

A Systematic Literature Review and Classification of Approaches for Keyword Search Over Graph-Shaped Data

Semantic Web
Vol. 16(5) 1–34
© The Author(s) 2025
Article reuse guidelines:
sagepub.com/journals-permissions
DOI: 10.1177/22104968251377243
journals.sagepub.com/home/swj

Leila Feddoul¹ , Frank Löffler^{1,2} and Sirko Schindler³

Abstract

Knowledge graphs provide machine-interpretable data that allow automatic data understanding and deduction of new facts. However, machines are not the only consumers of such semantic data. Human users could also benefit from graph-structured data by browsing and exploring it to detect interesting associations and draw conclusions. To achieve that, methods that allow for search over knowledge graphs are highly sought after. Keyword search is an intuitive and common way to retrieve relevant data (e.g., documents) and can also be leveraged to search over knowledge graphs. In this survey paper, we derive the typical architecture of a system for keyword search over graph-shaped data, we formally define the problem, we highlight related challenges, and we compare with existing relevant surveys to identify the gaps. We conduct a comprehensive review of studies dealing with the topic of keyword search over graph-shaped data (e.g., knowledge graphs) following a systematic method. Based on that, we derive and define different aspects for classifying existing works. We also give an overview of how those systems are evaluated and highlight possible future research directions.

Keywords

keyword search, knowledge graph, survey

Received: July 21, 2023; accepted: May 21, 2025

Editor: Mehwish Alam, Associate Professor of Language, Knowledge, and Artificial Intelligence, Télécom Paris, Institut Polytechnique de Paris, France

Solicited reviews: Stefano De Giorgis, Adjunct professor, Department of Modern Languages, Literatures, and Cultures, Università di Bologna, Bologna; Four anonymous reviews

1 Introduction

Knowledge graphs (KGs) have shown their importance in many application areas (e.g., search, question answering, or recommendations). Thus, the data represented as KGs is continuously increasing—a fact also evidenced by the continuous growth of KGs such as Wikidata (Cantallos et al., 2019). While their semantic data representation can be directly leveraged by machines to accomplish a specific task (e.g., inference), end users and domain experts should also be able to access and explore the KGs. This is a first step in allowing us to go beyond simple information lookup and enable the discovery of new facts and correlations given by the graph structure. We can distinguish between two scenarios: (1)

¹Heinz Nixdorf Chair for Distributed Information Systems, Friedrich Schiller University Jena, Jena, Germany

²Competence Center for Digital Research, Michael Stifel Center, Jena, Germany

³Institute of Data Science, German Aerospace Center (DLR), Jena, Germany

Corresponding Author:

Leila Feddoul, Heinz Nixdorf Chair for Distributed Information Systems, Friedrich Schiller University Jena, Jena, Germany.

Email: leila.feddoul@uni-jena.de



Creative Commons CC BY: This article is distributed under the terms of the Creative Commons Attribution 4.0 License

(<https://creativecommons.org/licenses/by/4.0/>) which permits any use, reproduction and distribution of the work without further

permission provided the original work is attributed as specified on the SAGE and Open Access page

(<https://us.sagepub.com/en-us/nam/open-access-at-sage>).

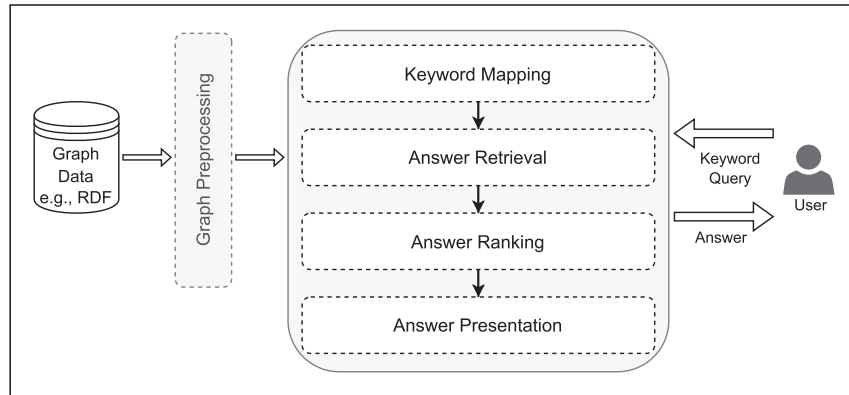


Figure 1. Typical Generic Components of a System for Keyword Search Over Graph-Shaped Data.

KG querying and search includes the usage of structured queries, interactive query builders, keyword search, or question answering. (2) *KG exploration* comprises visualization and browsing, faceted search, graph analytics, and summarization.

While structured queries allow the formulation of the information need in a precise way, the most familiar and intuitive method to search over data collections is keyword search (Agarwal et al., 2010). This reduces the time to learn new technologies and does not require an intricate knowledge of the underlying domain and the data schema. Therefore, keyword search is also used and adapted to work over KGs.

The typical generic components of a system for keyword search over graph-shaped data are depicted in Figure 1:

- *Graph preprocessing*: refers to the operations applied to prepare and index the data graph for efficient query processing (e.g., keyword-node index) and answer retrieval (e.g., shortest paths index or summary graph).
- *Keyword mapping*: this step consists of mapping words or a group of words in the query to a graph’s nodes and edges.
- *Answer retrieval*: this step consists of finding results relevant to the query (e.g., using graph traversal). The retrieval unit varies depending on the used technique (e.g., tree or entity).
- *Answer ranking*: this step ranks results in the order of relevance and is also dependent on the specific results nature.
- *Answer presentation*: refers to the way the results are provided to the user and the different features that are provided to support the search experience.

Implementing these common system components poses the following main challenges:

- *Finding accurate mappings between keyword terms and graph elements*: this includes a correct interpretation and disambiguation of the query keywords. This is a crucial step since it directly affects the efficiency and effectiveness of the subsequent answer retrieval phase.
- *Scaling to large graphs*: this includes the development of techniques to accelerate the retrieval of answers (e.g., suitable indexing solutions or greedy algorithms).
- *Ranking retrieved answers*: this includes the development of ranking functions that leverage different aspects depending on the type of the answer.

Those challenges lead to increasing research efforts in this area. To the best of our knowledge, there are four surveys that give an overview of the state of research on keyword search over graph-shaped data, focusing on different aspects. However, they show some shortcomings (e.g., outdated), which are discussed in Section 2. Therefore, we close this gap by performing a systematic review with the following contributions:

- We perform a **systematic literature review** to identify and classify relevant research on keyword search over graph-shaped data.
- We propose a **taxonomy** for classifying existing works based on different aspects: *search space*, *answer type*, *answer ranking*, and *answer scoring*.
- We provide a comprehensive review of **evaluation methodologies** used in the fields.
- We identify remaining research gaps and propose potential **future research directions**.

Table 1. Summary Table Comparing Existing Surveys With This One.

Survey	Date	Data	Classification	Evaluation Methods
Yu et al. (2010)	2010	Relational database	Answer type and scoring	✗
Wang and Aggarwal (2010)	2010	Graph	Data type	✗
Ghanbarpour and Naderi (2019)	2019	Graph	Answer ranking	✗
Yang et al. (2021)	2021	Graph	Answer type and scoring	✗
This survey	2025	Graph	Four distinct aspects	✓

With this survey, we aim at enabling academic researchers who deal with the same topic to understand the state-of-the-art and use it to identify research gaps, providing a comprehensive introduction and summary for newcomers, and helping tool developers and engineers by selecting systems that are suitable for integration into their applications.

The remainder of this survey is organized as follows: In Section 2, we position our work with respect to existing surveys. In Section 3, we define the task of keyword search over graph-shaped data as used within this survey. In Section 4, we describe the methodology for the systematic literature review, and in Section 5, we outline the different aspects used for the classification of relevant works. In Section 6, we present the different methods used to evaluate relevant systems, and in Section 7, we give a summary overview of them. In Section 8, we cluster relevant works and provide a detailed description for each one. Section 9 concludes the review and discusses future research directions.

2 Related Work

We identified four surveys focusing on keyword search over graph-shaped data, as shown in Table 1. However, two of them (Wang & Aggarwal, 2010; Yu et al., 2010) are quite outdated (2010), and one (Ghanbarpour & Naderi, 2019), dating from 2019, reports only on one aspect (answer ranking). One survey (Yu et al., 2010) considers only relational databases. However, this is still relevant as well since relational databases can also be modeled as graphs by considering the tuples as nodes and foreign key relations as edges (Feddoul, 2020). The most recent survey (Yang et al., 2021) classifies existing works based on the type of returned answer as already proposed in an earlier survey (Yu et al., 2010). While the latter also includes an overview of methods depending on the schema of relational databases, (Yang et al., 2021) concentrates on schema-free approaches. Both works use a classification that mixes two aspects in one concept, namely the answer type and the way of scoring the final answer. In Wang and Aggarwal (2010), works are classified in a rather generic manner based on the underlying data type, namely *Keyword Search on XML Data* and *Keyword Search over Relational Databases*, which are both considered schema-dependent, whereas *Keyword Search on Graphs* refers to schema-free approaches. In Ghanbarpour and Naderi (2019), the authors concentrate on the answer ranking aspect. The proposed factors for the answer ranking slightly overlap with the ones proposed in Section 5, but we adjust and add other ranking types while preserving the distinction between the important aspects. Similar to Yang et al. (2021), they also mix the two factors of answer scoring and ranking. In addition, we also notice the inconsistent use of category names (e.g., “centralized ranking function” to refer to the distinct-root semantics as in Yang et al., 2021).

In addition to the previously mentioned shortcomings, none of the existing surveys report on methodologies used for evaluating the systems, including, for example, the used ground truth, datasets, common benchmarks, and metrics. Furthermore, only one work considered the ranking (Ghanbarpour & Naderi, 2019). Overall, for a certain consistency and simplicity of definitions, it is beneficial to consider each classification aspect separately, as suggested and extended in our proposed categorization (cf. Section 5). In addition, our literature review demonstrates that this area of research is constantly evolving, resulting in new publications that were not included in previous surveys and more works considering KGs.

3 Preliminaries

We consider the task of using keywords to search over graph-shaped data where, given a graph G and a keyword query Q , the goal is to find proper answer(s) $A_1, A_2, \dots, A_n$ that satisfy the information need expressed in Q . The response might either be an unranked list of relevant answers, a ranked list by order of relevancy, or just the single-best answer. The relevant inputs and outputs are formally described as follows:

Graph Model. We consider the graph $G = (N, E, L_N, L_E)$, where N is a finite nonempty set of nodes, $E \subseteq N \times N$ is a finite set of edges connecting two nodes, L_N and L_E are finite sets of textual information (e.g., labels) describing nodes and edges, respectively. Depending on the situation, the graph can be directed or undirected, and its nodes and/or edges may have weights assigned.

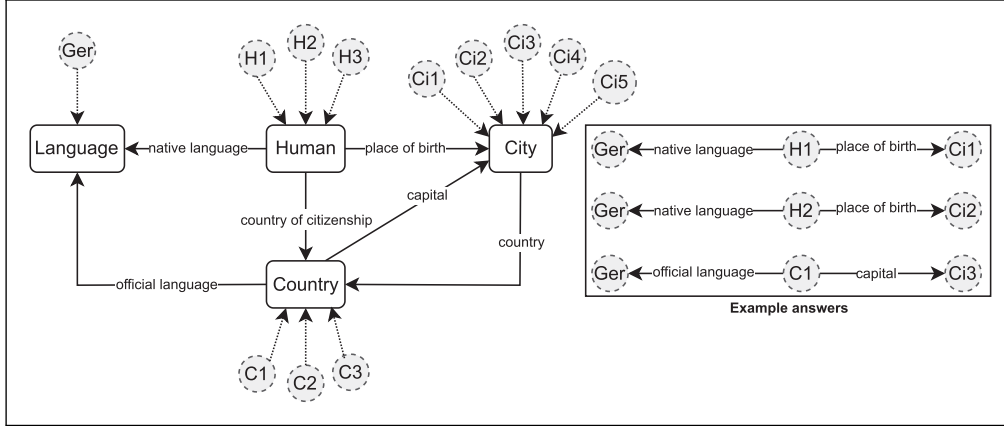


Figure 2. Example KG With Some Possible Answers to the Query $Q=\{german, city\}$. Boxes Represent the Concepts/Classes on the Schema Level, Dotted Circles are Individuals (Instance Level), and Dotted Arrows State That an Individual is an Instance of a Class. Ger is the Abbreviation of the Label German.

Note. KG = knowledge graph.

Keyword Query and Node. A keyword query Q consists of a nonempty list of terms $Q = k_1, k_2, \dots, k_n$. The set of nodes K_i with a textual content matching k_i for $i = 1, \dots, n$ is called *keyword nodes*. A keyword node can be associated with multiple keywords, and the same keyword can be mapped to various nodes. This matching can be extended to include edges as well, for example, by also considering them as nodes in the graph (Cheng et al., 2020; Tran et al., 2009). The mapping between keywords and nodes can be provided externally or be included as an additional task that is either a mere string matching or an entity linking (Balog, 2018), where multiple keywords could correspond to a single entity in the case of KGs.

Query Answer. A relevant answer A_i to the query Q is a connected subgraph $S = (N_S, E_S)$ (for every pair of nodes $n_1, n_2 \in N_S$ there is a path that connects them Valiente, 2021) that covers all the keywords (the subgraph contains for each keyword at least one corresponding *keyword node*). A relaxed variant will also allow for answers covering only a subset of keywords. This is a general description, without imposing any further restrictions on the characteristics of the answer. In Section 5, we give a detailed overview of possible answer types mentioned in the literature.

Figure 2 shows an example graph¹ with nodes and edges on both schema and instance levels. For the sake of brevity, we illustrate only the relations between concepts/classes and link them to their instances. This implies that each relation on the schema level has also its counterpart on the instance level (e.g., *H1 place of birth Ci1*). All nodes and edges have labels. The graph is directed, and the distance between a pair of nodes is given by the number of edges on the path between them (e.g., distance between *H1* and *Ci1* is one). The considered keyword query consists of two terms *german* and *city*. Each term has its corresponding keyword node either on the schema (*city*) or instance level (*german*). Three example answers that connect the keywords are provided. For the sake of simplicity, we omit the instance-class relations in the answer representation (e.g., *Ci1-City*). In the example, *Ci1*, *Ci2*, and *Ci3* are considered valid answer entities for the given query.

4 Survey Methodology

To collect relevant works in the area of keyword search over graph-shaped data, we performed a systematic literature review inspired by the guidelines provided by Brereton et al. (2007) and Kitchenham and Charters (2007). This allowed us to summarize and classify existing related work as well as determine remaining research gaps. The general process consists of three phases:

1. Plan review: includes the identification of the need for the review and the development of the review protocol, which describes the steps that will be followed while conducting the review.
2. Conduct review: includes the concrete steps of performing the review, starting from the identification of relevant publications.
3. Document review: includes writing a report about the review and publishing it.

The *identification of the need for the review* was already addressed in Section 2 and the *review protocol* will be clarified while describing the second phase. Furthermore, the last phase is addressed by this survey. For those reasons, we only describe the *conduct review* phase in more detail in the following:

A. Identify Relevant Research. In this step, potentially relevant candidate papers that match specific predefined keywords are collected. In practice, we first searched the DBLP² bibliographic database using the following queries³: *keyword search graph*, *keyword search RDF*, *keyword quer RDF*, *keyword quer graph*, *quer generation graph*, *quer generation RDF*, *keyword exploration graph*, and *keyword exploration RDF*. Then, duplicates are automatically detected and removed after saving the original search results for each query. Publication collection and deduplication were performed using the *web search* functionality of the reference management software JabRef 5.7⁴ and resulted in 206 candidates.

B. Select Primary Studies. In this step, a primary selection of relevant papers is performed based on predefined inclusion and exclusion criteria and a first reading of the title and abstract. In our case, this selection resulted in 68 papers. The applied criteria were as follows:

- *Inclusion criteria.*
 - * Publication year starting from 2013 (last 10 years).
 - * Open/institutional access.
 - * English language.
 - * Paper type: conference, journal, or workshop paper.
 - * Content fits to the topic of *keyword search over graph-shaped data*⁵.
- *Exclusion criteria.*
 - * Paper type: PhD symposium, Demo/poster, or extended abstract.
 - * Not peer reviewed (e.g., only published on an open access archive).
 - * Newer paper of the same approach or system exists (keep the newest).

C. Assess Study Quality and Extract/Synthesize Data. In this step, a more detailed reading of the papers takes place to allow their classification based on the criteria defined in Section 5. During this process, we may also decide against including some specific type of works as part of the works overview (cf. Section 8)⁶. The data extraction and synthesis are also done in parallel for each paper by: (1) collecting information (e.g., classification criteria and evaluation method) which allows further analysis, and (2) drafting a short descriptive summary.

At the end, we selected and classified 35 works.

We also included a subset of works that were published before 2013. This includes popular first works dealing with the keyword search problem over relational data (Bhalotia et al., 2002; He et al., 2007; Kacholia et al., 2005), in addition to one of the very first works operating over KGs (Tran et al., 2009), and three other papers (Dalvi et al., 2008; Golenberg et al., 2008; Kimelfeld & Sagiv, 2008) that also seemed relevant and cited one of the pioneer works in this field (Bhalotia et al., 2002).

5 Taxonomy Overview

The literature review allowed us to derive four important aspects for categorizing related works as shown in Figure 3, namely the *search space*, *answer type*, *answer ranking*, and *answer scoring*.

Search Space. Represents the nature of the underlying data searched to find relevant answers. We distinguish between four types:

- *Graph-based.* Search for answers over the whole graph.
- *Schema-based.* Operate on a provided data schema to find answers.
- *Summary-based.* Operate on a summarized version (e.g., replace a collection of instances with their corresponding class) of the data graph without relying on a predefined schema.
- *Virtual document-based.* Build documents from the textual content of the graph elements and retrieve relevant ones using information retrieval techniques.

Answer Type. Represents the nature of the final answer to be returned by the system. According to the analyzed works, we can distinguish between the following answer types:

Search space	Answer type	Answer ranking	Answer scoring
Graph-based	Subgraph	Compactness-based	Distinct
Schema-based	Tree	Importance-based	Root-keyword
Summary-based	Structured query	Textual-based	Keyword-pair
Virtual document-based	Other (e.g., triples)	Diversity-based	
		Semantics-based	
		Profile-based	

Figure 3. Dimensions for Classifying Approaches for Keyword Search Over Graph-Shaped Data.

- *Subgraph*. The answer is a set of connected keyword nodes. A special case is the *r-clique*, where the answer is a subgraph connecting keywords and the shortest distance between any pair of keyword nodes is equal to or lower than r (Ghanbarpour et al., 2020).
- *Tree*. The answer is a graph that has no cycles (tree) and that connects all keyword nodes (Valiente, 2021). A tree can be either directed or undirected and further characterized using the following properties:
 - * Rooted directed: a rooted directed tree is a tree in which a specific node is distinguished and called *root*, and there exist directed paths (pointing away from the root) from the root to each keyword node (Bhalotia et al., 2002).
 - * Distinct rooted: a distinct rooted tree is a tree with a distinct root with respect to a collection of top- k answers (He et al., 2007).
 - * Minimal: the minimality is defined here with respect to the keywords and states that all the leaves of the tree are keyword nodes (in other words, there are no leaf nodes not containing keywords (Lu et al., 2019), or the answer has no other subtree containing the keywords).
- *Structured query*. The answer is a query equivalent to a previously found intermediate answer (subgraph or tree). When executed, the query returns relevant concrete results (e.g., triples or entities in case of RDF graphs).
- *Other types*. There are also specific cases, where relevant answers are triples or entities of an RDF graph. Those are directly retrieved without having an intermediate step where a subgraph/tree/query is formed. This is usually the case when a graph's textual content is indexed (e.g., triple index), and the answers are retrieved using traditional ranking functions as used in document retrieval. Another specific answer type, introduced by just a single work, is the *Key-core* (Zhang et al., 2022), which represents a subgraph that is formed by connecting two minimal rooted directed trees containing all keywords. Furthermore, for each tree, the distance between each keyword node and the root, and the sum of keyword-node distances are smaller than two specific thresholds.

