

Exploring Relations and Graphs for Information Retrieval

Chris Kamphuis

Institute for Computing
and Information Sciences

**RADBOD
UNIVERSITY
PRESS**

Radboud
Dissertation
Series

Exploring Relations and Graphs for Information Retrieval

Chris Frans Henri Kamphuis

This work is part of the research program Commit2Data with project number 628.011.001 (SQIREL-GRAPHS), which is (partly) financed by the Netherlands Organisation for Scientific Research (NWO).

SIKS Dissertation Series No. 2025-37

The research reported in this thesis has been carried out under the auspices of SIKS, the Dutch Research School for Information and Knowledge Systems.



Exploring Relations and Graphs for Information Retrieval

Chris Frans Henri Kamphuis

Radboud Dissertation Series

ISSN: 2950-2772 (Online); 2950-2780 (Print)

Published by RADBOUD UNIVERSITY PRESS

Postbus 9100, 6500 HA Nijmegen, The Netherlands

www.radbouduniversitypress.nl

Design: Chris Frans Henri Kamphuis

Cover: Proefschrift AIO | Guntra Laivacuma

Printing: DPN Rikken/Pumbo

ISBN: 9789465151298

DOI: 10.54195/9789465151298

Free download at: <https://doi.org/10.54195/9789465151298>

© 2025 Chris Frans Henri Kamphuis

**RADBOUD
UNIVERSITY
PRESS**

This is an Open Access book published under the terms of Creative Commons Attribution-Noncommercial-NoDerivatives International license (CC BY-NC-ND 4.0). This license allows reusers to copy and distribute the material in any medium or format in unadapted form only, for noncommercial purposes only, and only so long as attribution is given to the creator, see <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

Exploring Relations and Graphs for Information Retrieval

Proefschrift

ter verkrijging van de graad van doctor
aan de Radboud Universiteit Nijmegen
op gezag van de rector magnificus prof. dr. J.M. Sanders,
volgens besluit van het college van promoties
in het openbaar te verdedigen

op maandag 3 november 2025

om 16:30 uur precies

door

Chris Frans Henri Kamphuis
geboren op 22 maart 1993
te Oldenzaal

Promotor:

Prof. dr. ir. A.P. (Arjen) de Vries

Manuscriptcommissie:

Prof. dr. H.F. (Hannes) Mühleisen

Prof. dr. ir. D. (Djoerd) Hiemstra

Prof. dr. M.A. (Martha) Larson

Prof. C. (Craig) Macdonald (University of Glasgow, Verenigd Koninkrijk)

Prof. dr. S. (Suzan) Verberne (Universiteit Leiden)

Contents

1	Introduction	1
1.1	Problem Description and Research Questions	2
1.2	Thesis Contributions and Structure	3
1.3	Other Publications	5
2	Background	7
2.1	Introduction	7
2.2	Information Retrieval	8
2.3	Data Retrieval	20
2.4	Graphs	24
2.5	Reproducible Science	29
3	IR using Relational Databases	33
3.1	Introduction	34
3.2	Related work	34
3.3	Prototype OldDog	40
3.4	Variants of BM25	42
3.5	Experiments	49
3.6	Results	50
3.7	Conclusion	52
4	From Tables to Graphs	53
4.1	Introduction	54
4.2	Related work	55
4.3	GeeseDB	60
4.4	Experiments	72
4.5	Conclusion	76
5	Creation of the Entity Graph	77
5.1	Introduction	78
5.2	Related Work	79

5.3	REL	80
5.4	From REL to REBL	81
5.5	Effects on Execution	84
5.6	Conclusion and Discussion	85
6	Using the Entity Graph	87
6.1	Introduction	88
6.2	Background	90
6.3	MMEAD	93
6.4	How To Use	95
6.5	Entity Expansion with MMEAD	99
6.6	Beyond Quantitative Results	104
6.7	Conclusion and Future Work	108
7	Conclusion	111
7.1	Contributions	111
7.2	Future Work	113
	Bibliography	115
	Summary	127
	Samenvatting	129
	Acknowledgements	131
	Research Data Management	133
	Curriculum Vitæ	135
	SIKS Dissertations	137

Chapter 1

Introduction

I *propose* to consider the question, “Can machines think?”

