. Wang, H. Liu, B. Hu, Y. Yan, Q. Zhang and Z. Zhang, Optimizing Long-tailed Link Prediction in Graph Neural Networks through Structure Representation Enhancement, in: *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, KDD '24, Association for Computing Machinery, New York, NY, USA, 2024, pp. 3222–3232.

- [109] E. Giamphy, J.-L. Guillaume, A. Doucet and K. Sanchis, A survey on bipartite graphs embedding, *Social Network Analysis and Mining* **13** (2023). doi:10.1007/s13278-023-01058-z.
- [110] T. Blevins, T. Limisiewicz, S. Gururangan, M. Li, H. Gonen, N.A. Smith and L. Zettlemoyer, Breaking the curse of multilinguality with cross-lingual expert language models, *arXiv preprint arXiv:2401.10440* (2024).
- [111] J. Pfeiffer, N. Goyal, X. Lin, X. Li, J. Cross, S. Riedel and M. Artetxe, Lifting the Curse of Multilinguality by Pre-training Modular Transformers, in: *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, M. Carpuat, M.-C. de Marneffe and I.V. Meza Ruiz, eds, Association for Computational Linguistics, Seattle, United States, 2022, pp. 3479–3495.
- [112] I. Gulrajani and D. Lopez-Paz, In Search of Lost Domain Generalization, in: *International Conference on Learning Representations*, 2021. <https://openreview.net/forum?id=lQdXeXDoWtI>.
- [113] S. Fan, X. Wang, C. Shi, P. Cui and B. Wang, Generalizing Graph Neural Networks on Out-of-Distribution Graphs, *IEEE Trans. Pattern Anal. Mach. Intell.* **46**(1) (2024), 322–337. <https://doi.org/10.1109/TPAMI.2023.3321097>.
- [114] R. Roelofs, *Measuring Generalization and overfitting in Machine learning*, University of California, Berkeley, 2019.
- [115] B. Yang, W. Yih, X. He, J. Gao and L. Deng, Embedding Entities and Relations for Learning and Inference in Knowledge Bases, 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.
- [116] Y. Lin, Z. Liu, M. Sun, Y. Liu and X. Zhu, Learning entity and relation embeddings for knowledge graph completion, in: *Proceedings of the AAAI conference on artificial intelligence*, Vol. 29, 2015.
- [117] E. Jiménez-Ruiz, T. Saveta, O. Zamazal, S. Hertling, M. Roder, I. Fundulaki, A.N. Ngomo, M. Sherif, A. Annane, Z. Bellahsene et al., Introducing the HOBBIT platform into the ontology alignment evaluation campaign, in: *13th International Workshop on Ontology Matching (OM)*, Vol. 2288, 2018, pp. 49–60.
- [118] A. Rossi, D. Barbosa, D. Firmani, A. Matinata and P. Merialdo, Knowledge graph embedding for link prediction: A comparative analysis, *ACM Transactions on Knowledge Discovery from Data (TKDD)* **15**(2) (2021), 1–49.
- [119] T. Dettmers, P. Minervini, P. Stenatorp and S. Riedel, Convolutional 2d knowledge graph embeddings, in: *Proceedings of the AAAI conference on artificial intelligence*, Vol. 32, 2018.
- [120] B. Fatemi, S. Ravanbakhsh and D. Poole, Improved knowledge graph embedding using background taxonomic information, in: *Proceedings of the AAAI conference on artificial intelligence*, Vol. 33, 2019, pp. 3526–3533.
- [121] H. Xiao, M. Huang and X. Zhu, TransG : A Generative Model for Knowledge Graph Embedding, in: *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Association for Computational Linguistics, 2016, pp. 2316–2325.
- [122] G. Niu, B. Li, Y. Zhang, S. Pu and J. Li, AutoETER: Automated entity type representation for knowledge graph embedding, *arXiv preprint arXiv:2009.12030* (2020).
- [123] K. Xu, L. Wang, M. Yu, Y. Feng, Y. Song, Z. Wang and D. Yu, Cross-lingual knowledge graph alignment via graph matching neural network, *arXiv preprint arXiv:1905.11605* (2019).
- [124] T. Qian, Y. Liang and Q. Li, Solving cold start problem in recommendation with attribute graph neural networks, *arXiv preprint arXiv:1912.12398* (2019).
- [125] B. Shi, H. Wang, Y. Li and S. Deng, RelaGraph: Improving embedding on small-scale sparse knowledge graphs by neighborhood relations, *Information Processing & Management* **60**(5) (2023), 103447.