Figure 4 depicts some example answers of the query $Q = \{german, city\}$ over the KG of Figure 2. The first three answers are distinct rooted trees with respect to the five examples. Moreover, all four first answers are rooted directed trees that are also minimal. The last answer contains a cycle and thus is not a tree but a subgraph. The provided SPARQL query is also a possible answer type, which is equivalent to the tree *Language—Human—City* on the schema level. It returns the entities *Ci1* and *Ci2* that represent the cities where the German native speaker humans *H1* and *H2* were born.

Answer Ranking. In the case where more than one relevant result is retrieved, answers are ranked based on different criteria. In general, we distinguish between six ranking categories, which are either individually applied or combined to rank answers:

- *Compactness-based*. Use the size (e.g., tree height or number of nodes) of the retrieved structure as a quality criterion.
- *Importance-based*. Use specific criteria (e.g., node in-degree) that may reflect the popularity and importance of the answer elements.

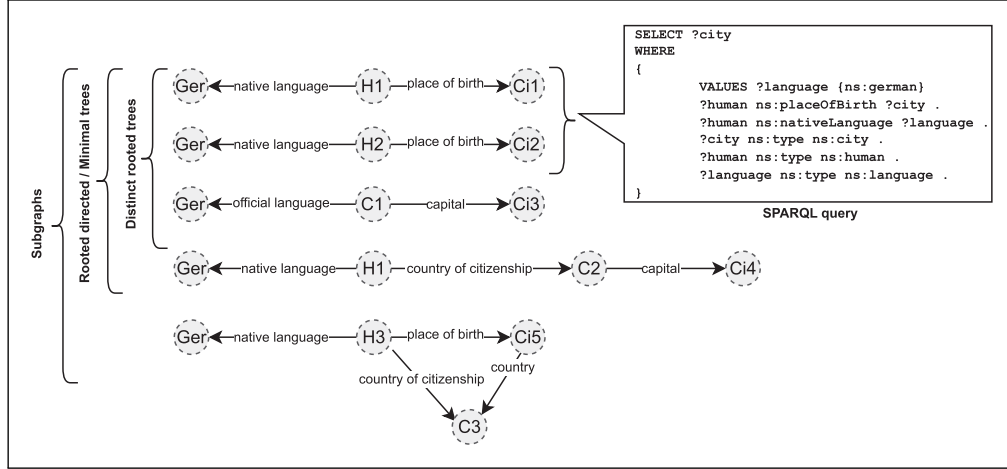


Figure 4. Different Example Answers of the Query $Q=\{\text{german}, \text{city}\}$ Over the KG of Figure 2 Together With Their Types. *ns:* Is an Example Namespace.

Note. KG = knowledge graph.

- *Textual-based.* Use matches between query keywords and the textual content of the answer as a ranking factor (e.g., TF/IDF; Leskovec et al., 2020).
- *Diversity-based.* Include criteria that aim at penalizing results showing a certain aspect of redundancy.
- *Semantics-based.* Use semantic characteristics of KGs to rank answers (e.g., semantic distance).
- *Profile-based.* Use user profiles⁷ to provide answers closer to the user intent.

Answer Scoring. This aspect applies to the cases where the final or intermediate answers are graph-shaped (e.g., tree, subgraph, or structured query). In general, a weight is assigned to each element in the answer (node or edge) based on one of the previously mentioned answer ranking categories. However, a score for the whole answer is needed to be able to rank them. Thus, the final score of a graph-shaped answer is usually calculated by aggregating node and edge weights (answer scoring). Since the usual aim is to find answers with minimum weights (or cost) (e.g., Steiner tree Dreyfus and Wagner 1971), the score of an answer is, in general, inversely proportional to its weight. Thus, answers with higher scores (lower weights) are considered more relevant and thus ranked higher. We distinguish between three *answer scoring* schemes that can be used to calculate total weights for both nodes and edges (graph elements):

- *Distinct scoring.* The total weight is calculated as the sum of individual weights: $\sum_{x \in S} \text{weight}(x)$, where x is a graph element and S is either the set of edges or the set of nodes. Here, each included graph element is counted a single time.
- *Root-keyword scoring.* The total weight is calculated as follows: $\sum_{\text{keyN} \in N_s} \sum_{x \in \text{path}(r, \text{keyN})} \text{weight}(x)$, where keyN is a keyword node, N_s is the set of answer's nodes, r is the root, x is a graph element, and $\text{path}(r, \text{keyN})$ is the set of the specific elements between the root and a specific keyword node. In summary, the answer weight is the sum of the graph element weights belonging to the path between the root⁸ and each keyword node. In this case, graph elements could be counted multiple times, for example, in the case of having keyword nodes that are not tree leaves. This considers the distribution of each keyword node with respect to the root.
- *Keyword-pair scoring.* The total weight of an answer is calculated as follows $\sum_{i=1}^{|Q|} \sum_{j=i+1}^{|Q|} \text{weight}(\text{keyN}_i, \text{keyN}_j)$, given a query Q . In other words, it is the sum of the weights between each pair of keyword nodes. It is worth mentioning that in specific cases, such as in Bryson et al. (2020), the weight used in the keyword-pair scoring is replaced with a value reflecting the connectivity between two nodes. Thus, it is interpreted as directly proportional to the total score.

In the same final score, two different scoring schemes can be combined, one for edges and the other for nodes. Furthermore, some of the works have restricted the nodes incorporated while calculating the edge weight and consider only specific ones (e.g., roots and leaves containing keywords).

During the literature review, we also came across several specific groups of works that deal with the following aspects: *parallel processing* (Guan et al., 2018; Hao et al., 2015; Liu et al., 2016; Virgilio & Maccioni, 2014; Yang et al., 2019),

probabilistic/uncertain RDF graphs (Lian et al., 2015), *temporal graphs* (Gkirtzou et al., 2015; Liu et al., 2018b), *distributed graphs* (Yuan et al., 2017), *incomplete graphs* (Hao et al., 2021), or *search results diversification*⁹ (Bikakis et al., 2013; Park, 2018; Wang et al., 2017b; Zhong et al., 2018). We considered them as special use cases, and thus decided to only shortly report them here without further analysis and refrain from including them in the overview tables (cf. Section 7).

6 System Evaluation Methods

In general, two aspects are evaluated: effectiveness and efficiency. While the latter deals with the execution time and memory usage, effectiveness judges the ability of the system to retrieve results that are relevant to the query and requires an underlying dataset that is searched through using keywords, a set of user queries, and corresponding results together with their relevance judgments. We also notice that some systems have used approach-specific aspects and metrics for the evaluation: #nodes explored/touched (Kacholia et al., 2005), index performance (He et al., 2007; Tran et al., 2009), cache misses (Dalvi et al., 2008), repetition rate (Zhang et al., 2014), answer compactness (diameter) (Kargar & An, 2015), or #intermediate results (Lu et al., 2019).

Table 2 gives an overview of the data¹⁰ and metrics used for evaluating the effectiveness of the surveyed systems (sorted in ascending chronological order). All the systems evaluate efficiency (e.g., execution time), except for Kharrat et al. (2016) and Sitthisarn (2014) that only assess effectiveness.

Furthermore, for one system (Kimelfeld & Sagiv, 2008), no evaluation was conducted, and the other works with no entries did not evaluate effectiveness.

Queries and datasets. We observe that the number of queries used for evaluation varies between 5 and 467. The number of queries is larger when they are randomly generated or an existing benchmark is used, while we notice only dozens of queries by manual creation, since it is a time and personnel-consuming task. Furthermore, the most used datasets are: DBLP, IMDb,¹¹ and DBpedia (Lehmann et al., 2015).

Ground truth. For each query, a list of potentially relevant results is needed. This is either manually created by result rating from users or by manually creating a corresponding structured query, or automatically generated (other system, structured query, or a tool for benchmark construction). In general, we observe that only 10 systems evaluated their approaches using existing test collections (Coffman Coffman & Weaver, 2014), DBpedia-Entity v1/v2 (Balog & Neumayer, 2013; Hasibi et al., 2017), and QALD¹² that provide both queries and relevant results.

Metrics. Most of the metrics used to measure the effectiveness are: *precision*, *recall*, *F1-measure*, *precision at k* ($P@k$), *discounted cumulative gain* (*DCG*) and its variants, *average precision* (*AP*), *mean AP* (*MAP*), and *mean reciprocal rank* (*MRR*). However, we notice that some single systems used other metrics that are shortly mentioned while summarizing the works in Section 8 (e.g., rank difference, the fraction of queries with results, Kendall-Tau, or the mean of the user scores). In the following, we briefly define the most popular metrics:

Precision, Recall, and F-measure. In general, the precision is defined as the ratio between the number of retrieved elements that are relevant to a specific query and the total number of retrieved elements (Cleverdon, 1967):

$$P = \frac{|\text{relevant} \cap \text{retrieved}|}{|\text{retrieved}|}. \quad (1)$$

Recall is defined as the ratio between the number of retrieved elements that are relevant and the total number of all existing relevant elements (Cleverdon, 1967):

$$R = \frac{|\text{relevant} \cap \text{retrieved}|}{|\text{relevant}|}. \quad (2)$$

These definitions may be slightly adapted by some works depending on the nature of the retrieved answer. For example, Dosso and Silvello (2020) define a triple-based variant of precision and recall to compare a generated and a ground truth subgraph, by defining recall as the ratio between the relevant triples and the total number of triples in the ground truth. Another adaptation is given by Navarro et al. (2021), where precision and recall are calculated for each generated SPARQL answer. Here, recall is a function of the intersection between the generated and ground truth query.

The F1-measure is calculated by combining both precision and recall as follows (Zhang & Zhang, 2009):

$$F1 = 2 \cdot \frac{P \cdot R}{P + R}. \quad (3)$$

Table 2. Overview of the Data and Metrics Used for Evaluating the Effectiveness of Surveyed Systems.

System	#Queries	Query Creation	Ground Truth	Dataset	Metrics
Bhalotia et al. (2002)	7	M	M	DBLP, IIT Bombay dataset	Rank difference
Kacholia et al. (2005)	5	M	G (SQL)	DBLP, IMDB, US patent database	Precision, recall
He et al. (2007)	—	—	—	—	—
Kimelfeld and Sagiv (2008)*	—	—	—	—	—
Dalvi et al. (2008)	12	M	M	DBLP, IMDB	Recall
Golenberg et al. (2008)	?	?	?	?	Ranking correlation
Tran et al. (2009)	39	M	M	DBLP, TAP	MRR
Zhong and Liu (2009)	12	M	?	DBLP	Minimum/avg. scores of top-10, 25, and 50
Sitthisarn (2014)	20	M	M	OntoLife, RDFResume, DBLP, FOAF, Researcher	MRR
Zhang et al. (2014)	10	R	M	DBLP, IMDb	DCG
Bae et al. (2014)	—	—	—	—	—
Kargar and An (2015)	19	M,R	M	DBLP, IMDb, Mondial	P@10, P@2, %ground truth answers
Wang et al. (2015)	20	M,R	M	DBpedia	1 match (top-10)
Park and Lim (2015)	30	M	M	DBLP, IMDb, Mondial	P@10, P@20
Mass and Sagiv (2016)	50	Coffman	Coffman	Mondial, Wikipedia, IMDb	MAP (top-1000)
Zheng et al. (2016)	10	M	?	DBLP	P@5, AP@5, MRR
Kharrat et al. (2016)	50	M	?	?	Fuzzy precision, fraction queries with result, F1
Han et al. (2017)	180	QALD-6, Free917	QALD-6, Free917	DBpedia, Freebase	Precision, recall, F1
Ouksili et al. (2017)	10	R	M	AIFB, the conference dataset	P@k, nDCG@k (k=10,20,5)
García et al. (2017)	50	Coffman	Coffman	Mondial, IMDb	—
Sinha et al. (2018)	20	M	M	Jamendo, DBpedia	Kendall-Tau, P@5, Rank-DCG
Rihany et al. (2018)	20	M	M	AIFB, DBpedia (subset)	P@10, P@5
Menendez et al. (2019)	75	Coffman, QALD-2	Coffman, QALD-2	IMDb, MusicBrainz	MAP
Lu et al. (2019)	—	—	—	—	—
Dosso and Silvello (2020)	177	M, DBpedia-Entity v1 (QALD-2)	M, DBpedia-Entity v1 (QALD-2)	LinkedMDB, IMDB, DBpedia, BSBM, LUBM	Recall, P@1, P@5, tb-DCG
Bryson et al. (2020)	200	R	—	DBLP, IntAct PPI	Expression level, h-index
Ghanbarpour et al. (2020)	10	Kargar and An (2015)	G (system)	IMDb, DBLP	Average answer weight
Kadilierakis et al. (2020)	467	DBpedia-Entity v2	DBpedia-Entity v2	DBpedia	nDCG@100, nDCG@10
Cheng et al. (2020)	438	DBpedia-Entity v2	DBpedia-Entity v2	Dbpedia	P@1
Jiang et al. (2021)	—	—	—	—	—
Izquierdo et al. (2021)	191	Coffman, (Dosso & Silvello, 2020)	Coffman, (Dosso & Silvello, 2020)	LUBM, BSBM, IMDb, Mondial, DBpedia	MAP, MRR, top-1
Navarro et al. (2021)	136	QALD-7	QALD-7	DBpedia	Precision, recall, F1
Zhang et al. (2021)	10	?	G (system)	DBLP, KMap	Relationship completeness
Shi et al. (2021)	281	DBpedia-Entity v2	DBpedia-Entity v2	DBpedia	Mean of user scores
Zhang et al. (2022)	5	R	M	DBpedia, DBLP	—

Note. ? = not mentioned in the paper; * = no evaluation; - = does not apply; M = manual; R = random; G = generated. Dataset refers to the underlying data queried using keywords.

The previous metrics are *set-based* measures. Next, we define *rank-based* metrics that also consider the order in which the documents are retrieved.

P@k. A good fraction of works used the precision at a specific rank position k as metric (Schütze et al., 2008). This is more practical since it does not evaluate all retrieved results, but only the top- k :

$$P@k = \frac{|\text{relevant} \cap \text{retrieved}@k|}{k}. \quad (4)$$

AP. The AP is given by the following formula (Liu, 2009), where n is the number of retrieved elements and rel_k is the relevance of the element at position k (1 or 0):

$$\text{AP} = \frac{1}{|\text{relevant}|} \sum_{k=1}^n \text{P@k} \cdot \text{rel}_k. \quad (5)$$

MAP. The MAP is calculated as the ratio between the sum of AP scores for each query and the total number of queries (Q) (Katerenchuk & Rosenberg, 2016):

$$\text{MAP} = \frac{\sum_{q=1}^Q \text{AP}(q)}{Q}. \quad (6)$$

DCG. This measure also takes into account the case where the relevance judgments are given by a graded scale (instead of a binary one) and considers both the relevance and position (Järvelin & Kekäläinen, 2002; Katerenchuk & Rosenberg, 2016). The DCG at a rank position n is defined as follows:

$$\text{DCG}_n = \sum_{i=1}^n \frac{\text{rel}_i}{\log_2(i+1)}, \quad (7)$$

where rel_i is the relevance of the retrieved element at the position i .

However, the DCG is not normalized, which makes it difficult to compare different queries with varying result sizes. To deal with this, the normalized DCG (nDCG) is introduced. The normalization (nDCG _{n}) is performed by dividing the DCG by an ideal DCG (IDCG _{n}). The latter is derived by taking the top- n relevance judgments ranked in decreasing order of relevance and calculating the DCG _{n} of this ideal order. The nDCG _{n} measures could be averaged over all queries to obtain an average nDCG _{n} of the system.

One work (Sinha et al., 2018) used an improved variant of the DCG, called Rank-DCG (Katerenchuk & Rosenberg, 2016). Dosso and Silvello (2020) propose a triple-based DCG (tp-DCG) based on a new notion of relevance (graph-based) that is calculated as the fraction of relevant triples in a subgraph-shaped answer.

MRR. The reciprocal rank aims at evaluating the system with respect to its capability of rapidly finding the first relevant result. The MRR is the average of reciprocal ranks over all queries Q , where pos_i is the position of the very first relevant result to the query i (Craswell, 2009):

$$\text{MRR} = \frac{1}{Q} \sum_{i=1}^Q \frac{1}{\text{pos}_i}. \quad (8)$$

7 System Summary Overview

Tables 3 and 4 provide an overview of all systems described in detail in Section 8 sorted in ascending chronological order. In addition to the aspects introduced in Section 5, the following criteria are also reported: the underlying *data*, the *goal* in terms of returned results, the availability of an *index*, the *search algorithm* used, the type of the *intermediate result*, and the specific *ranking*. Most of the approaches either work over any graph-structured data or over KGs (RDF). Furthermore, some of the works also require the existence of a data schema (ontology) to answer the query, and only one work operates on attributed graphs. The majority of works are graph-based (22), followed by summary-based (5), schema-based (4), and virtual document-based (4) approaches. Overall, the aim is to retrieve the top- k relevant answers, mostly ranked in the order of relevance to the query. However, a fraction of works only seek to return one relevant result or deal with the naive scenario of retrieving all relevant answers.

Even if we notice an inconsistency in the naming of the search algorithms used, works that aim at extracting tree or subgraph-shaped results usually rely on graph traversal algorithms (e.g., backward, breadth-first, depth-first, or random walk). The latter are either adapted to work with a specific index structure (Dalvi et al., 2008; Kacholia et al., 2005; Zhong & Liu, 2009), or a specific result (Ghanbarpour et al., 2020; Lu et al., 2019; Zhang et al., 2022), or data graph shape (Han et al., 2017; Zheng et al., 2016). Some works also exploit approximation algorithms that aim at computing the top- k best solutions of a specific optimization problem (Kargar & An, 2015; Park & Lim, 2015) or optimized enumeration algorithms (Golenberg et al., 2008; Kimelfeld & Sagiv, 2008). Most of the systems that aim at generating SPARQL queries construct a graph-shaped intermediate result that is translated to a structured query. The final answer type is, in most cases, a tree (15), followed by structured queries (9), and subgraphs (7). Only two works return triples and one retrieves entities in addition to triples. A newly defined answer type (key-core) is used by a single system.

Table 3. Overview of Surveyed Methods for Keyword Search over Graph-Shaped Data.

System	Data	Search Space	Goal	Search Algorithm	Intermediate Result	Final Answer Type	Answer Ranking	Answer Scoring
Bhalotia et al. (2002)	Graph	Graph-based	Top-k	Backward	—	Rooted directed tree	IMP	Distinct (E, N)
Kacholia et al. (2005)	Graph	Graph-based	Top-k	Bidirectional	—	Rooted directed tree	IMP	Root-keyword (E), distinct (N)
He et al. (2007)	Graph	Graph-based	Top-k	Index-based bidirectional	—	Distinct rooted directed tree	COMP, IMP, TXT	Root-keyword (E), distinct (N)
Kimelfeld and Sagiv (2008)	Graph	Graph-based	All	Threaded enumeration	—	Minimal tree	—	—
Dalvi et al. (2008)	Graph	Summary-based	Top-k	Multi-granular graph-based	—	Minimal rooted directed tree	IMP	Root-keyword (E), distinct (N)
Golenberg et al. (2008)	Graph	Graph-based	Top-k, All	Enumeration (polynomial delay)	—	Minimal rooted directed tree	COMP, IMP, DIV	Distinct (E, N)
Tran et al. (2009)	RDF	Summary-based	Top-k	Backward	Subgraph	Structured query	COMP, IMP, TXT	Root-keyword (E, N)
Zhong and Liu (2009)	Graph	Summary-based	Top-k	Composed subgraph	—	Distinct rooted tree	COMP	Root-keyword (E)
Sixtharn (2014)	RDF + Onto	Schema-based	Top-1	Node merging/expansion	Rooted tree	Structured query	COMP, IMP, TXT	—
Zhang et al. (2014)	Graph	Graph-based	Top-k	Shortest path	—	R-clique	COMP, IMP, SEM	Distinct (E, N)
Bae et al. (2014)	RDF	Graph-based	Top-k	Depth-first	—	Tree	COMP, IMP	—
Kargar and An (2015)	Graph	Graph-based	Top-k	Lawler's procedure	R-clique	Structured query	COMP	Keyword-pair (N)
Wang et al. (2015)	RDF	Graph-based	Top-k	Depth-first	Minimal tree	Distinct rooted tree	COMP, IMP, TXT	—
Park and Lim (2015)	Graph	Graph-based	Top-k	Threshold algorithm	—	Minimal tree	COMP, IMP, TXT	Root-keyword (N)
Mass and Sagiv (2016)	Graph	VD-based	Top-k	?	—	Subgraph	TXT	—
Zheng et al. (2016)	RDF	Graph-based	Top-k	Bipartite graph-based	Subgraph	Structured query	COMP	Keyword-pair (E, N)
Kharrat et al. (2016)	RDF + Onto	Graph-based	Top-1	—	Subgraph	Structured query	SEM	Distinct (E)
Han et al. (2017)	RDF	Graph-based	Top-1	Best-first	Subgraph	Minimal subgraph	COMP, TXT	—
Oukili et al. (2017)	RDF + Onto	Graph-based	Top-k	Shortest path	—	Structured query	—	—
Garcia et al. (2017)	RDF	Schema-based	Top-1	?	Minimal tree	Minimal subgraph	COMP, PROF	—
Sinha et al. (2018)	RDF	Summary-based	Top-k	?	—	Minimal subgraph	COMP, TXT	—
Rihany et al. (2018)	RDF	Graph-based	Top-k	Dijkstra	—	Structured query	IMP	—
Menendez et al. (2019)	RDF	Schema-based	Top-1	?	Minimal tree	Minimal rooted undirected tree	—	—
Lu et al. (2019)	Graph	Graph-based	Top-k	Canonical form-based	—	Minimal tree	COMP, TXT	—
Dosso and Silvello (2020)	RDF	VD-based	Top-k	Breadth-first	—	Minimal tree	IMP	Keyword-pair (N)
Bryson et al. (2020)	Attr. graph	Graph-based	Top-k	Random walk with restart	—	Subgraph	COMP, IMP	Keyword-pair (E)
Ghanbarpour et al. (2020)	Graph	Graph-based	Top-k, All	R-clique finding	—	Minimal covered r-clique	TXT	—
Kadlierakis et al. (2020)	RDF	VD-based	All	—	—	Triple, entity	—	—
Cheng et al. (2020)	RDF	Graph-based	Top-1	Best-First	—	Minimal tree	—	—
Jiang et al. (2021)	RDF + Onto	Summary-based	Top-k	Test different algorithms	—	—	IMP, TXT	—
Izquierdo et al. (2021)	RDF	Graph-based	Top-k	—	Tree	Structured query	—	—
Navarro et al. (2021)	RDF + Onto	Schema-based	Top-k	Breadth-first	Subgraph	Structured query	—	—
Zhang et al. (2021)	RDF	VD-based	Top-k	?	—	Tree	COMP	Distinct (E)
Shi et al. (2021)	RDF	Graph-based	Top-1	Shortest path	—	Minimal tree	IMP, SEM	Distinct (E, N)
Zhang et al. (2022)	Graph	Graph-based	Top-1	Key-core computation	—	Key-core	—	—

Note. ? = not mentioned in the paper; - = does not apply; **Onto** = ontology; **Attr** = attributed; **E** = edge; **N** = node; **VD** = virtual document; **COMP** = compactness; **IMP** = importance; **TXT** = textual; **DIV** = diversity; **SEM** = semantics; **PROF** = profile.

Table 4. Overview of Surveyed Methods for Keyword Search Over Graph-Shaped Data (Indexing and Ranking).

System	Index	Ranking
Bhalotia et al. (2002)	Keyword-element	Node prestige (in-degree), Edge proximity (in-degree)
Kacholia et al. (2005)	Keyword-element	Node prestige (PageRank), edge proximity (in-degree)
He et al. (2007)	Shortest paths	Distance, PageRank*, and TF/IDF*
Kimelfeld and Sagiv (2008)	?	no
Dalvi et al. (2008)	Keyword-element	Node prestige (PageRank), edge proximity (in-degree)
Golenberg et al. (2008)	?	Tree height, Node in-degree/out-degree, redundancy penalty
Tran et al. (2009)	Keyword-element, summary graph	Path length, popularity, keyword matching (TF/IDF*)
Zhong and Liu (2009)	Keyword-subgraph, keyword-element, element-subgraph	Distance
Sitthisarn (2014)	Shortest paths (BLINKS) with distance, keyword-element	Shortest path between root and keyword, term frequency, property frequency
Zhang et al. (2014)	No	Term frequency, PageRank, node degree, answer size
Bae et al. (2014)	Predicates, subject/object, literal, incoming/outgoing relations, similarity score	Distance (number of subjects in the path), semantic similarity (WordNet)
Kargar and An (2015)	Shortest paths, keyword-element	Node degree, shortest path
Wang et al. (2015)	Keyword-element, shortest paths	Tree diameter
Park and Lim (2015)	Keyword-element, shortest paths with node relevance	Node relevance, shortest path
Mass and Sagiv (2016)	Graph nodes/edges, node-text, virtual documents	Term frequency, distance, node degree
Zheng et al. (2016)	No	Term frequency
Kharrat et al. (2016)	?	Shortest path
Han et al. (2017)	No	Triple assembly cost
Ouksili et al. (2017)	Keyword-element	Term frequency, answer size
García et al. (2017)	Keyword-element	–
Sinha et al. (2018)	No	Distance
Rihany et al. (2018)	No	Distance, keyword matching score
Menendez et al. (2019)	No	InfoRank
Lu et al. (2019)	No	No
Dosso and Sivallo (2020)	Subgraph-text, subject	BM25, Markov random field model
Bryson et al. (2020)	No	Node score
Ghanbarpour et al. (2020)	No	Node degree, distance
Kadilierakis et al. (2020)	Triple-text	BM25, DFR, LM-Dirichlet, LM Jelinek-Mercer
Cheng et al. (2020)	No	No
Jiang et al. (2021)	Hierarchy of summary graphs	–
Izquierdo et al. (2021)	Keyword-element	Keyword matching, InfoRank
Navarro et al. (2021)	Keyword-element	no
Zhang et al. (2021)	Keyword-element, community	Structural compactness
Shi et al. (2021)	?	Semantic distance, PageRank
Zhang et al. (2022)	No	No

Note. ? = not mentioned in the paper; - = does not apply; * = recommended but not used; BLINKS = Bi-Level Indexing for Keyword Search.

Most of the works use a compactness-based ranking (18), followed by importance-based (15) and textual-based (12) ones. However, semantics-based (3), diversity-based (1), and profile-based (1) rankings are only rarely used. Almost half (16) of the works combine at least two ranking methods. The distinct answer scoring is mostly used (7), followed by the root-keyword (6) and the keyword-pair (4) scoring. Here also, we notice that two scoring types may be combined when, for each graph element, a different scoring is used. In addition, the specific answer scoring type may be applied only for either nodes or edges. Furthermore, for each system, we also document the type of index used and the specific rankings (cf. Table 4). An index is usually used to support keyword to graph element matching, or to accelerate the graph traversal by storing a specific set of shortest paths between two nodes. Almost half of the works (14) use a keyword to graph element index (keyword-element).

In general, we observe that most works have focused on traversing the whole graph to find tree-shaped answers for now. This shows that the overall trend is to enhance efficiency. However, answer ranking is not sufficiently investigated, which is demonstrated by the fact that in most cases, the most intuitive ranking type is used (e.g., compactness-based).

8 Detailed Overview and Classification of Existing Works

In the following, we give detailed descriptions of the collected relevant works. We structure them into subsections based on the *search space* aspect, since it appeared to best cluster relevant works into distinct groups. Other aspects will be clearly mentioned for each work in the short summary, together with the method of evaluation used. In total, we report on 22 graph-based, four schema-based, five summary-based, and four virtual document-based methods.

8.1 Graph-based Methods

Graph-based methods do not aim at reducing the size of the graph but directly search for answers over it. Since most of the works are graph-based, we further categorize them with respect to the *answer type* to improve readability. We also notice that all the answer types are covered in this case.

8.1.1 Tree. Bhalotia et al. (2002) presented Browsing AND Keyword Searching (BANKS), a system enabling keyword search over relational databases.

Method. The database is modeled as a directed graph by considering the tuples as nodes and the foreign key relations as edges. The query answer is defined as a *rooted directed tree*. BANKS models each relation between two nodes using two types of edges: forward (initial edges pointing from foreign key to primary key) and backward (additional edges in the opposite direction). This ensures finding a rooted tree directed away from the root. BANKS aims at finding an approximate set of top- k minimum Steiner trees (hard problem: NP-complete). It starts by mapping query terms to nodes using an index (text contained in graph elements), then the backward expanding search algorithm concurrently starts a shortest path finding algorithm from each keyword node with the aim of finding a node connecting all keywords (information node). To avoid waiting for all answer trees to be generated to sort them, an approximate solution is proposed. Generated trees are incrementally added (after duplicate detection¹³) to a small fixed-size heap (ordered by relevance score). When the latter is full and there are more candidate answers, the highest-ranked tree is returned and replaced in the heap. While the proposed heuristic does not necessarily ensure the retrieval of trees in decreasing order, the authors state that it works well even using small heaps.

Ranking. An *importance-based* ranking was used where each node is assigned a weight based on its importance (node prestige). The node weight is calculated as the fraction of the number of incoming edges and the maximum node weight, giving a higher score to nodes with more incoming edges. For the edges, a weighting based on the edges' importance (proximity) is proposed. Each forward edge is given a weight of 1, whereas a backward link (v, u) is given a weight that is directly proportional to the number of links to v coming from nodes having the same type as u . The edge weight is normalized using the minimum edge weight. The total score of an answer is then calculated by combining both node and edge scores by addition (*distinct scoring* for nodes and edges) or multiplication using a factor λ to control the effect of edge and node scores ($\lambda = 0.2$ performed best). The edge score gives higher relevance to smaller trees. Only root and keyword nodes are considered while calculating the total score of the nodes, and a node containing multiple search terms is counted multiple times.

Evaluation. Evaluation of performance, effectiveness, together with parameters effect were conducted using a part of the DBLP dataset (paper titles, their authors, and citations) and a dataset about the master/PhD dissertations in the IIT Bombay. The authors used an internally created gold standard with seven queries and their corresponding relevant answers. The rank difference between the ground truth answers and the actual answers was used as an evaluation metric.

Kacholia et al. (2005) state that the backward search in BANKS may not efficiently perform in the cases where: (1) a query keyword has many matching nodes in the graph, or (2) some of the visited nodes have a large in-degree, since it prevents exploring other directions until all incoming edges are visited.

Method. To overcome the previous shortcoming, this work proposes a bidirectional search paradigm to retrieve *rooted directed trees* which does not traverse the graph only backwards but also forwards, starting from potential answer roots in the direction of keyword nodes. For this purpose, the authors propose a node prioritization heuristic (spreading activation) that regulates the traversal. Keyword nodes are assigned an initial score given by $a_{n,i} = \text{nodePrestige}(u)/|K_i|, \forall n \in K_i$, where K_i is the set of nodes matching a keyword k_i . This prioritizes keyword nodes having a higher prestige and penalizes those matching a large number of nodes. Each node propagates an attenuated activation to its neighbors $\mu = 0.5$ and keeps the remaining $1 - \mu$. In case a node receives activation from different edges starting from a keyword k_i , the maximum activation $a_{(u,i)}$ is considered. The overall activation of a node is calculated as the sum of the activation scores originating from each keyword, to prioritize nodes that are close to a larger number of keywords.

Ranking. Since the focus was on the search algorithm, and not on the ranking of answers, the same *importance-based* scoring as proposed by BANKS was used with the multiplication-based relevance score. The only difference is the used final answer scoring for edges, which is *root-keyword* based, in contrast to BANKS.

Evaluation. The evaluation focuses on the efficiency aspect and compares the bidirectional search with BANKS and the Sparse algorithm (Hristidis et al., 2003) using the following metrics and five manually created queries: number of nodes explored and the nodes touched, as well as the time taken. Results show that the bidirectional search is faster than both baseline algorithms. A very small effectiveness analysis (precision and recall (cf. Section 6)) was conducted using the same queries and their corresponding results generated using an SQL query. Both backward and bidirectional performed equally with respect to effectiveness.

He et al. (2007) propose Bi-Level INDEXing for Keyword Search (BLINKS) and aim at exploiting indexes to precompute and store possible shortest paths. This is to deal with the fact that both backward and bidirectional search algorithms' performance relies more or less on the graph structure and the position of the different keywords in the graph, which makes the performance unpredictable.

Method. The index is filled using backward search and consists of two kinds of lists: (1) a *keyword-node list* that stores for each keyword the list of nodes that can reach it ordered by distance, and (2) a *node-keyword map* which stores the shortest distance between nodes and keywords, to optimize forward expansions.

BLINKS does not aim at indexing all possible paths (single-level index) since it will result in very large indexes for large-scale graphs. However, it proposes a bi-level index by dividing the whole graph into blocks (subgraphs): the *top-level block index* maps keyword/nodes to blocks, and the *intra-block index* stores detailed information within each block (e.g., intra-block keyword-node lists). The backward search strategy is also enhanced using *equi-distance* and *cost-balanced* expansions. BLINKS answers are *distinct rooted directed trees* to avoid answers that deliver only a few additional information compared to the rest, and technically allow more effective indexing.

Ranking. Again, the focus was not on the ranking, but on the query processing and indexing strategy. The authors propose a scoring function for the answer that combines three aspects (*compactness-based*, *importance-based*, and *textual-based*): sum of the shortest path distances from the root for each keyword node, sum of the matching score for each keyword with its corresponding node (e.g., TF/IDF), and the score of the answer root (e.g., PageRank; Page et al., 1999). However, for convenience, only the sum of the shortest path distances from the root for each keyword node is considered. The same aggregation is used for edge weights (*root-keyword scoring*).

Evaluation. Only efficiency experiments (comparison with the bidirectional algorithm Kacholia et al., 2005) were conducted, including execution time with different parameters (e.g., graph partitioning method) and index performance using 10 queries over the DBLP XML¹⁴ and IMDb¹⁵ datasets. Results reveal that BLINKS is overall faster than the baseline.

Kimelfeld and Sagiv (2008) use a different family of algorithms to efficiently retrieve all answers.

Method. The authors use threaded enumerators, which enable one enumeration algorithm to directly use elements produced by another one (or itself) instead of waiting until the latter finishes. The desired answers are *minimal trees*, which means that there exists no other subtree in the answer that connects the respective keyword nodes. Three types of answers are considered: *directed*, *undirected*, and *strong* (undirected with all keywords as leaves). The work focuses on the efficiency aspect and aims at proposing answer enumeration algorithms that solve the problem with polynomial delay between two output answers, except for the first answer, where the delay is exponential. For that, algorithms for each answer type are investigated, with the following different purposes with respect to how the answers are presented: *sorted*, *heuristically sorted*, and *unsorted*. The authors consider the sorted variant as the most wanted one and propose an algorithm with polynomial delay given an acyclic data graph and a directed answer. Algorithms for the unsorted scenario for all types of answers with polynomial delay are also proposed. Furthermore, a proposition is given about how the latter

can be enhanced to output the answers in a heuristic order. Answer ranking was not addressed, and there is no practical evaluation of the proposed approaches; instead, theoretical proofs are elaborated.

Golenberg et al. (2008) combine the methods proposed by Bhalotia et al. (2002); Kacholia et al. (2005) by using shortest-path iterators to find the first answer, and includes the work done in Kimelfeld and Sagiv (2008) to generate the rest of the nonredundant answers using a more practical and faster approach.

Method. Kimelfeld and Sagiv (2008) retrieve answers (*minimal rooted directed trees*) with a polynomial delay (ranking them by order of increasing weights) and is based on heuristics for generating Steiner trees. To simplify the problem, Golenberg et al. (2008) propose the usage of two rankings, a simple one while exploring the graph, and another one after generating candidate answers. A condition for this to work is that the primary ranking should be highly correlated with the final one. The problem is further simplified to increase efficiency and restricted to the enumeration of answers in a two-approximate order (“if one answer precedes another, then the first is worse than the second by at most a factor of 2”).

Ranking. In general, the authors combine *compactness-based*, *importance-based*, and *diversity-based* rankings. The primary ranking weight used during the exploration is the height of a subtree (“maximum length of any path from the root to a leaf”). The final ranking of generated answers is calculated as a combination of: an *absolute relevance score* based on the in-degree/out-degree of the nodes belonging to an edge, and the *redundancy penalty* which aims at producing diversified answers by “penalizing answers based on their degree of similarity to the ones that have already been given a final rank.” The total score of an answer is inversely proportional to its total weight using a *distinct scoring* for nodes and edges.

Evaluation. The efficiency of the algorithm was evaluated with an increasing number of query keywords. The correlation between the primary and final ranking, together with the effect of the redundancy penalty on ranking quality, is also evaluated. No information was provided on the used evaluation dataset/queries.

Kargar and An (2015) aim to find a practical algorithm to solve the hard problem of retrieving top- k answers using an approximation with polynomial delay. In contrast to the previous works, it does not directly retrieve trees but *r-cliques* that are transformed to tree answers at the end.

Method. First, needed indexes are prepared, which include a keyword-node index and a shortest path index between each pair of nodes within a certain distance. After finding keyword nodes in the graph, top- k *r-cliques* are calculated by adapting Lawler’s procedure (Lawler, 1972). First, a polynomial algorithm is called to retrieve the best one answer over the whole search space. Then, the latter is divided into subspaces according to the best answer, and the locally best answers in each subspace are found using the same polynomial algorithm. The subspace from which the best answer originates is further divided to find locally best answers. This process continues until all top- k answers are retrieved. Finally, a Steiner tree is produced over each *r-clique*.

Ranking. The edge weight used during the evaluation is based on the in-degree of the two edge nodes. The final answer score is calculated using a *keyword-pair scoring* considering the edge weight as pairwise node weight.

Evaluation. The efficiency evaluation is done using 15 queries over three datasets (DBLP XML, IMDb (MovieLens¹⁶), and Mondial¹⁷), where two approach versions (original and another one starting from nodes containing rarest keywords) were compared with BANKS, Dynamic (Ding et al., 2007), and an algorithm that produces communities as answers (Qin et al., 2009). In addition to the runtime also answer compactness was reported. Both proposed versions are faster than the other systems. The effectiveness was evaluated by first calculating the percentage of produced ground truth (generated using a naive algorithm that produces all *r-cliques*) answers using their approximate approach. In addition, a small user study was conducted with four manually created user queries over DBLP, and eight users were asked to rate the answers. The proposed system, BANKS, and dynamic achieve better precision than the community algorithm over all queries using the $P@k$ with k is equal to 10 and 2 (cf. Section 6).

Park and Lim (2015) propose an index-based approach for answering keyword queries over graphs that can be modified to also accept answers that have more than one node corresponding to a specific keyword.

Method. Similar to Wang et al. (2015), the approach starts by constructing an inverted index that stores for each keyword in the graph corresponding nodes (containing the keyword or connected to other nodes containing the keyword) together with their relevance score to the keyword. This index is used to find top- k *distinct rooted trees* using the threshold algorithm (Fagin et al., 2003; Güntzer et al., 2001).

Ranking. Each node is assigned a *textual-based* and *compactness-based* score that is calculated as the multiplication between two elements: a matching score calculated based on the number of occurrences of the keyword in the node, and a distance score calculated as the shortest path between the considered node and the one containing the keyword. The total relevance score of an answer tree is the sum of the relevance scores between the root and the keyword nodes (*root-keyword scoring*).

Evaluation. The system was compared (shortest distance set to 4) with BLINKS using 30 queries manually constructed over three datasets (Mondial, IMDb, and DBLP XML). The answer trees of the different approaches were rated by five

adjudicators, and the P@10 and P@20 were calculated. The proposed approach outperforms BLINKS for 73% of the queries, but the latter is still faster.

Lu et al. (2019) address the efficiency issue by ignoring some redundant intermediate results during graph expansion.

Method. The authors consider directed graphs and define a query answer as a *minimal rooted undirected tree* with a fixed size (number of nodes). The rationale behind their algorithm (canonical form-based search) is to consider only intermediate result trees that have a canonical form Zaki (2005) while searching the graph. Answer ranking was not addressed.

Evaluation. Experiments are run to evaluate the efficiency of the algorithm using the TPC-H benchmark¹⁸ (decision support database) and the IMDb¹⁹ database and 10 randomly generated queries. The work is compared with the other two systems that generate all answers without redundancy checking (Hristidis & Papakonstantinou, 2002; Kargar et al., 2015). Results reveal that the proposed approach is faster than the baseline for retrieving trees of size six. Furthermore, their algorithm always produces less number of intermediate results compared to the baselines.

Cheng et al. (2020) tackle the case where subtrees connecting all the keywords may not exist, or they exist, but they are far away from each other, resulting in a noncompact answer. While previous works retrieve top- k answers, here, only the best top-1 answer is found.

Method. A method is proposed that still ensures the compactness of the answers by not requiring that all keywords should be connected (relaxed answer). The authors formulate an optimization problem that aims at finding an optimal answer with a tradeoff between compactness (the smallest answer diameter) and the degree of relaxation (covering more keywords). For traversing the graph, a best-first strategy (CORE) is used, where the best search direction that may yield a better answer (minimal relaxed) compared to the current one is followed. The exploration terminates if the remaining possible answers cannot be better than the current best answer. The algorithm requires determining a specific desired answer diameter. Answer ranking was not addressed.

Evaluation. The evaluation was conducted over the RDF versions of Mondial (40 queries from Coffman's benchmark; Coffman & Weaver, 2014), LinkedMDB²⁰ (200 random natural language questions transformed to keywords), and DBpedia (438 queries from DBpedia-Entity v2; Hasibi et al., 2017). The baselines used are: a previous algorithm version of the same authors (CertQR+), and another algorithm for calculating Group Steiner Trees (GST) (Li et al., 2016) with edge weights assigned with one. By calculating the compactness of GST answers, it is shown that it fails to find answers in a lot of cases and also finds noncompact answers. In a second step, it is shown that their relaxed method finds complete answers in the case of a number of keywords lower than 10 and 100. The runtime of their approach outperforms CertQR+. Since CORE only computes the top-1 optimal answer, P@1 was used for calculating effectiveness over the DBpedia benchmark ("P@1 = 1 if a computed answer contained a gold-standard answer entity, otherwise P@1 = 0") and compare CORE with CertQR+ and the GST algorithm. Results show that both CORE and CertQR+ have comparable results, and the GST-based approach slightly outperforms both approaches with respect to the mean P@1.

Shi et al. (2021) aim at improving the performance of algorithms that calculate semantically cohesive *minimal trees*. The latter should consist of entities with minimal pairwise semantic distance. Similar to (Cheng et al., 2020), the top-1 answer is retrieved.

Method. The authors state that using the semantic distance as an additional factor increases the difficulty of the Steiner tree problem and thus propose two approximation algorithms: quality-oriented and efficiency-oriented. The answer is defined as a *minimal tree* following the same definitions by some of the previous works, and the aim is to find the single best answer. Each vertex is assigned a weight and each pair of vertices is assigned a semantic distance function. The problem of computing an optimum answer is approximated by first computing a set of relevant paths that lead to each keyword from a specific root node. After that, the candidate paths are transformed into an answer by merging the paths and removing unnecessary edges and vertices to make it a minimal tree. This tree is not necessarily optimal, but it is shown that it has a guaranteed approximation ratio. The quality-oriented algorithm finds a minimum cost path for each root, whereas the efficiency-oriented algorithm relaxes this definition by computing the path only for roots that are keywords.

Ranking. The total cost of an answer is calculated using a weighted additive function that aggregates the total PageRank-based vertices weights (*importance-based*) with the total pairwise semantic distance (*semantics-based*) of vertices following a *distinct scoring*. A concrete implementation of the semantic distance is out of scope, but it should be independent of the graph structure.

Evaluation. An evaluation was conducted using three synthetic KGs generated by the LUBM benchmark and three DBpedia subgraphs, including the whole graph. For each LUBM dataset, 250 queries were generated by varying two parameters: the number of keywords and the average number of vertices matched with each keyword. For DBpedia, queries from DBpedia-Entity v2 were filtered by removing keywords without matches. For the evaluation, existing metrics were reused: (1) a PageRank-based vertex scoring and (2) a semantic distance function calculated as the Jaccard distance between the sets of types of entities (Cheng et al., 2017). The latter requires entities having multiple types, which is not

available for LUBM. For that, the angular distance between entity embedding vectors using RDF2Vec (Ristoski et al., 2019) is calculated (structural semantics). The metrics used are: the runtime compared to BANKS-II (Kacholia et al., 2005) and DPBF (Ding et al., 2007), the ratio between the cost of the approximate and optimum answers (only on small graphs, using an exact version from their algorithm) and the cohesiveness ratio between a baseline answer (BANKS-II and DPBF) and their answer, where cohesiveness is defined as the sum of the pairwise semantic distance. The runtime of the efficiency-oriented algorithm was comparable to BANKS-II. The quality-oriented algorithm showed poor performance by reaching timeouts (200 s) most of the time. The effectiveness was evaluated by conducting a user study using 14–20 queries from the DBpedia-Entity v2 dataset with 16 participants (281 queries). Each participant was presented with the query and corresponding answers retrieved by different methods (DPBF, efficiency-oriented approach) and was asked to rate each answer on a scale (1–4).

8.1.2 Subgraph. Zhang et al. (2014) aim at retrieving *r-cliques* that are ranked based on the combination of various aspects.

Method. The first step is the weights' assignment to nodes and edges in the graph. After that, the search space is narrowed down by selecting only the top-*k* nodes that match the keywords using the node weights. Then, for each pair of selected nodes, the shortest path between them is calculated, and the top-*k* *r-cliques* are generated.

Ranking. The node weights are calculated as a function of the keyword frequency and PageRank, and the edge weights are calculated as a function of the degree of their respective nodes. The total score of an answer is a linear combination of the sum of the node weights and edge weights (*distinct scoring*), multiplied by other factors (answer size, keywords count in the query). The ranking is a combination of *compactness-based*, *textual-based*, and *importance-based* factors.

Evaluation. The evaluation was performed using five randomly generated queries from each of the DBLP and IMDb datasets. Relevance judgments were carried out by five users that were asked to score the nodes in the results based on the relevance to the query. Results show that the proposed approach is faster and more effective compared to the Dup-Free algorithm (Kargar et al., 2013).

Zheng et al. (2016) propose a method that transforms the RDF graph to a bipartite graph.

Method. The method aims at optimizing the efficiency of the keyword search task by taking advantage of the nature of the adjacency matrix of a bipartite graph. First, the RDF graph is transformed to a bipartite graph that consists of two sets represented as nodes with labels: one set with entities/classes and another one with properties. First, the keyword query is expanded using synonyms from WordNet. The result is then matched to graph elements and a connecting subgraph is extracted.

Ranking. The list of possible subgraphs is ranked using a *textual-based* score based on the number of keywords contained in a specific subgraph.

Evaluation. The evaluation is performed by comparing with three other systems (KREAG: Li & Qu, 2011; BLINKS and EASE: Li et al., 2008) using DBLP with 10 manually created queries. Both the execution time and effectiveness are reported using P@5, AP at 5 (AP@5), and the MRR as metrics (cf. Section 6). Results show that the proposed approach is faster and more accurate.

Ouksili et al. (2017) focus on the problem where keywords cannot be mapped to the underlying graph elements because of missing information (e.g., relation). In contrast to Zheng et al. (2016), the existence of a data schema (ontology) is required.

Method. The main idea is to leverage external knowledge defined in the form of patterns, which will add new missing triples. The pattern defines an equivalence between a property and a property path (e.g., two authors are collaborators if they have written the same paper). Those patterns are manually defined and are either domain-specific or generic by using standard properties (e.g., owl:sameAs). The approach starts with a fragment extraction step that extracts matching graph elements, and next, the set of elements is expanded using a pattern store. Possible fragment combinations are constructed by applying a Cartesian product. For each possible combination, the smallest *minimal subgraph* is retrieved.

Ranking. The used ranking is both *textual-based* using a term-frequency based matching score, and *compactness-based* using the size of the subgraph (sum of nodes and edges). Patterns are considered when computing the size of a subgraph by reducing the size by 1 if the subgraph contains a path from a pattern.

Evaluation. Two datasets were used: AIFB²¹ and a dataset about conferences. The actual approach using patterns was compared with two other configurations without external knowledge. The effectiveness is evaluated with P@*k* and the normalized DCG at *k* (NDCG@*k*) with *k* is equal to 5, 10, and 20 (cf. Section 6) using 10 randomly generated queries and the ratings of four users. The usage of patterns outperforms the other configurations. The effect of the query size on the execution time is also evaluated.

Rihany et al. (2018) have the same motivation as Ouksili et al. (2017) and focus on enhancing the keyword to graph element matching by using WordNet as an external knowledge base.

Method. The approach starts with keyword matching, where the first graph elements with exact matching labels are retrieved. If this step fails, related terms are retrieved from WordNet using different relations (e.g., synonyms). Since for each keyword a list of matching elements is retrieved, possible combinations are derived using the Cartesian product. For each possible combination, the smallest *minimal subgraph* is retrieved using the Dijkstra algorithm (Dijkstra, 1959).

Ranking. The result subgraphs are ranked based on a function that depends on the number of matching elements and the number of connecting nodes. This function favors small answers with more exact matching nodes (*textual-based* and *compactness-based*).

Evaluation. The evaluation was done using 10 queries from each of AIFB and a subset of DBpedia about movies. It is shown that the approach without external knowledge usage is faster and evaluates effectiveness using a subset of 10 queries and three users. Results reveal that P@10 and P@5 are always above 0.92 for both datasets.

Bryson et al. (2020) are motivated by the fact that shortest path-based methods will assign the same score to connecting subgraphs having the same length.

Method. The authors propose a function for assigning scores to nodes based on a random walk with restart. Different from previously mentioned works, they define *attributed graphs* as underlying data and aim at finding a top-1 or top-*k* subgraphs that cover all keywords. Since this is a rather hard problem, an approximation is proposed. The algorithm starts from the nodes with the so-called *rarest keyword* (appearing in the fewest nodes in the whole graph) and constructs its surrounding subgraph. This is a greedy approach that reduces the search space from the beginning and thus increases performance. Afterwards, a random walk with restart is run, which will assign a score to each neighbor node, and the best node is selected until the best answer is found. The entire process is repeated to retrieve the top-*k* answers. Since the random walks are the bottleneck of the algorithm, it is approximated using the Monte Carlo method (Bahmani et al., 2010), which does not require knowing the whole graph at each iteration. A parallel version of the same algorithm is also designed since each walk is independent of the other.

Ranking. The used ranking is *importance-based* and assigns scores to nodes relative to others by taking the whole graph into consideration. In contrast to PageRank, the proposed approach is query-dependent. The total answer score is calculated using the *keyword-pair scoring* considering node weights.

Evaluation. The method is evaluated using DBLP XML and IntAct PPI²² (molecular interaction data) datasets and compared with six systems (shortest-path-based, embedding-based, and PageRank-based). The quality of the answers is evaluated using different dataset-dependent metrics and the runtime of their distributed version. For IntAct PPI, the expression level of each protein is calculated, where “a cost value is assigned to each protein; the lower the cost, the higher the expression level” using 100 randomly generated queries with 2 to 6 genes. For the DBLP dataset, the h-index is used as a metric of quality with 100 randomly generated skill sets (2 to 6 skills). For the discovered subgraph, the average h-index of the members is calculated. Their approach results in answers with a higher expression level and a higher average h-index compared to the others.

Ghanbarpour et al. (2020) aim at retrieving minimal r-cliques called *minimal covered r-clique*. The motivation behind that is to overcome the limitations of finding Steiner trees, distinct root trees (losing results having the same root, considering the distance to the root not between each two keyword nodes), r-cliques, and r-radius Steiner graphs (Kargar & An, 2015; Kargar et al., 2014; Li et al., 2008) (not minimal) based on previous definitions.

Method. The work uses adapted versions (BKS and its improvement called BKS_R) of the Bron-Kerbosch (Bron & Kerbosch, 1973) and Tomita (Tomita et al., 2006) algorithms that originally aimed at finding maximal cliques (with the largest possible number of vertices). The algorithms are also improved to return nonduplicate minimal covered r-cliques. Furthermore, the authors propose algorithm versions (BKSM and BKS_{RM}) that are dedicated to also work on parallel configurations. This is possible since the answers are constructed independently of each other. An approximate version to generate top-*k* results with a polynomial delay is also proposed using an adapted version of Kargar and An (2015).

Ranking. The authors use an *importance-based* edge weight that is a function of the in-degree of its nodes as defined by Kacholia et al. (2005). Tests are performed using a uniform *compactness-based* scoring, by assigning one to all edges. The final answer score is calculated using a *keyword-pair scoring* considering the edge weight as pairwise node weight.

Evaluation. The different approximate versions of the proposed approach are evaluated against the r-clique and r-clique-rare algorithms (Kargar & An, 2015) over IMDb (MovieLens) and DBLP XML data using the same fove queries for each dataset. Results show that BKS and BKS_R are faster than r-clique and close to r-clique-rare, but still do not outperform it for both datasets. However, the parallel setting outperforms all the other algorithms. The other evaluated aspect is the compactness using the average percentage of cliques with the maximum *r*. In contrast to the proposed approach, for r-clique and r-clique-rare, not all retrieved results have a maximum distance *r*. To evaluate the quality of the retrieved results, the ground-truth weight is built by selecting the top-50 minimum weight answers generated by the nonapproximate version of BKS_R and calculating their average weight. Results reveal that BKS and BKS_R retrieve results with average uniform weights equal to the ground truth average, in contrast to r-clique and r-clique-rare.

8.1.3 Structured Query. Wang et al. (2015) propose an index-based approach that decreases memory consumption while performing keyword search over graphs.

Method. First, the RDF graph is transformed to an attributed graph that consists of interlinked subject nodes with text information (predicate, class, and predicate linking to another entity). The approach starts with an offline phase where the index is constructed. The index stores for each keyword in the data graph the corresponding list of pairs of nodes (node1 containing the keyword, node2 reachable from node1), and the shortest distance between the pair of nodes over 2 hops. The index is leveraged in an online phase to extract the top- k *minimal trees* using a depth-first search algorithm and then transform them to a corresponding SPARQL query.

Ranking. The ranking is *compactness-based* using the tree diameter, which is calculated as the maximum distance among the shortest distances of all pairs of nodes.

Evaluation. The evaluation mainly concentrated on the efficiency aspect, comparing the indexing time, memory usage, and the runtime with two other index-based approaches (Rclique Kargar & An, 2011, Gdensity Li et al., 2012). Results show that the proposed system is faster and requires less memory usage over the two datasets DBLP XML and DBLP2²³. The effectiveness was manually checked and compared with Rclique. The keyword queries were manually created by first randomly selecting 20 subgraphs with diameter four over DBpedia, constructing corresponding SPARQL queries, and formulating keyword queries by selecting one or two keywords from each node. Afterwards, it is manually checked if any of the top- k SPARQL queries generated by the systems match the ground truth.

Kharrat et al. (2016) propose a SPARQL query generation method specific to DBpedia. In contrast to Wang et al. (2015), it aims at retrieving only the top-1 answer and requires the existence of a data schema (ontology).

Method. The first step is query preprocessing using WordNet, followed by the keyword matching of query terms to DBpedia concepts/properties/instances. The next step is a query expansion with semantically related concepts using some properties (e.g., `rdfs:seeAlso` or `owl:sameAs`). Afterwards, ambiguous terms are filtered by calculating the pairwise relatedness between keywords. The latter is based on the intersection between the set of matching terms (DBpedia IRIs, text in this case) for each keyword. Then, all possible combinations are generated from the matching elements, candidate subgraphs for each combination are constructed using predefined graph patterns, ranked, and the top-1 result is translated to a SPARQL query.

Ranking. The subgraphs are ranked based on a *compactness-based* measure, which is calculated as the sum of the shortest paths between each pair of matching elements (*keyword-pair scoring*).

Evaluation. The effectiveness of the system is evaluated using 50 queries, and the following metrics: The fuzzy precision, which is calculated as the sum of the average correctness rate for each query divided by the queries returning results. The correctness rate of an answer is calculated as the set of correct terms that match the user's intention (subject, predicate, and object) divided by the set of all terms. The other used metrics are recall (queries returning results divided by the total number of queries) and the F1-measure (cf. Section 6) as a combination of the fuzzy precision and recall. No information was provided about the origin of the queries, the used ground truth, and the dataset. The reached results are as follows: 0.52 (F1), 0.43 (recall), and 0.5 (fuzzy precision).

Han et al. (2017) propose an efficient and accurate keyword search system over RDF data by leveraging bipartite graphs together with translation-based graph embeddings. Similar to Kharrat et al. (2016), the goal is to find just the top-1 result.

Method. The method starts with a *segmentation and annotation* phase where the user query is first tokenized, annotated using types (entity, class, and relation), and then matched to a specific graph element. The result is a ranked list of possible query annotations. The second step is the *query graph assembly*, where the detected graph elements need to be linked with each other. For that, the found possible matches are modeled as a bipartite graph where each pair of entities/classes and each set of possible matching relations is represented as a node, and there is a crossing weighted edge between the set of entities/classes and the set of edges. The algorithm tries to find the optimal connection (subset of crossing edges) called *assembly query graph* with the minimum cost and transforms it to a SPARQL query. The constraint here is that the query should ideally contain classes, entities, and as many needed relations, which are not always present given the nature of keyword queries, and in this case, the graph elements cannot be connected. The authors point to this issue and shortly state that, in this case, the missing relation should be first determined using a suitable algorithm for minimum spanning tree finding.

Ranking. Since the goal is to find just one best answer, the algorithm does not output a list of top- k possible answers. However, to know the optimal answer, a certain function should be used to score the possible subgraphs. The authors use a *semantics-based* scoring where edge weights are calculated using TransE²⁴ (Bordes et al., 2013), a translation-based graph embedding model. The total cost of a candidate subgraph is the sum of its edge costs.

Evaluation. The system is compared in terms of efficiency and effectiveness with two approaches: PBF (Ding et al., 2007), a graph-based approach, and SUMG (Tran et al., 2009), a summary graph-based approach. Two datasets were used: 6th Open Challenge on Question Answering over Linked Data (QALD-6; Unger et al., 2016) with 100 queries, and

Free917 (Cai & Yates, 2013) with 80 selected queries converted to keywords. Their approach using graph embeddings is also compared with two traditional link prediction methods. The system performed better than both other approaches and ranked the third place in QALD-6, competing with question answering systems. The most time-consuming part was the keyword-graph element mapping, but on average, the system was faster than the two baselines.

Izquierdo et al. (2021) propose a different method that does not rely on traversing the graph.

Method. The very first preprocessing step is to calculate offline the so-called *K-minimum value (KMV)-synopses* over the whole graph, where “A *KMV-synopsis* of a set defines a random sample of size k .” They are calculated because they can be used later to estimate the Jaccard similarity of sets. First, keywords are matched to the literals’ content of the RDF graph. If for one keyword several matches are possible, the one with the highest score, consisting of the combination of literal and node matching (InfoRank; Menendez et al., 2019) scores, is selected. Then, a so-called initial *query forest* (collection of disconnected trees) is built, consisting of keywords and the list of domains and ranges of the properties represented as nodes related to the keyword literals. The goal is to connect it to form a *query tree* (Steiner tree) by applying three types of operations: *node fusion*, *edge addition*, and *tree expansion*, before translating it to a SPARQL query. The node fusion step tries to combine the forest nodes that link to a basically similar set of entities. The similarity of those sets is calculated using the Jaccard similarity. If the latter is high, the nodes are combined, which indicates that in the next step, the goal is to find entities that have both merged properties that match the initial corresponding keywords. Next, edge addition is applied by identifying object properties that could relate two entities using an estimated set similarity that is computed using the KMV-synopses. The last step is the tree expansion which is a relaxed version of the edge addition, which requires having either domain or range, not necessarily both, to be similar to the other sets. The two last steps are repeated until more connected trees are formed. A cleaning step removes unnecessary edges so that only leaves corresponding to matching nodes are left. Only one tree having the larger score is selected and translated into a *structured query* (SPARQL).

Ranking. The used ranking is both *textual-based* using keyword matching scores and *importance-based* using InfoRank.

Evaluation. The evaluation compares the approach with: (1) a schema-based RDF keyword search tool (García et al., 2017; previous work of some of the same authors), and (2) a virtual document-based approach (Dosso & Silvello, 2020). For (1), an adapted version of Coffman’s benchmark using Mondial and IMDb databases with a total of 64 queries is used. The ground-truth answers were generated using a tool for automatic benchmark construction for keyword search (Neves et al., 2021). The effectiveness is measured using the following metrics: MAP (cf. Section 6), Top-1, and the MRR. The reached scores are as follows for Mondial and IMDb, respectively: MAP of 0.96 and 0.76, Top-1 of 0.96 and 0.76, and an MRR of 0.79 and 0.72, and thus also outperforms the baseline. The average execution time was 11.4 s and 0.60 s for IMDb and Mondial, with no big differences compared to the baseline. For (2), the same benchmark was used as in the virtual document-based baseline, including the following datasets: LUBM,²⁵ BSBM,²⁶ IMDb,²⁷ and DBpedia²⁸ and also using the same metrics. The proposed approach outperforms the baseline for all metrics.

8.1.4 Other. Bae et al. (2014) propose a system that exploits indexing and semantic similarity scores to find the triples relevant to a specific keyword (not query).

Method. The first step is the offline preparation of five indexes: predicates, subject/object, incoming/outgoing vertices, literals, and similarity scores between two literals. The approach starts by finding the literal vertex corresponding to a certain keyword, filters candidate nodes without a common predicate with the input keyword node, traverses the graph in a depth-first manner, and considers only nodes (literals) with a semantic score greater than a specific threshold. The list of relevant triples is presented to the user.

Ranking. The ranking of target nodes (or triples) is performed based on the similarity score, which is calculated as a combination of a distance score (number of subjects relating source and target), and a semantic similarity depending on the most specific class subsuming the classes of the pair of nodes (Lin & Sandkuhl, 2008), calculated over WordNet (Fellbaum & George, 1998). The ranking is a combination of *compactness-based* and *semantics-based* factors.

Evaluation. Only an efficiency evaluation was performed using 10 manually created keywords over three subsets of DBpedia (Lehmann et al., 2015). The approach was compared with the RDF management system Jena (McBride, 2001) by retrieving the first two results (using SPARQL queries for Jena). Results show that the proposed system is faster.

Zhang et al. (2022) do not focus on finding top- k answers but aim at presenting the user with a view that better shows how the answers are correlated and how they may overlap. Furthermore, authors also aim at extending the way a user can interact with the retrieved answers by allowing answer manipulation, for example, intersection or union.

Method. To achieve the previous purpose, a specific answer type is defined. The authors represent a so-called *key-core*, which is a subgraph containing highly related answers. Answer ranking was not addressed.

Evaluation. The evaluation is performed using two datasets, DBpedia and DBLP, using 500 randomly selected queries from the datasets’ vocabularies. The effectiveness is manually evaluated on five example queries. For testing efficiency, their approach based on graph traversal is compared against a defined naive baseline (retrieve all key-components and then

compute the key-core by union of all key-components) while also varying the graph size, the number of keywords, and other used thresholds such as the distance maximal to the keywords.

8.2 Schema-based Methods

In contrast to graph-based methods, schema-based ones require the existence of a data schema and thus work on a very small graph. We notice that the answer type produced by this method is always a structured query.

Sitthisarn (2014) propose a bidirectional approach for keyword search with defined answer types over RDF graphs with an available data schema (ontology).

Method. The proposed system constructs two indexes: the keyword-element index, and an extended version of the single-level index by BLINKS that stores for each node the distance to the other node and the direction. The algorithm takes as input a set of nodes corresponding to keywords and a desired answer type (root of interest). The first step is the node merging, where it is checked whether each pair of keyword nodes can be directly connected, have the same class, or one is a subclass of the other. In the latter cases, the nodes are marked as nonactive. The remaining active nodes are bidirectionally expanded to connecting nodes in a round-robin manner until finding a connection with a root. The results are ranked and the best *rooted tree* is transformed to a SPARQL query.

Ranking. The results are ranked by combining *compactness-based*, *textual-based*, and *importance-based* factors based on a previous version of the same authors (Sitthisarn et al., 2011). The overall tree score is given by the multiplication of three scores, each is a function of: the shortest path between the root and keywords, keyword frequency, and the predicate frequency.

Evaluation. The effectiveness is evaluated with 20 manually created queries and five ontologies. The result of the current approach and its previous version (not bidirectional) are compared with a manually created SPARQL query using the MRR as metric. Results show that the new approach outperforms the old one.

García et al. (2017) propose a tool to simplify access to specific data from an industry use case. Similar to Sitthisarn (2014), it first retrieves a tree-shaped intermediate result before generating the top-1 structured query. However, it does not require a data schema (ontology).

Method. The overall aim is to construct one SPARQL query to answer a keyword query. The algorithm operated on the data schema where keywords are matched to literals after removing stop words. The matched structure is called *nucleus* and consists of a class, a list of properties, and property values. Each *nucleus* is assigned a matching score that is used to construct a *minimal tree* that covers all keywords. While a prototypical user interface was proposed, answer ranking was not addressed.

Evaluation. The evaluation was conducted over an industrial database (hydrocarbon exploration) and using Coffman's benchmark (Mondial and IMDb (IMDb-MO)²⁹). The relational databases were first converted to triples, and the retrieved results using the proposed approach were compared with the ground truth. Results were only analyzed without using any evaluation metric, where 64% (Mondial), 75% (IMDb) of the 50 queries were correctly answered. Using six queries for the industrial database, it is shown that the execution time is reasonable.

Menendez et al. (2019) focus on improving the ranking of results by proposing a new definition of node importance in RDF graphs.

Method. The method aims at generating one structured query (SPARQL) as in García et al. (2017); Sitthisarn (2014) to produce a list of ranked entities. First, the system finds nodes (instances and classes) whose labels match the keywords. Next, for each query keyword, the node with a higher *accum* score is selected (number of keywords that appear in the label). In the case of equal *accum* scores, the disambiguation is performed using another specific node scoring function (*info_score*). The next step is the linking of the classes corresponding to the selected nodes over the graph schema. This consists of finding the *minimal Steiner tree*, but no details are given about how this tree can be found. The authors only claim that this step will give the information that, for example, two selected nodes are connected through one specific property. This connected tree is transformed to a SPARQL query and the results are ranked based on the sum of the *info_score* of the tree nodes.

Ranking. An *importance-based* ranking of entities is used by proposing a new metric *InfoRank* for scoring instances, classes, and object properties that is based on the *informativeness* of an instance, which is defined as the number of its data properties (literals). Each object property is assigned a score that is the maximum of the sum of the *InfoRank* of the subject and object (over all triples involving the property), divided by the sum of the *InfoRank* of all other existing properties. Next, the PageRank of each instance is calculated using the object property score as edge weight. The final score assigned to an instance (*info_score*) is the PageRank after x iterations multiplied by the *InfoRank* of the instance. No ranking of answer trees is proposed.

Evaluation. The proposed ranking score was evaluated by comparing it with the PageRank of some classes, instances, and properties over two datasets³⁰ : IMDb (IMDb-MO) and MusicBrainz³¹ (enriched with DBpedia). This has shown that the *info_score* gives higher scores to the most important classes in both considered domains compared to PageRank. Other experiments were conducted using Coffman's IMDb (50 queries adapted to RDF) and QALD-2³² MusicBrainz (25 queries) with different ranking measures (e.g., PageRank or HITS; Kleinberg, 1999). *InfoRank*-based ranking achieved the highest MAP.

Navarro et al. (2021) aim at generating the set of *structured queries* (SPARQL) that answer a keyword query. As in Sitthisarn (2014), the data schema is also required.

Method. Their approach starts with an entity identification module based on the Stanford Regexner Annotator (Manning et al., 2014), which annotates keywords with candidate classes/individuals from the underlying ontology/RDF repositories. Filter expressions are also detected using a gazetteer. Afterwards, entity combinations are generated, since a keyword can be associated with multiple entities. Less specific annotations (ontology hierarchy) are considered as redundant and thus removed. The answer search algorithm operates on the ontology of the target KG, extended with the detected individuals and filters. For each different candidate entity of a specific keyword, such an ontology-based graph model will be built. The graph exploration is done by starting Breadth-First Search from a specific keyword node, and the neighbors exploration continues until all candidate keyword elements are included. No cost function is used during the graph exploration. The result is a *rooted tree* that is further cleaned by removing unwanted edges and leaves. Each keyword interpretation results in a corresponding tree. Finally, each tree is translated to a SPARQL query. Only one representative is selected in case of redundant SPARQL queries. Answer scoring is not addressed.

Evaluation. The evaluation was conducted using the QALD-7 dataset for question answering over DBpedia, where each question is associated with a set of keywords. Some types of questions were discarded (e.g., boolean) and the keywords were manually reviewed and adapted. The final test set contained 136 queries. Precision, recall, and F1-score are used as metrics to evaluate the generated SPARQL queries against the provided ones. In this case, different cases (e.g., correct/imprecise/partial answer) are defined. The evaluation reveals a precision of 0.52 and a recall of 0.60, considering the best answer for each query. The runtime was also reported with an average time of ≈ 57 ms over all queries (without including synonyms from WordNet in the index).

8.3 Summary-based Methods

Summary-based methods also aim at searching over a reduced-size graph but without relying on the existence of a data schema. The works under this category aim at retrieving trees (Dalvi et al., 2008; Zhong & Liu, 2009), subgraphs (Sinha et al., 2018), or structured queries (Tran et al., 2009).

Dalvi et al. (2008) propose an approach to deal with large graphs that may not fit into memory.

Method. This work represents the graph in a hybrid manner (multi-granular) and aims at retrieving *minimal rooted directed trees*: (1) an in-memory summary graph, where nodes are clustered to create supernodes, and (2) a disk-resident part that contains the nodes appertaining to each cluster, together with their adjacency information.

Furthermore, two alternative algorithms that adapt existing exploration algorithms are suggested to work over multigranular graphs, with the goal of reducing input output calls.

Ranking. The same *importance-based* ranking as in Kacholia et al. (2005) was used.

Evaluation. The authors implemented their approaches by extending BANKS. DBLP (8 queries) and IMDb (4 queries) were used as evaluation datasets and compared different configurations of the proposed algorithm with a schema-based approach (the Sparse algorithm; (Hristidis et al., 2003)) with respect to the execution time, recall, and cache misses per query.

Zhong and Liu (2009) concentrate on improving the efficiency for complex queries where the keywords have a lot of matching nodes in the graph (*frontier*).

Method. This paper proposes a solution to reduce the frontier that aims at performing search only on parts where actually the candidate answers are more likely to reside. The rationale behind their approach is the fact that if we have multiple matching nodes for each keyword, only a few of the different keyword nodes will be tightly related to each other. A graph index that maps keywords to a set of subgraphs from the whole graph is used, together with additional indexes such as keyword-vertex and vertex-subgraph indexes. While creating the subgraph index, a bound for the size of the extracted subgraphs is set. In this way, answering a query would mean searching only subgraphs that actually contain all keywords. The data graph is considered as undirected and defines possible answers as *distinct rooted trees*. A so-called composed subgraph search that it performed over a combination of matched subgraphs is used. The graph exploration algorithm is similar to the backward search (Bhalotia et al., 2002).

Ranking. A *compactness-based* ranking is used with the distance as a ranking factor, and the total score of an answer is defined as the sum of the shortest distances between the root and each keyword node (*root-keyword scoring*) without scoring the nodes.

Evaluation. This work evaluated effectiveness and efficiency and compared with BLINKS using the DBLP XML datasets with 12 queries. As a metric for calculating effectiveness, the minimum scores and average scores of the top 10, 25, and 50 answers found by the baseline and their approach are compared, which shows comparable results. The analysis of the execution time shows that their method is overall faster compared to the baseline.

Tran et al. (2009) propose another paradigm for keyword search over graphs (RDF).

Method. Instead of presenting the user with the top- k final answers that may originate from various queries, top- k possible *structured queries* are returned. Furthermore, the returned answers are not trees, but subgraphs, and user keywords could also be mapped to edges. A backward-based search algorithm that operates on a summary graph instead of the whole graph is used to improve efficiency. There is no distinction between incoming and outgoing edges during the exploration, and when reaching a node, all its neighbors are taken into consideration, preventing cyclic expansion. However, the direction of the edges is important while generating the SPARQL queries corresponding to the subgraphs. The summary graph is derived from the data graph by representing each collection of entities by their corresponding class and aggregating all entities without types using a class called *Thing*. The edges between the summary graph nodes mirror the edges between their corresponding instances.

Ranking. The score of an answer subgraph is calculated as the sum of its path scores following the *root-keyword scoring* while aggregating both edge and node weights. The cost of a path is computed as the aggregation of its element costs. Three different ranking functions are introduced: (1) a *compactness-based* ranking using the path length that is defined as the number of elements in a path, (2) an *importance-based* ranking using the popularity of path elements (nodes and edges), where node/edge popularity is given as a function of the number of original instances/edges that were clustered into the corresponding class/edge divided by the number of nodes/edges in the summary graph. The latter is defined in a way that elements with higher popularity should have lower cost and thus contribute to answers with higher scores, and (3) *textual-based* ranking using the keyword matching score which is incorporated by dividing one of the previously defined cost functions by a matching score that is either ranging between 0 and 1 if the element is corresponding to a keyword or is set to one otherwise.

This should reflect the fact that higher matching scores reduce the path cost. The authors use a matching score (between keywords and labels in the graph) that combines both syntactic and semantic similarity (e.g., WordNet) and recommend TF-IDF in case of labels with a sufficient number of terms.

Evaluation. The approach was evaluated using 30 queries together with information needs for DBLP and nine queries for TAP (KG describing sports, geography, music, and other domains). Twelve participants rated the returned structured queries (query is correct if it matches the information need), and the MRR was used as a metric for the assessment of effectiveness. Results reveal that the textual-based ranking function using keyword matching performed best compared to the other two previously proposed functions. The authors also perform an evaluation of the query processing time, comparing it with other approaches (e.g., Kacholia et al., 2005 and other indexing-based methods) together with an investigation of the impact (on the time) of the number of queries to retrieve, and index performance (also using LUBM). Results show that overall, the proposed approach is faster than the other baselines.

Sinha et al. (2018) propose a method that takes advantage of the user profile information. This is so far the only work that exploits profiles to rank answers.

Method. The proposed approach starts by extracting a structural summary graph from the data that consists of classes, relations between them derived from the relations of corresponding entities, and the union of entity labels. First, keywords are matched to corresponding labels in the graph, and for each possible matching, a *minimal subgraph* that connects corresponding classes is constructed. For each possible matching, only the smallest connecting subgraph is considered. Subgraphs producing empty results are discarded.

Ranking. The returned subgraphs are ranked based on their similarity with available profile graphs. The latter are either added explicitly or implicitly after a search activity. The similarity is a combination of four metrics: concept similarity, relation similarity, entity property similarity, and entity connection similarity.

Evaluation. The evaluation was conducted using 10 queries for each of the Jamendo³³ (music) and DBpedia³⁴ datasets. The ground truth was constructed by two users who rated every pattern graph based on its relevance to the result subgraph. The metrics used for evaluation are Kendall-Tau (Agresti, 2010), P@ k and Rank-DCG (cf. Section 6). Results reveal a Kendall-Tau above 0.5 and a Rank-DCG above 0.6 for most queries. The average P@5 is at least 0.5 for both datasets. In addition, efficiency experiments show that the system responds in a reasonable time that allows user interaction.

Jiang et al. (2021) propose an index (BiG-Index) to speed up finding answers for keywords over KGs and test the efficiency of some existing works (He et al., 2007; Kargar & An, 2011) using it.

Method. The rationale behind the approach is first to replace the labels of instances with their corresponding generalized ontology labels. Afterwards, this generalized graph version is summarized by merging similar subgraphs into one representative. The last two steps are recursively repeated to build a hierarchy of graphs (indexes). Furthermore, details on how to answer a query using the proposed index are provided. Answer ranking was not addressed.

Evaluation. The performance of some algorithms (BLINKS and r-clique; Kargar & An, 2011) for keyword search over graphs is compared both with and without using the index. The evaluation is performed using both real and synthetic datasets (YAGO3³⁵, DBpedia, and IMDb with a total of 18 synthetic queries). Since the methodology requires the usage of an ontology, both DBpedia and IMDb were used with an ontology generated from YAGO3. Results reveal a decrease of the runtime by 50.5% and 29.5% for BLINKS and r-clique, respectively.

8.4 Virtual Document-based Methods

Virtual-based methods transform the graph content into text documents and apply techniques for full-text search. In the following existing works, the retrieved answers are either trees (Dosso & Silvello, 2020; Mass & Sagiv, 2016; Zhang et al., 2021) or triples and entities (Kadilierakis et al., 2020).

Mass and Sagiv (2016) propose a virtual document-based method that tries to deal with efficiency by reducing the search space and enhancing effectiveness by using a scoring function that combines both query-dependent and -independent scores.

Method. The method starts by preparing a virtual document representation for the whole graph. For each node, a virtual document is created by concatenating the textual information (title and attributes) of the node itself and its neighbors within a specific distance. First, a two-level node filtering is performed to reduce the search space. Nodes whose virtual document contains every query keyword are ranked based on a specific score, and the top- k are selected (*selected roots*). Or each keyword, the same filtering is performed over the virtual documents of the roots to end up with a list of top- k *selected keyword nodes*. Next, *minimal trees* connecting the selected roots and keyword nodes and constructed, and duplicate trees are removed. The answer trees are then transformed to virtual documents by grouping the textual information for each node and then ranking.

Ranking. The ranking used by the selection of nodes and ranking final answers is calculated using a Markov random field model (Metzler & Croft, 2005) by combining query-dependent and -independent features. Those features are functions of the frequency of a query term in a virtual document, the node importance, which is proportional to its degree, and the edge weights that are set to one to favor smaller answers.

Evaluation. The evaluation was performed using Coffman's benchmark, where the MAP (top-1000) was compared with four other systems (e.g., BANKS). Their approach outperforms all the baselines and shows a good trade-off between effectiveness and efficiency.

Dosso and Silvello (2020) propose an approach based on virtual documents, since the authors claim that this is a promising method in terms of efficiency. As in Mass and Sagiv (2016), the retrieved answer type is *minimal tree*.

Method. This work introduces Topological Syntactical Aggregator+BM25 (TSA+BM25; Robertson & Zaragoza, 2009) and Topological Syntactical Aggregator+Virtual Documents Pruning (TSA+VDP). TSA is a very first offline phase where the virtual documents are created by clustering triples with related concepts. In practice, it first builds subgraphs having a specific distance around single topics (classes) together with their corresponding text documents (literals, IRIs, and predicate strings). Only predicates with a frequency higher than a threshold are taken into account. In the online phase, an initial list of ranked documents relevant to the user query is retrieved. The next step aims at merging overlapping subgraphs corresponding to the top- k candidate documents. If the intersection of the subgraphs' triples is greater than a threshold, they are merged. Next, virtual documents corresponding to the merged graphs are created, a second BM25 ranking is performed, and the relevant subgraphs are returned to the user (TSA+BM25). The next step is VDP, which considers the union of the top- k subgraphs of the last step as a new graph. *Minimal trees* with a defined radius containing all keywords are produced using Breadth-First Search, and nonrelevant triples (not containing any keyword) are pruned. The last step is a final ranking of results.

Ranking. The initial and second rankings are both *textual-based* using BM25. The final ranking combines *compactness-based* and *textual-based* factors using the Markov random field model that considers unigrams and bigrams within the virtual document and the distance of the words from the root in the graph.

Evaluation. The evaluation was conducted using real (LinkedMDB and IMDb with 100 queries, and DBpedia-Entity v1 (Balog & Neumayer, 2013) with 50 queries) and synthetic databases (LUBM with 14 queries and BSBM (Bizer & Schultz, 2009) with 13 queries). Only one best answer is considered as ground-truth answer. The latter is the result of a manually created SPARQL query corresponding to the keywords. The authors propose a new definition of precision, recall, and the DCG based on a new metric *signal-to-noise ratio* that gives a score to a returned graph based on the intersection of its

triples and the ones from the ground truth. Their approach is compared with other three (Elbassuoni & Blanco, 2011; Le et al., 2014; Mass & Sagiv, 2016) virtual-document based baselines, using average tp-DCG (cf. Section 6), recall, P@1, P@5, and runtime over all queries. Overall, TSA-based systems outperform the baselines, but all in general have low precision. Considering the online runtime, the proposed approach was among the fastest systems.

Kadilierakis et al. (2020) aim at adapting an existing information retrieval system (Elasticsearch³⁶) and thereby use traditional document indexing and retrieval instead of traversing the underlying graph (Dosso & Silvello, 2020) or generating structured queries.

Method. The textual content of a triple is considered as a virtual document that is indexed and retrieved. Two index versions are evaluated depending on the extent of the stored textual content: (1) a baseline index stores only textual content of the considered triple, for example, text of the URI, and (2) an extended index considering also information on other properties of the triple’s resources, for example, `rdfs:label`.

Ranking. The authors use the *textual-based* ranking functions provided by Elasticsearch to rank triples and derive a ranking for the entities involved in a triple using the DCG formula.

Evaluation. The DBpedia-Entity v2 was used as a test collection using the `nDCG@100` and `nDCG@10` as metrics. Its goal is to test the influence of different configurations of Elasticsearch (e.g., query type) together with the index content on the effectiveness of the search. Their system (Elas4RDF) was compared to the unsupervised methods tested in the context of DBpedia-Entity v2 and reached comparable results when used with the extended index and the BM25 ranking. The average query time was also reported: ≈ 0.7 s and ≈ 1.6 s for the baseline and the extended index, respectively. A user interface for the same system is proposed in Nikas et al. (2020) with a focus on functionality and usability evaluation.

Zhang et al. (2021) propose a pipeline consisting of an offline and online stage. The key contribution is the usage of community detection techniques to create a group of entities (subgraphs) belonging to the same topic (e.g., computer security in DBpedia). This is used as an index to accelerate answer retrieval.

Method. During the offline phase, the two indexes are created: (1) the entity index is built by transforming each entity together with its data properties into a virtual document whose terms are indexed, and (2) the community index, which maps each entity to its community of entities.

The online phase starts with mapping query keywords to candidate entities using the entity index. For each entity, the common community containing all entities is constructed, which is, in the end, a subgraph of the whole RDF entity graph. Finally, a ranked list of trees connecting all keywords (Steiner tree) is computed. No further details are given on the tree’s construction algorithm or the specific structural *compactness-based ranking*.

Evaluation. Three aspects were compared with an index-based system (EASE) using five queries for each dataset (DBLP and KMap³⁷): index building (time and storage), runtime, and effectiveness. The latter is based on answer completeness with respect to the relations connecting keywords using the answer of a method that directly works over the whole graph (not index-based) as a baseline. Both systems have comparable results.

9 Conclusion and Future Directions

In this survey, we have systematically selected, classified, and provided an overview of 35 research papers. We have derived four overall aspects for classifying related works: (1) search space, (2) answer type, (3) answer ranking, and (4) answer scoring. Each of those aspects is further specified by defining different possible aspect types (e.g., compactness-based answer ranking). In the following, we highlight some potential directions for future research using a structure that is based on the typical components of a system for keyword search over graph-shaped data.

9.1 Keyword Mapping

Enhancing Entity Linking. Most of the works attempt to map keywords with all possible graph elements that contain one of the keywords in their textual description (string matching). This approach not only affects efficiency by inducing a huge number of possible candidate nodes/edges (especially for big graphs) that should be connected in the next step, but also induces a lot of nonrelevant results. Future research can focus on leveraging entity linking techniques that drastically reduce the number of candidate graph elements in an early stage.

Handling Ambiguous Queries. Disambiguation is considered as a subtask in a typical entity linking pipeline (Balog, 2018), where a specific entity is selected based on the context. This is usually challenging, given the short nature of keywords. The enhancement of entity disambiguation techniques would also help reducing the number of candidates and thus indirectly improving efficiency and effectiveness.

Incorporating User Intent. Predicting the user intent can help to better understand the query and thus provide guidance for an accurate mapping of keywords. Keyword search systems that work over KGs can, in addition, leverage the connected nature of the graphs to support techniques for user intent prediction.

Query Rewriting. The usage of retrieval-augmented generation (RAG) (Lewis et al., 2020) techniques (e.g., GraphRAG; Edge et al., 2024) can perform query rewriting to improve query understanding in case of ambiguous or incomplete keyword queries. This can be achieved by first expanding the query with external context (e.g., relevant document corpus or the KG itself) and then using large language models (LLMs) (Ozdemir, 2023) to rephrase the query in a more precise way.

9.2 Answer Retrieval

Improving Efficiency. As KGs grow in size and complexity, current keyword search algorithms may not be able to scale. To deal with that, the following directions could be investigated: distributed indexing, parallel processing, and graph summarization. To deal with that, techniques for graph summarization (Liu et al., 2018a) can be used. The latter should have the potential of reducing the size of the graph while preserving important information. However, only a few systems (Dalvi et al., 2008; Jiang et al., 2021; Sinha et al., 2018; Tran et al., 2009; Zhong & Liu, 2009) follow a summary-based approach, and this without evaluating the quality of the generated summaries. Therefore, there is still room for further investigations to deal with some common summarization challenges, for example, information loss. Other directions are the usage of distributed indexing and parallel processing. The latter is still also not sufficiently studied (Guan et al., 2018; Hao et al., 2015; Liu et al., 2016; Virgilio & Maccioni, 2014; Yang et al., 2019).

Leveraging LLMs. With the recent emergence of LLMs, a potential research direction could investigate the ability to automatically generate answers such as structured queries, for example, SPARQL given a keyword query and a target KG.

9.3 Answer Ranking

Incorporating Semantic Information. Existing surveys emphasize the general need of enabling semantic search over graph data (Wang & Aggarwal, 2010), and propose to, for example, leverage ontologies to improve efficiency and effectiveness (Yang et al., 2021). Furthermore, most of the existing techniques working over KGs still consider only structural or textual metrics to judge the relevance of a result to a certain query. Future works can focus on leveraging the semantic information encapsulated in KGs and propose ranking functions that go beyond structural properties of the graph using, for example, semantic similarity or KG embeddings (Wang et al., 2017a).

9.4 Answer Presentation

Existing works usually focus on proposing functioning systems aiming at increasing efficiency and effectiveness. However, one aspect is not sufficiently studied, namely, finding new and suitable ways to present the results to the end user and improve the browsing experience.

Supporting Explainability. The aim is to make the search results more transparent by explaining why a specific relevant result was generated. This will allow a better understanding of the relation between the different pieces of information corresponding to the different keywords. A potential future direction is to think about new functionalities to enable explainable information retrieval over KGs.

Enabling Exploratory Search. Users are not always familiar with the domain in question, and they do not always have a specific goal in mind (White & Roth, 2009). Therefore, they usually start with a tentative query and continue exploring to better understand a topic or discover new interesting insights and relations. Future research can focus on providing a range of features to support exploration and creative information-seeking behavior using simple keywords and KGs.

Generating Answers in Natural Language. Systems performing search over KGs typically retrieve entities or triples, which may not be easy to interpret for end users. RAG techniques and specifically GraphRAG can transform the retrieved graph elements into natural language responses that are also enriched using the KG as context.

9.5 Evaluating Search Performance

As already mentioned in Section 6, only a few systems were evaluated using existing test collections. On the other hand, we also notice a lack of standardized benchmarks dedicated to keyword search over KGs (Feddou et al., 2023; only three³⁸). Two of them should be adapted before usage since they were originally created to serve other tasks (question answering) or to work with other data formats (relational databases). Future research can focus on developing established evaluation datasets for keyword search over KGs.

9.6 Other Specific Cases

The literature review also revealed the existence of some additional niche areas with very few contributions³⁹, namely *probabilistic/uncertain graphs*, *distributed graphs*, *incomplete graphs*, or *temporal graph*. Future research can further investigate the application of keyword search over the previously mentioned special graphs.

Acknowledgements

We thank all members and partners of the involved projects: simpLEX, Canarèno, and KollOM-Fit. We would further like to thank Birgitta König-Ries for the guidance and feedback.

Funding

The authors disclosed receipt of the following financial support for the research, authorship, and/or publication of this article: This work has been funded by the German Federal Ministry of the Interior and Community and the Thuringian Ministry of Finance.

Declaration of Conflicting Interests

The authors declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

ORCID iDs

Leila Feddoul  <https://orcid.org/0000-0001-8896-8208>

Frank Löffler  <https://orcid.org/0000-0001-6643-6323>

Sirko Schindler  <https://orcid.org/0000-0002-0964-4457>

Notes

1. The schema is inspired by Wikidata.
2. <https://dblp.org/>
3. The keyword *quer* is deliberately used to include both singular and plural forms in the following queries.
4. <https://www.jabref.org/>
5. We did not include works with very specific use cases (e.g., working over multiple data sources or dealing with spatial keywords) or works performing keyword search over XML (Bray et al., 2008) or over text documents.
6. Those papers include works dealing with, for example, distributed graphs, incomplete graphs, or the usage of parallel processing. They will be just shortly mentioned as possible other special aspects in Section 5.
7. User profile refers usually to a collection of data that reflects user interest. In Sinha et al. (2018), for example, the user profile is a set of pattern graphs that were collected based on previous user interaction with the system (e.g., user selects a graph-shaped answer as relevant to a specific query). This way, query answers are ranked based on their similarity with the profile graphs.
8. In the case of retrieving subgraphs, a specific node that connects all the keywords (e.g., connecting element in Tran et al. 2009), is considered as equivalent to the tree root.
9. These approaches aim to diversify answers at the same time when searching the graphs and thus propose algorithms that are directly geared toward producing diverse results. This is different from the diversity-based ranking, since the latter aims at reducing redundancy of already generated answers using an algorithm that is not necessarily tuned to take this aspect into consideration.
10. All datasets are referenced in the corresponding paper summary in Section 8.
11. <https://www.imdb.com/>
12. <https://qald.aksw.org/>
13. Before adding a result to the heap, a duplicate tree detection (e.g., trees with the same undirected version) is conducted, and the tree with the highest relevance is kept. If the duplicate of the new result have already been output it is discarded even if its relevance is better.
14. <https://dblp.uni-trier.de/xml/>
15. <https://www.imdb.com/>
16. <https://grouplens.org/datasets/movielens/>
17. <https://www.dbis.informatik.uni-goettingen.de/Mondial/>
18. <https://www.tpc.org/tpch/>
19. <https://relational.fit.cvut.cz/dataset/IMDb>
20. <https://www.cs.toronto.edu/~oktie/linkedmdb/>, linkedmdb-latest-dump.zip
21. <http://km.aifb.uni-karlsruhe.de/ws/eon2006/ontoeval.zip>
22. <https://www.ebi.ac.uk/intact/home>
23. <https://snap.stanford.edu/data/com-DBLP.html>

24. <https://github.com/thunlp/Fast-TransX>
25. <http://swat.cse.lehigh.edu/projects/lubm/>
26. <http://wbsg.informatik.uni-mannheim.de/bizer/berlinsparqlbenchmark/>
27. <https://datasets.imdbws.com/>
28. <http://downloads.dbpedia.org/wiki-archive/data-set-37.html>
29. <https://sites.google.com/site/ontopiswc13/home/imdb-mo>
30. <https://sites.google.com/view/quira/>
31. <http://musicbrainz.org>
32. <https://github.com/ag-sc/QALD>
33. <http://dbtune.org/jamendo/>
34. <http://downloads.dbpedia.org/wiki-archive/dbpedia-dataset-version-2015-10.html>
35. <http://www.mpi-inf.mpg.de/yago>
36. <https://www.elastic.co/guide/en/elasticsearch/reference/current/index.html>
37. <https://old.datahub.io/dataset/knowledge-map>
38. Coffman, DBpedia-Entity v1/v2, and QALD
39. This claim is based on our findings while performing the literature review as stated in the last paragraph of Section 5.

References

- Agarwal, G., Kabra, G., & Chang, K. C. C. (2010). Towards rich query interpretation: Walking back and forth for mining query templates. In *WWW'10: Proceedings of the 19th international conference on world wide web* (pp. 1–10). Association for Computing Machinery. <https://doi.org/10.1145/1772690.1772692>.
- Agresti, A. (2010). *Analysis of ordinal categorical data* (vol. 656). John Wiley & Sons. <https://doi.org/10.1002/9780470594001>
- Bae, M., Kang, S., & Oh, S. (2014). Semantic similarity method for keyword query system on RDF. *Neurocomputing*, 146, 264–275. <https://doi.org/10.1016/j.neucom.2014.04.062>
- Bahmani, B., Chowdhury, A., & Goel, A. (2010). Fast incremental and personalized pagerank. *Proceedings of the VLDB Endowment*, 4(3), 173–184. <https://doi.org/10.14778/1929861.1929864>
- Balog, K. (2018). Entity linking. In *Entity-oriented search. The information retrieval series* (vol 39, pp. 147–188). Springer International Publishing. https://doi.org/10.1007/978-3-319-93935-3_5
- Balog, K., & Neumayer, R. (2013). A test collection for entity search in DBpedia. In G. J. F. Jones, P. Sheridan, D. Kelly, M. de Rijke, & T. Sakai (Eds.), *The 36th international ACM SIGIR conference on research and development in information retrieval, SIGIR'13, Dublin, Ireland—July 28–August 01, 2013* (pp. 737–740). ACM. <https://doi.org/10.1145/2484028.2484165>
- Bhalotia, G., Hulgeri, A., Nakhe, C., Chakrabarti, S., & Sudarshan, S. (2002). Keyword searching and browsing in databases using banks. In *Proceedings 18th international conference on data engineering* (pp. 431–440). IEEE. <https://doi.org/10.1109/ICDE.2002.994756>
- Bikakis, N., Giannopoulos, G., Liagouris, J., Skoutas, D., Dalamagas, T., & Sellis, T. K. (2013). RDivF: Diversifying keyword search on RDF graphs. In T. Aalberg, C. Papatheodorou, M. Dobrev, G. Tsakonas, & C. J. Farrugia (Eds.), *Research and advanced technology for digital libraries—international conference on theory and practice of digital libraries, TPD 2013, Valletta, Malta, September 22–26, 2013. Proceedings, lecture notes in computer science* (vol. 8092, pp. 413–416). Springer. https://doi.org/10.1007/978-3-642-40501-3_49
- Bizer, C., & Schultz, A. (2009). The Berlin SPARQL benchmark. *International Journal on Semantic Web and Information Systems*, 5(2), 1–24. <https://doi.org/10.4018/jswis.2009040101>
- Bordes, A., Usunier, N., Garcia-Durán, A., Weston, J., & Yakhnenko, O. (2013). Translating embeddings for modeling multi-relational data. In C. J. C. Burges, L. Bottou, Z. Ghahramani, & K. Q. Weinberger (Eds.), *Advances in neural information processing systems 26: 27th annual conference on neural information processing systems 2013. Proceedings of a meeting held December 5–8, 2013, Lake Tahoe, Nevada, United States* (pp. 2787–2795). Curran Associates Inc. <https://proceedings.neurips.cc/paper/2013/hash/1cecc7a77928ca8133fa24680a88d2f9-Abstract.html>
- Bray, T., Paoli, J., Sperberg-McQueen, C. M., Maler, E., & Yergeau, F. (2008). Extensible markup language (xml) 1.0 (5th ed.). W3C. <https://www.w3.org/TR/REC-xml>
- Brereton, P., Kitchenham, B. A., Budgen, D., Turner, M., & Khalil, M. (2007). Lessons from applying the systematic literature review process within the software engineering domain. *Journal of Systems and Software*, 80(4), 571–583. <https://doi.org/https://doi.org/10.1016/j.jss.2006.07.009>
- Bron, C., & Kerbosch, J. (1973). Algorithm 457: Finding all cliques of an undirected graph. *Communications of the ACM*, 16(9), 575–577. <https://doi.org/10.1145/362342.362367>
- Bryson, S., Davoudi, H., Golab, L., Kargar, M., Lytvyn, Y., Mierzejewski, P., Szlichta, J., & Zihayat, M. (2020). Robust keyword search in large attributed graphs. *Information Retrieval Journal*, 23(5), 502–524. <https://doi.org/10.1007/s10791-020-09379-9>

- Cai, Q., & Yates, A. (2013). Large-scale semantic parsing via schema matching and lexicon extension. In *Proceedings of the 51st annual meeting of the association for computational linguistics, ACL 2013, 4–9 August 2013, Sofia, Bulgaria, Volume 1: Long papers* (pp. 423–433). The Association for Computer Linguistics. <https://aclanthology.org/P13-1042/>
- Cantalalops, M. M., Sánchez-Alonso, S., & García-Barriocanal, E. (2019). A systematic literature review on wikidata. *Data Technologies and Applications*, 53(3), 250–268. <https://doi.org/10.1108/DTA-12-2018-0110>
- Cheng, G., Li, S., Zhang, K., & Li, C. (2020). Generating compact and relaxable answers to keyword queries over knowledge graphs. In J. Z. Pan, V. A. M. Tamma, C. d’Amato, K. Janowicz, B. Fu, A. Polleres, O. Seneviratne, & L. Kagal (Eds.), *The semantic web—ISWC 2020—19th international semantic web conference, Athens, Greece, November 2–6, 2020, proceedings, part I, lecture notes in computer science* (vol. 12506, pp. 110–127). Springer. https://doi.org/10.1007/978-3-030-62419-4_7
- Cheng, G., Shao, F., & Qu, Y. (2017). An empirical evaluation of techniques for ranking semantic associations. *IEEE Transactions on Knowledge and Data Engineering*, 29(11), 2388–2401. <https://doi.org/10.1109/TKDE.2017.2735970>
- Cleverdon, C. (1967). The Cranfield tests on index language devices. In *ASLIB proceedings* (vol. 19, pp. 173–194). MCB UP Ltd.
- Coffman, J., & Weaver, A. C. (2014). An empirical performance evaluation of relational keyword search techniques. *IEEE Transactions on Knowledge and Data Engineering*, 26(1), 30–42. <https://doi.org/10.1109/TKDE.2012.228>
- Craswell, N. (2009). Mean reciprocal rank. In L. Liu & M.T. Özsu (Eds.), *Encyclopedia of database systems* (pp. 1703–1703). Springer. https://doi.org/10.1007/978-0-387-39940-9_488
- Dalvi, B. B., Kshirsagar, M., & Sudarshan, S. (2008). Keyword search on external memory data graphs. *Proceedings of the Very Large Data Bases Endowment (PVLDB)*, 1(1), 1189–1204. <https://doi.org/10.14778/1453856.1453982>
- Dijkstra, E. W. (1959). A note on two problems in connexion with graphs. *Numerische Mathematik*, 1(1), 269–271. <https://doi.org/10.1007/BF01386390>
- Ding, B., Yu, J. X., Wang, S., Qin, L., Zhang, X., & Lin, X. (2007). Finding top-*k* min-cost connected trees in databases. In R. Chirkova, A. Dogac, M. T. Özsu, & T. K. Sellis (Eds.), *Proceedings of the 23rd international conference on data engineering, ICDE 2007, The Marmara Hotel, Istanbul, Turkey, April 15–20, 2007* (pp. 836–845). IEEE Computer Society. <https://doi.org/10.1109/ICDE.2007.367929>
- Dosso, D., & Silvello, G. (2020). Search text to retrieve graphs: A scalable RDF keyword-based search system. *IEEE Access*, 8, 14089–14111. <https://doi.org/10.1109/ACCESS.2020.2966823>
- Dreyfus, S. E., & Wagner, R. A. (1971). The steiner problem in graphs. *Networks*, 1(3), 195–207. <https://doi.org/10.1002/net.3230010302>
- Edge, D., Trinh, H., Cheng, N., Bradley, J., Chao, A., Mody, A., Truitt, S., & Larson, J. (2024). From local to global: A graph RAG approach to query-focused summarization. <https://arxiv.org/abs/2404.16130>
- Elbassuoni, S., & Blanco, R. (2011). Keyword search over RDF graphs. In C. Macdonald, I. Ounis, & I. Ruthven (Eds.), *Proceedings of the 20th ACM conference on information and knowledge management, CIKM 2011, Glasgow, United Kingdom, October 24–28, 2011* (pp. 237–242). ACM. <https://doi.org/10.1145/2063576.2063615>
- Fagin, R., Lotem, A., & Naor, M. (2003). Optimal aggregation algorithms for middleware. *Journal of Computer and System Sciences*, 66(4), 614–656. [https://doi.org/10.1016/S0022-0000\(03\)00026-6](https://doi.org/10.1016/S0022-0000(03)00026-6)
- Feddoul, L. (2020). Semantics-driven keyword search over knowledge graphs. In H. Alani & E. Simperl (Eds.), *Proceedings of the doctoral consortium at ISWC 2020 co-located with 19th international semantic web conference (ISWC 2020), Athens, Greece, November 3rd, 2020, CEUR workshop proceedings* (vol. 2798, pp. 17–24). CEUR-WS.org. <https://ceur-ws.org/Vol-2798/paper3.pdf>
- Feddoul, L., Löffler, F., & Schindler, S. (2023). Keysearchwiki: An automatically generated dataset for keyword search over Wikidata. In L. Kaffee, S. Razniewski, K. Alghamdi, & H. Arnaout (Eds.), *Proceedings of the Wikidata workshop 2023 co-located with 22nd international semantic web conference (ISWC 2023), Athens, Greece, November 13, 2023, CEUR workshop proceedings* (vol. 3640). CEUR-WS.org. <https://ceur-ws.org/Vol-3640/paper6.pdf>
- Fellbaum, C., & George, A. (1998). *WordNet: An electronic lexical database*. MIT Press.
- García, G., Izquierdo, Y., Menendez, E., Dartayre, F., & Casanova, M. A. (2017). RDF keyword-based query technology meets a real-world dataset. In V. Markl, S. Orlando, B. Mitschang, P. Andritsos, K. Sattler, & S. Breß (Eds.), *Proceedings of the 20th international conference on extending database technology, EDBT 2017, Venice, Italy, March 21–24, 2017* (pp. 656–667). OpenProceedings.org. <https://doi.org/10.5441/002/edbt.2017.86>
- Ghanbarpour, A., & Naderi, H. (2019). Survey on ranking functions in keyword search over graph-structured data. *Journal of Universal Computer Science*, 25(4), 361–389. http://www.jucs.org/jucs_25_4/survey_on_ranking_functions
- Ghanbarpour, A., Niknafs, K., & Naderi, H. (2020). Efficient keyword search over graph-structured data based on minimal covered r- cliques. *Frontiers of Information Technology & Electronic Engineering*, 21(3), 448–464. <https://doi.org/10.1631/FITEE.1800133>
- Gkirtzou, K., Karozos, K., Vassalos, V., & Dalamagas, T. (2015). Keywords-to-sparql translation for RDF data search and exploration. In S. Kapidakis, C. Mazurek, & M. Werla (Eds.), *Research and advanced technology for digital libraries—19th international conference on theory and practice of digital libraries, TPDL 2015, Poznań, Poland, September 14–18, 2015. Proceedings, lecture notes in computer science* (vol. 9316, pp. 111–123). Springer. https://doi.org/10.1007/978-3-319-24592-8_9

- Golenberg, K., Kimelfeld, B., & Sagiv, Y. (2008). Keyword proximity search in complex data graphs. In J. T. Wang (Ed.), *Proceedings of the ACM SIGMOD international conference on management of data, SIGMOD 2008, Vancouver, BC, Canada, June 10–12, 2008* (pp. 927–940). ACM. <https://doi.org/10.1145/1376616.1376708>
- Guan, J., Wang, J., & Yu, L. (2018). Multi-keyword parallel search algorithm for streaming RDF data. In Z. Xu, X. Gao, Q. Miao, Y. Zhang, & J. Bu (Eds.), *Big Data—6th CCF conference, Big Data 2018, Xi'an, China, October 11–13, 2018, proceedings, communications in computer and information science* (vol. 945, pp. 494–511). Springer. https://doi.org/10.1007/978-981-13-2922-7_33
- Güntzer, U., Balke, W., & Kießling, W. (2001). Towards efficient multi-feature queries in heterogeneous environments. In *2001 international symposium on information technology (ITCC 2001), 2–4 April 2001, Las Vegas, NV, USA* (pp. 622–628). IEEE Computer Society. <https://doi.org/10.1109/ITCC.2001.918866>
- Han, S., Zou, L., Yu, J. X., & Zhao, D. (2017). Keyword search on RDF graphs—a query graph assembly approach. In E. Lim, M. Winslett, M. Sanderson, A. W. Fu, J. Sun, J. S. Culpepper, E. Lo, J. C. Ho, D. Donato, R. Agrawal, Y. Zheng, C. Castillo, A. Sun, V. S. Tseng, & C. Li (Eds.), *Proceedings of the 2017 ACM on conference on information and knowledge management, CIKM 2017, Singapore, November 6–10, 2017* (pp. 227–236). ACM. <https://doi.org/10.1145/3132847.3132957>
- Hao, Y., Cao, H., Qi, Y., Hu, C., Brahma, S., & Han, J. (2015). Efficient keyword search on graphs using mapreduce. In *2015 IEEE international conference on Big Data (IEEE BigData 2015), Santa Clara, CA, USA, October 29–November 1, 2015* (pp. 2871–2873). IEEE Computer Society. <https://doi.org/10.1109/BigData.2015.7364106>
- Hao, Y., Cao, X., Sheng, Y., Fang, Y., & Wang, W. (2021). KS-GNN: keywords search over incomplete graphs via graphs neural network. In M. Ranzato, A. Beygelzimer, Y. N. Dauphin, P. Liang, & J. W. Vaughan (Eds.), *Advances in neural information processing systems 34: Annual conference on neural information processing systems 2021, NeurIPS 2021, December 6–14, 2021, virtual* (pp. 1700–1712). Curran Associates Inc. <https://proceedings.neurips.cc/paper/2021/hash/0d7363894acdee742caf7fe4e97c4d49-Abstract.html>
- Hasibi, F., Nikolaev, F., Xiong, C., Balog, K., Bratsberg, S. E., Kotov, A., & Callan, J. (2017). Dbpedia-entity v2: A test collection for entity search. In N. Kando, T. Sakai, H. Joho, H. Li, A. P. de Vries, & R. W. White (Eds.), *Proceedings of the 40th international ACM SIGIR conference on research and development in information retrieval, Shinjuku, Tokyo, Japan, August 7–11, 2017* (pp. 1265–1268). ACM. <https://doi.org/10.1145/3077136.3080751>
- He, H., Wang, H., Yang, J., & Yu, P. S. (2007). BLINKS: ranked keyword searches on graphs. In C. Y. Chan, B. C. Ooi, & A. Zhou (Eds.), *Proceedings of the ACM SIGMOD international conference on management of data, Beijing, China, June 12–14, 2007* (pp. 305–316). ACM. <https://doi.org/10.1145/1247480.1247516>
- Hristidis, V., Gravano, L., & Papakonstantinou, Y. (2003). Efficient ir-style keyword search over relational databases. In J. C. Freytag, P. C. Lockemann, S. Abiteboul, M. J. Carey, P. G. Selinger, & A. Heuer (Eds.), *Proceedings of 29th international conference on very large data bases, VLDB 2003, Berlin, Germany, September 9–12, 2003* (pp. 850–861). Morgan Kaufmann. <https://doi.org/10.1016/B978-012722442-8/50080-X>
- Hristidis, V., & Papakonstantinou, Y. (2002). DISCOVER: keyword search in relational databases. In *Proceedings of 28th international conference on very large data bases, VLDB 2002, Hong Kong, August 20–23, 2002* (pp. 670–681). Morgan Kaufmann. <https://doi.org/10.1016/B978-155860869-6/50065-2>
- Izquierdo, Y., García, G. M., Menendez, E., Leme, L. A. P. P., Neves, A. B., Lemos, M., Finamore, A. C., Oliveira, C., & Casanova, M. A. (2021). Keyword search over schema-less RDF datasets by SPARQL query compilation. *Information Systems*, 102, 101814. <https://doi.org/10.1016/j.is.2021.101814>
- Järvelin, K., & Kekäläinen, J. (2002). Cumulated gain-based evaluation of ir techniques. *ACM Transactions on Information Systems*, 20(4), 422–446. <https://doi.org/10.1145/582415.582418>
- Jiang, J., Choi, B., Xu, J., & Bhowmick, S. S. (2021). A generic ontology framework for indexing keyword search on massive graphs. *IEEE Transactions on Knowledge and Data Engineering*, 33(6), 2322–2336. <https://doi.org/10.1109/TKDE.2019.2956535>
- Kacholia, V., Pandit, S., Chakrabarti, S., Sudarshan, S., Desai, R., & Karambelkar, H. (2005). Bidirectional expansion for keyword search on graph databases. In K. Böhm, C. S. Jensen, L. M. Haas, M. L. Kersten, P. Larson, & B. C. Ooi (Eds.), *Proceedings of the 31st international conference on very large data bases, Trondheim, Norway, August 30–September 2, 2005* (pp. 505–516). ACM. <http://www.vldb.org/archives/website/2005/program/paper/wed/p505-kacholia.pdf>
- Kadilierakis, G., Fafalios, P., Papadakis, P., & Tzitzikas, Y. (2020). Keyword search over RDF using document-centric information retrieval systems. In A. Harth, S. Kirrane, A. N. Ngomo, H. Paulheim, A. Rula, A. L. Gentile, P. Haase, & M. Cochez (Eds.), *The semantic web—17th international conference, ESWC 2020, Heraklion, Crete, Greece, May 31–June 4, 2020, proceedings, lecture notes in computer science* (vol. 12123, pp. 121–137). Springer. https://doi.org/10.1007/978-3-030-49461-2_8
- Kargar, M., & An, A. (2011). Keyword search in graphs: Finding r-cliques. *Proceedings of the VLDB Endowment*, 4(10), 681–692. <https://doi.org/10.14778/2021017.2021025>
- Kargar, M., & An, A. (2015). Finding top-k, r-cliques for keyword search from graphs in polynomial delay. *Knowledge and Information Systems*, 43(2), 249–280. <https://doi.org/10.1007/s10115-014-0736-0>

- Kargar, M., An, A., Cercone, N., Godfrey, P., Szlichta, J., & Yu, X. (2015). Meaningful keyword search in relational databases with large and complex schema. In J. Gehrke, W. Lehner, K. Shim, S. K. Cha, & G. M. Lohman (Eds.), *31st IEEE international conference on data engineering, ICDE 2015, Seoul, South Korea, April 13–17, 2015* (pp. 411–422). IEEE Computer Society. <https://doi.org/10.1109/ICDE.2015.7113302>
- Kargar, M., An, A., & Yu, X. (2013). *Duplication free and minimal keyword search in large graphs* (Technical Report CSE-2013-02). York University.
- Kargar, M., An, A., & Yu, X. (2014). Efficient duplication free and minimal keyword search in graphs. *IEEE Transactions on Knowledge and Data Engineering*, 26(7), 1657–1669. <https://doi.org/10.1109/TKDE.2013.85>
- Katerenchuk, D., & Rosenberg, A. (2016). RankDCG: Rank-ordering evaluation measure. In N. Calzolari, K. Choukri, T. Declerck, S. Goggi, M. Grobelnik, B. Maegaard, J. Mariani, H. Mazo, A. Moreno, J. Odijk, & S. Piperidis (Eds.), *Proceedings of the tenth international conference on language resources and evaluation LREC 2016, Portorož, Slovenia, May 23–28, 2016* (pp. 3675–3680). European Language Resources Association (ELRA).
- Kharrat, M., Jedidi, A., & Gargouri, F. (2016). SPARQL query generation based on RDF graph. In A. L. N. Fred, J. L. G. Dietz, D. Aveiro, K. Liu, J. Bernardino, & J. Filipe (Eds.), *Proceedings of the 8th international joint conference on knowledge discovery, knowledge engineering and knowledge management (IC3K 2016)—volume 1: KDIR, Porto—Portugal, November 9–11, 2016* (pp. 450–455). SciTePress. <https://doi.org/10.5220/0006091904500455>
- Kimelfeld, B., & Sagiv, Y. (2008). Efficiently enumerating results of keyword search over data graphs. *Information Systems*, 33(4–5), 335–359. <https://doi.org/10.1016/j.is.2008.01.002>
- Kitchenham, B. A., & Charters, S. (2007). *Guidelines for performing systematic literature reviews in software engineering* (Technical Report EBSE 2007-001). Keele University and Durham University Joint Report. https://www.elsevier.com/_data/promis_misc/525444systematicreviewsguide.pdf
- Kleinberg, J. M. (1999). Authoritative sources in a hyperlinked environment. *Journal of the ACM*, 46(5), 604–632. <https://doi.org/10.1145/324133.324140>
- Lawler, E. L. (1972). A procedure for computing the k best solutions to discrete optimization problems and its application to the shortest path problem. *Management Science*, 18(7), 401–405. <https://doi.org/10.1287/mnsc.18.7.401>
- Le, W., Li, F., Kementsidis, A., & Duan, S. (2014). Scalable keyword search on large RDF data. *IEEE Transactions on Knowledge and Data Engineering*, 26(11), 2774–2788. <https://doi.org/10.1109/TKDE.2014.2302294>
- Lehmann, J., Isele, R., Jakob, M., Jentzsch, A., Kontokostas, D., Mendes, P. N., Hellmann, S., Morsey, M., van Kleef, P., Auer, S., & Bizer, C. (2015). DBpedia—A large-scale, multilingual knowledge base extracted from wikipedia. *Semantic Web*, 6(2), 167–195. <https://doi.org/10.3233/SW-140134>
- Leskovec, J., Rajaraman, A., & Ullman, J. D. (2020). *Mining of massive datasets*. 3rd ed. Cambridge University Press. <https://doi.org/10.1017/9781108684163>
- Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, Wt., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. In H. Larochelle, M. Ranzato, R. Hadsell, M. Balcan, & H. Lin (Eds.), *Advances in neural information processing systems* (vol. 33, pp. 9459–9474). Curran Associates, Inc. https://proceedings.neurips.cc/paper_files/paper/2020/file/6b493230205f780e1bc26945df7481e5-Paper.pdf
- Li, G., Ooi, B. C., Feng, J., Wang, J., & Zhou, L. (2008). EASE: an effective 3-in-1 keyword search method for unstructured, semi-structured and structured data. In J. T. Wang (Ed.), *Proceedings of the ACM SIGMOD international conference on management of data, SIGMOD 2008, Vancouver, BC, Canada, June 10–12, 2008* (pp. 903–914). ACM. <https://doi.org/10.1145/1376616.1376706>
- Li, H. Y., & Qu, Y. Z. (2011). Kreag: Keyword query approach over RDF data based on entity-triple association graph. *Jisuanji Xuebao (Chinese Journal of Computers)*, 34(5), 825–835.
- Li, N., Yan, X., Wen, Z., & Khan, A. (2012). Density index and proximity search in large graphs. In X. Chen, G. Lebanon, H. Wang, & M. J. Zaki (Eds.), *21st ACM international conference on information and knowledge management, CIKM'12, Maui, HI, USA, October 29–November 2, 2012* (pp. 235–244). ACM. <https://doi.org/10.1145/2396761.2396794>
- Li, R., Qin, L., Yu, J. X., & Mao, R. (2016). Efficient and progressive group steiner tree search. In F. Özcan, G. Koutrika, & S. Madden (Eds.), *Proceedings of the 2016 international conference on management of data, SIGMOD conference 2016, San Francisco, CA, USA, June 26–July 1, 2016* (pp. 91–106). ACM. <https://doi.org/10.1145/2882903.2915217>
- Lian, X., Chen, L., & Huang, Z. (2015). Keyword search over probabilistic RDF graphs. *IEEE Transactions on Knowledge and Data Engineering*, 27(5), 1246–1260. <https://doi.org/10.1109/TKDE.2014.2365791>
- Lin, F., & Sandkuhl, K. (2008). A survey of exploiting WordNet in ontology matching. In M. Bramer (Ed.), *Artificial intelligence in theory and practice II, IFIP 20th world computer congress, TC 12: IFIP AI 2008 Stream, September 7–10, 2008, Milano, Italy, IFIP* (vol. 276, pp. 341–350). Springer. https://doi.org/10.1007/978-0-387-09695-7_33
- Liu, C., Yao, L., Li, J., Zhou, R., & He, Z. (2016). Finding smallest k-compact tree set for keyword queries on graphs using mapreduce. *World Wide Web*, 19(3), 499–518. <https://doi.org/10.1007/s11280-015-0337-1>

- Liu, T. Y. (2009). Learning to rank for information retrieval. *Foundations and Trends® in Information Retrieval*, 3(3), 225–331. <https://doi.org/10.1561/15000000016>
- Liu, Y., Safavi, T., Dighe, A., & Koutra, D. (2018a). Graph summarization methods and applications: A survey. *ACM Computing Surveys (CSUR)*, 51(3), 1–34. <https://doi.org/10.1145/3186727>
- Liu, Z., Wang, C., & Chen, Y. (2018b). Keyword search on temporal graphs. In *34th IEEE international conference on data engineering, ICDE 2018, Paris, France, April 16–19, 2018* (pp. 1807–1808). IEEE Computer Society. <https://doi.org/10.1109/ICDE.2018.00261>
- Lu, X., Theodoratos, D., & Dimitriou, A. (2019). Leveraging pattern mining techniques for efficient keyword search on data graphs. In L. H. U, J. Yang, Y. Cai, K. Karlapalem, A. Liu, & X. Huang (Eds.), *Web information systems engineering—WISE 2019 workshop, demo, and tutorial, Hong Kong and Macau, China, January 19–22, 2020, revised selected papers, communications in computer and information science* (vol. 1155, pp. 98–114). Springer. https://doi.org/10.1007/978-981-15-3281-8_10
- Manning, C. D., Surdeanu, M., Bauer, J., Finkel, J. R., Bethard, S., & McClosky, D. (2014). The Stanford CoreNLP natural language processing toolkit. In *Proceedings of the 52nd annual meeting of the association for computational linguistics, ACL 2014, June 22–27, 2014, Baltimore, MD, USA, system demonstrations* (pp. 55–60). The Association for Computer Linguistics. <https://doi.org/10.3115/v1/p14-5010>
- Mass, Y., & Sagiv, Y. (2016). Virtual documents and answer priors in keyword search over data graphs. In T. Palpanas & K. Stefanidis (Eds.), *Proceedings of the workshops of the EDBT/ICDT 2016 joint conference, EDBT/ICDT workshops 2016, Bordeaux, France, March 15, 2016, CEUR workshop proceedings* (vol. 1558). CEUR-WS.org. <http://ceur-ws.org/Vol-1558/paper20.pdf>
- McBride, B. (2001). Jena: Implementing the RDF model and syntax specification. In S. Decker, D. A. Fensel, A. P. Sheth, & S. Staab (Eds.), *Proceedings of the second international workshop on the semantic web—SemWeb’2001, Hongkong, China, May 1, 2001, CEUR workshop proceedings* (vol. 40). CEUR-WS.org. <http://CEUR-WS.org/Vol-40/mcbride.pdf>
- Menendez, E. S., Casanova, M. A., Leme, L. A. P. P., & Boughanem, M. (2019). Novel node importance measures to improve keyword search over RDF graphs. In S. Hartmann, J. Küng, S. Chakravarthy, G. Anderst-Kotsis, A. M. Tjoa, & I. Khalil (Eds.), *Database and expert systems applications—30th international conference, DEXA 2019, Linz, Austria, August 26–29, 2019, proceedings, part II, lecture notes in computer science* (vol. 11707, pp. 143–158). Springer. https://doi.org/10.1007/978-3-030-27618-8_11
- Metzler, D., & Croft, W. B. (2005). A Markov random field model for term dependencies. In R. A. Baeza-Yates, N. Ziviani, G. Marchionini, A. Moffat, & J. Tait (Eds.), *SIGIR 2005: Proceedings of the 28th annual international ACM SIGIR conference on research and development in information retrieval, Salvador, Brazil, August 15–19, 2005* (pp. 472–479). ACM. <https://doi.org/10.1145/1076034.1076115>
- Navarro, F. A., Martinez-Costa, C., & Fernández-Breis, J. T. (2021). Semankey: A semantics-driven approach for querying RDF repositories using keywords. *IEEE Access*, 9, 91282–91302. <https://doi.org/10.1109/ACCESS.2021.3091413>
- Neves, A. B., Leme, L. A. P. P., Izquierdo, Y. T., & Casanova, M. A. (2021). Automatic construction of benchmarks for RDF keyword search systems evaluation. In J. Filipe, M. Smialek, A. Brodsky, & S. Hammoudi (Eds.), *Proceedings of the 23rd international conference on enterprise information systems, ICEIS 2021, Online Streaming, April 26–28, 2021* (vol. 1, pp. 126–137). SCITEPRESS. <https://doi.org/10.5220/0010519401260137>
- Nikas, C., Kadilierakis, G., Fafalios, P., & Tzitzikas, Y. (2020). Keyword search over RDF: Is a single perspective enough?. *Big Data and Cognitive Computing*, 4(3), 22. <https://doi.org/10.3390/bdcc4030022>
- Ouksili, H., Kedad, Z., Lopes, S., & Nugier, S. (2017). Using patterns for keyword search in RDF graphs. In Y. E. Ioannidis, J. Stoyanovich, & G. Orsi (Eds.), *Proceedings of the workshops of the EDBT/ICDT 2017 joint conference (EDBT/ICDT 2017), Venice, Italy, March 21–24, 2017, CEUR workshop proceedings* (vol. 1810). CEUR-WS.org. http://ceur-ws.org/Vol-1810/GraphQ_paper_08.pdf
- Ozdemir, S. (2023). *Quick start guide to large language models: Strategies and best practices for using ChatGPT and other LLMs*. Addison-Wesley data and analytics series. Pearson Education (US).
- Page, L., Brin, S., Motwani, R., & Winograd, T. (1999). *The pagerank citation ranking: Bringing order to the web* (Technical Report 1999-66). Stanford InfoLab. <http://ilpubs.stanford.edu:8090/422/>. Previous number = SIDL-WP-1999-0120.
- Park, C. (2018). Effective keyword search on graph data using limited root redundancy of answer trees. *International Journal of Web Information Systems*, 14(3), 299–316. <https://doi.org/10.1108/IJWIS-10-2017-0070>
- Park, C., & Lim, S. (2015). Efficient processing of keyword queries over graph databases for finding effective answers. *Information Processing & Management*, 51(1), 42–57. <https://doi.org/10.1016/j.ipm.2014.08.002>
- Qin, L., Yu, J. X., Chang, L., & Tao, Y. (2009). Querying communities in relational databases. In Y. E. Ioannidis, D. L. Lee, & R. T. Ng (Eds.), *Proceedings of the 25th international conference on data engineering, ICDE 2009, March 29, 2009–April 2, 2009, Shanghai, China* (pp. 724–735). IEEE Computer Society. <https://doi.org/10.1109/ICDE.2009.67>
- Rihany, M., Kedad, Z., & Lopes, S. (2018). Keyword search over RDF graphs using WordNet. In M. Hojeij, B. Finance, Y. Taher, K. Zeitouni, R. Haque, & M. Dbouk (Eds.), *Proceedings of the 1st international conference on Big Data and cyber-security intelligence, BDCSIntell 2018, Hadath, Lebanon, December 13–15, 2018, CEUR workshop proceedings* (vol. 2343, pp. 75–82). CEUR-WS.org. <http://ceur-ws.org/Vol-2343/paper15.pdf>

- Ristoski, P., Rosati, J., Noia, T. D., Leone, R. D., & Paulheim, H. (2019). RDF2Vec: RDF graph embeddings and their applications. *Semantic Web*, 10(4), 721–752. <https://doi.org/10.3233/SW-180317>
- Robertson, S., & Zaragoza, H. (2009). The probabilistic relevance framework: BM25 and beyond. *Foundations and Trends in Information Retrieval*, 3(4), 333–389. <https://doi.org/10.1561/15000000019>
- Schütze, H., Manning, C. D., & Raghavan, P. (2008). *Introduction to information retrieval* (vol. 39). Cambridge University Press.
- Shi, Y., Cheng, G., Tran, T., Kharlamov, E., & Shen, Y. (2021). Efficient computation of semantically cohesive subgraphs for keyword-based knowledge graph exploration. In J. Leskovec, M. Grobelnik, M. Najork, J. Tang, & L. Zia (Eds.), *WWW '21: The web conference 2021, virtual event/Ljubljana, Slovenia, April 19–23, 2021* (pp. 1410–1421). ACM/IW3C2. <https://doi.org/10.1145/3442381.3449900>
- Sinha, S. B., Lu, X., & Theodoratos, D. (2018). Personalized keyword search on large RDF graphs based on pattern graph similarity. In B. C. Desai, S. Flesca, E. Zumpato, E. Masciari, & L. Caroprese (Eds.), *Proceedings of the 22nd international database engineering & applications symposium, IDEAS 2018, Villa San Giovanni, Italy, June 18–20, 2018* (pp. 12–21). ACM. <https://doi.org/10.1145/3216122.3216167>
- Sithisarn, S. (2014). A semantic keyword search based on the bidirectional fix root query graph construction algorithm. In T. Supnithi, T. Yamaguchi, J. Z. Pan, V. Wuwongse, & M. Buranarach (Eds.), *Semantic technology—4th joint international conference, JIST 2014, Chiang Mai, Thailand, November 9–11, 2014. Revised selected papers, lecture notes in computer science* (vol. 8943, pp. 387–394). Springer. https://doi.org/10.1007/978-3-319-15615-6_29
- Sithisarn, S., Lau, L., & Dew, P. M. (2011). Semantic keyword search for expert witness discovery. In *2011 international conference on semantic technology and information retrieval* (pp. 18–25). IEEE. <https://doi.org/10.1109/STAIR.2011.5995759>
- Tomita, E., Tanaka, A., & Takahashi, H. (2006). The worst-case time complexity for generating all maximal cliques and computational experiments. *Theoretical Computer Science*, 363(1), 28–42. <https://doi.org/10.1016/j.tcs.2006.06.015>
- Tran, T., Wang, H., Rudolph, S., & Cimiano, P. (2009). Top-k exploration of query candidates for efficient keyword search on graph-shaped (RDF) data. In Y. E. Ioannidis, D. L. Lee, & R. T. Ng (Eds.), *Proceedings of the 25th international conference on data engineering, ICDE 2009, March 29, 2009–April 2, 2009, Shanghai, China* (pp. 405–416). IEEE Computer Society. <https://doi.org/10.1109/ICDE.2009.119>
- Unger, C., Ngomo, A. N., & Cabrio, E. (2016). 6th open challenge on question answering over linked data (QALD-6). In H. Sack, S. Dietze, A. Tordai, & C. Lange (Eds.), *Semantic web challenges—third SemWebEval challenge at ESWC 2016, Heraklion, Crete, Greece, May 29–June 2, 2016, revised selected papers, communications in computer and information science* (vol. 641, pp. 171–177). Springer. https://doi.org/10.1007/978-3-319-46565-4_13
- Valiente, G. (2021). *Algorithms on trees and graphs: With Python code*. Texts in computer science. Springer International Publishing.
- Virgilio, R. D., & Maccioni, A. (2014). Distributed keyword search over RDF via mapreduce. In V. Presutti, C. d'Amato, F. Gandon, M. d'Aquin, S. Staab, & A. Tordai (Eds.), *The semantic web: Trends and challenges—11th international conference, ESWC 2014, Anissaras, Crete, Greece, May 25–29, 2014. Proceedings, lecture notes in computer science* (vol. 8465, pp. 208–223). Springer. https://doi.org/10.1007/978-3-319-07443-6_15
- Wang, H., & Aggarwal, C. C. (2010). A survey of algorithms for keyword search on graph data. In CC. Aggarwal & H. Wang (Eds.), *Managing and mining graph data, advances in database systems* (vol. 40, pp. 249–273). Springer. https://doi.org/10.1007/978-1-4419-6045-0_8
- Wang, Q., Mao, Z., Wang, B., & Guo, L. (2017a). Knowledge graph embedding: A survey of approaches and applications. *IEEE Transactions on Knowledge and Data Engineering*, 29(12), 2724–2743. <https://doi.org/10.1109/TKDE.2017.2754499>
- Wang, Y., Wang, K., Fu, A. W., & Wong, R. C. (2015). Keylabel algorithms for keyword search in large graphs. In *2015 IEEE international conference on Big Data (IEEE BigData 2015), Santa Clara, CA, USA, October 29–November 1, 2015* (pp. 857–864). IEEE Computer Society. <https://doi.org/10.1109/BigData.2015.7363833>
- Wang, Y., Zhong, M., Zhu, Y., Li, X., & Qian, T. (2017b). Diversified top-k keyword query interpretation on knowledge graphs. In L. Chen, C. S. Jensen, C. Shahabi, X. Yang, & X. Lian (Eds.), *Web and Big Data—first international joint conference, APWeb-WAIM 2017, Beijing, China, July 7–9, 2017, proceedings, part I, lecture notes in computer science* (vol. 10366, pp. 541–555). Springer. https://doi.org/10.1007/978-3-319-63579-8_41
- White, R. W., & Roth, R. A. (2009). *Exploratory search: Beyond the query-response paradigm*. Synthesis lectures on information concepts, retrieval, and services. Morgan & Claypool Publishers. <https://doi.org/10.2200/S00174ED1V01Y200901ICR003>
- Yang, J., Yao, W., & Zhang, W. (2021). Keyword search on large graphs: A survey. *Data Science and Engineering*, 6(2), 142–162. <https://doi.org/10.1007/s41019-021-00154-4>
- Yang, Y., Agrawal, D., Jagadish, H. V., Tung, A. K. H., & Wu, S. (2019). An efficient parallel keyword search engine on knowledge graphs. In *35th IEEE international conference on data engineering, ICDE 2019, Macao, China, April 8–11, 2019* (pp. 338–349). IEEE. <https://doi.org/10.1109/ICDE.2019.00038>
- Yu, J. X., Qin, L., & Chang, L. (2010). Keyword search in relational databases: A survey. *IEEE Data Engineering Bulletin*, 33(1), 67–78. <http://sites.computer.org/debull/A10mar/you-paper.pdf>

- Yuan, Y., Lian, X., Chen, L., Yu, J. X., Wang, G., & Sun, Y. (2017). Keyword search over distributed graphs with compressed signature. *IEEE Transactions on Knowledge and Data Engineering*, 29(6), 1212–1225. <https://doi.org/10.1109/TKDE.2017.2656079>
- Zaki, M. J. (2005). Efficiently mining frequent embedded unordered trees. *Fundamenta Informaticae*, 66(1–2), 33–52. <http://content.iospress.com/articles/fundamenta-informaticae/fi66-1-2-03>
- Zhang, E., & Zhang, Y. (2009). F-measure. In *Encyclopedia of database systems* (pp. 1147–1147). Springer. https://doi.org/10.1007/978-0-387-39940-9_483
- Zhang, H., Dong, B., Feng, B., & Wei, B. (2021). A keyword query approach based on community structure of RDF entity graph. In *IEEE 45th annual computers, software, and applications conference, COMPSAC 2021, Madrid, Spain, July 12–16, 2021* (pp. 1143–1148). IEEE. <https://doi.org/10.1109/COMPSAC51774.2021.00157>
- Zhang, Z., Xia, D., & Xie, X. (2014). Keyword search on graphs based on content and structure. In Z. Cai, C. Wang, S. Cheng, H. Wang, & H. Gao (Eds.), *Wireless algorithms, systems, and applications—9th international conference, WASA 2014, Harbin, China, June 23–25, 2014. Proceedings, lecture notes in computer science* (vol. 8491, pp. 760–772). Springer. https://doi.org/10.1007/978-3-319-07782-6_68
- Zhang, Z., Yu, J. X., Wang, G., Yuan, Y., & Chen, L. (2022). Key-core: Cohesive keyword subgraph exploration in large graphs. *World Wide Web*, 25(2), 831–856. <https://doi.org/10.1007/s11280-021-00926-y>
- Zheng, Z., Ding, Y., Wang, Z., & Wang, Z. (2016). A novel method of keyword query for RDF data based on bipartite graph. In *22nd IEEE international conference on parallel and distributed systems, ICPADS 2016, Wuhan, China, December 13–16, 2016* (pp. 466–473). IEEE Computer Society. <https://doi.org/10.1109/ICPADS.2016.0069>
- Zhong, M., & Liu, M. (2009). Efficient keyword proximity search using a frontier-reduce strategy based on d -distance graph index. In B. C. Desai, D. Saccà, & S. Greco (Eds.), *International database engineering and applications symposium (IDEAS 2009), September 16–18, 2009, Cetraro, Calabria, Italy, ACM international conference proceeding series* (pp. 206–216). ACM. <https://doi.org/10.1145/1620432.1620453>
- Zhong, M., Wang, Y., & Zhu, Y. (2018). Coverage-oriented diversification of keyword search results on graphs. In J. Pei, Y. Manolopoulos, S. W. Sadiq, & J. Li (Eds.), *Database systems for advanced applications—23rd international conference, DASFAA 2018, Gold Coast, QLD, Australia, May 21–24, 2018, proceedings, Part II, lecture notes in computer science* (vol. 10828, pp. 166–183). Springer. https://doi.org/10.1007/978-3-319-91458-9_10