Alan Turing - 1950

Like Turing in the quote cited above, I too propose to consider the question, “Can machines think?” Instead of approaching this through a thought experiment as Turing did, nowadays, one can approach this question by asking it to a search engine. When issuing this query to popular web search systems, we get varying results: the first result on Google is a passage generated from the article written by Turing, while the first result on Bing is a passage from a website that concludes machines cannot think.^{1,2}

When looking for *information*, we use systems that process queries every single day. While Google and Bing are all-purpose web engines that mainly focus on finding and retrieving information from the internet, people also use specialized search systems in their day-to-day lives: Amazon and eBay when we are looking for a product to buy, Scholar and ResearchGate for scientific resources, YouTube and TikTok for videos, or Facebook and LinkedIn when we are searching for people. It might even be possible that you are reading this text after you found this document through search.

When searching for the query, “Can machines think?”, searching through text documents might be sufficient for the person who searches. However, more than only considering text can be needed when someone’s information

¹However, if a machine cannot think, can we trust the result presented by this algorithm?

²These results were retrieved in October of 2022

need is more complicated. For example, when one wants to buy a product on Amazon, aspects other than text also need to be considered. Maybe you want to buy an iPhone; information on the price, which edition is the most recent, or which color it has are all essential to determine which one you want. You may also want to consider the rating provided by people who previously bought an iPhone.

If someone searches for people on LinkedIn, they might be more interested in persons they are connected to than strangers. If you are looking for someone to do a job, it is ideal that a shared connection can vouch for them. In this case, how people relate to each other in their network might indicate *relevance*. Not only the structure of how people relate to each other determines relevance; their experience, where they work, or reviews of their previous work will also matter.

Although it might be possible to encode all this information as written text, often, it is more convenient to save this information in a more structured approach. Where information retrieval researchers research the retrieval of “information” through text data, data management researchers research the retrieval of structured data [van Rijsbergen, 1979]. This thesis considers both methods simultaneously: systems that can work with structured and unstructured information are investigated.

1.1 Problem Description and Research Questions

Although information retrieval and data retrieval are research fields investigated by different disciplines, they are closely related, and systems that use both have been researched and developed in the past. Techniques developed in one community might help the other, as things like storing and quickly retrieving data are essential for both information and data retrieval.

Modern database system are often organized using the relational model. In such databases the data is organized as a sets of tuples. In the past, such databases have been used for information retrieval research. The database community has shown an increased interest in graph databases in recent years. In these databases the data is organized by modeling the data as a graph. As these databases are becoming more popular for data retrieval tasks where the data is highly interconnected, they might also benefit similar tasks in the information retrieval field where data (especially relevant results) is often highly interconnected. This thesis will investigate how these databases, with dedicated graph query languages, can be used

for information retrieval tasks. This leads us to this thesis’s main research question:

Research Question: How can information retrieval benefit from graph databases and graph query languages?

Three sub-research questions are defined to guide us in answering the main research question:

- *Research Question 1: What are the benefits of using relational databases for information retrieval?*
- *Research Question 2: Can we extend the benefits from using relational databases for information retrieval to using graph databases while being able to express graph-related problems easier?*
- *Research Question 3: When does information retrieval research benefit from graph data?*

Chapters 3, 4 and 6 will take these research questions and try to answer them respectively. Chapter 5 will present work used in Chapter 6. Then in Chapter 7, we will take the answers to these research questions and use them to answer the main research question.

1.2 Thesis Contributions and Structure

- Chapter 2 describes the necessary background information to provide context to the other chapters. The background that is described in this chapter concerns the “general” background knowledge for this thesis. Individual chapters also have related work sections that concern the background knowledge in those chapters. The information described in those related work sections contains overlapping information, i.e., the knowledge needed to understand the context of those chapters. This is done to make it possible to understand a chapter without first reading the chapters that come before it.
- Chapter 3 concerns information retrieval research using relational databases. Although, traditionally, inverted indexes are used for information retrieval, we show that by employing relational databases, some advantages are gained. First, attempts to use relational databases for information retrieval through the years will be described. One of the latter attempts used a column-oriented relational database system for information retrieval, achieving competitive efficiency compared

to a traditional system built using inverted indexes. This approach is re-implemented as a prototype system, which is then used for a reproduction study. This chapter establishes the usefulness of relational databases for information retrieval by demonstrating its benefits. The content in this chapter is based on the following published works:

- C. Kamphuis and A. P. de Vries. The OldDog Docker Image for OSIRRC at SIGIR 2019. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 47–49, Aachen, 2019b. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2409/docker07.pdf>
- C. Kamphuis, A. P. de Vries, L. Boytsov, and J. Lin. Which BM25 Do You Mean? A Large-Scale Reproducibility Study of Scoring Variants. In *Advances in Information Retrieval, ECIR '20*, pages 28–34, Cham, 2020. Springer International Publishing. ISBN 978-3-030-45442-5
- Chapter 4 takes the concept of employing relational databases for information retrieval and extends the relational model to the graph model. GeeseDB is introduced, a graph database prototype system built on top of an embedded column-oriented relational engine. Using the graph model makes it possible to express more complex information retrieval problems than when the relational model is used. These more complex models often use multi-stage retrieval approaches. GeeseDB is built on top of an embedded database system, so moving data between ranking stages can be done efficiently. The content of this chapter has previously been described in the following published works:
 - C. Kamphuis and A. P. de Vries. Reproducible IR needs an (IR) (graph) query language. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 17–20, Aachen, 2019a. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2409/position03.pdf>
 - C. Kamphuis and A. P. de Vries. GeeseDB: A Python Graph Engine for Exploration and Search. In *Proceedings of the 2nd International Conference on Design of Experimental Search & Information REtrieval Systems, DESIRES '21*, pages 10–18, Aachen, 2021. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2950/paper-11.pdf>
- Chapter 5 introduces the concept of entity linking. Specifically, the Radboud Entity Linker [van Hulst et al., 2020] system is discussed. When trying to utilize REL to annotate a large corpus with entity linking, issues were found that prohibited the annotations process to such an extent that it was not possible to do within a reasonable time. In order to be able to annotate the corpus, the REL toolkit internals

were upgraded, and a batch extension was developed. Altogether this led to more efficient software such that REL can annotate larger corpora. The content in this chapter has previously been described in the following published work:

- C. Kamphuis, F. Hasibi, J. Lin, and A. P. de Vries. REBL: Entity Linking at Scale. In *Proceedings of the 3rd International Conference on Design of Experimental Search & Information REtrieval Systems*, DESIRES '22, Aachen, 2022. CEUR-WS.org. URL <https://desires.dei.unipd.it/2022/papers/paper-08.pdf>
- Chapter 6 employs the software of chapter 5 to annotate a large web corpus. We developed a specification for sharing entity link annotations. Following this specification, we made annotations for the MS MARCO [Bajaj et al., 2016] corpora publicly available. Using these annotations, we show that, through query expansion, we can increase recall effectiveness for first-stage rankers on this dataset. Next, a demonstration shows how entity links can also be used for geographical information applications. This content has been described in the following published work:
 - C. Kamphuis, A. Lin, S. Yang, J. Lin, A. P. de Vries, and F. Hasibi. MMEAD: MS MARCO Entity Annotations and Disambiguations. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '23, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 978-1-4503-9408-6. doi: 10.1145/3539618.3591887
- Chapter 7 serves as a conclusion and summarizes the content discussed in the thesis. We will reflect on the research questions that motivated this dissertation, and discuss what future research is needed.

1.3 Other Publications

During the employment at Radboud University, the following work was also published:

- C. Kamphuis, F. Hasibi, A. P. de Vries, and T. Crijns. Radboud University at TREC 2019. In *NIST Special Publication 1250: The Twenty-Eighth Text REtrieval Conference Proceedings (TREC 2019)*, TREC '19, Gaithersburg, Maryland, 2019. [SI]: NIST. URL <https://trec.nist.gov/pubs/trec28/papers/RUIR.N.C.pdf>
- J. Lin, J. Mackenzie, C. Kamphuis, C. Macdonald, A. Mallia, M. Siedlaczek, A. Trotman, and A. P. de Vries. Supporting Interoperability Between Open-Source Search Engines with the Common Index File Format. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '20, page 2149–2152, New York, NY, USA, 2020a. Association for Computing Machinery. ISBN 9781450380164. doi: 10.1145/3397271.3401404

- P. Boers, C. Kamphuis, and A. P. de Vries. Radboud University at TREC 2020. In *NIST Special Publication 1266: The Twenty-Ninth Text REtrieval Conference Proceedings (TREC 2020)*, TREC'20, Gaithersburg, Maryland, 2020. [SI]: NIST. URL <https://trec.nist.gov/pubs/trec29/papers/RUIR.N.pdf>
- T. Schoegje, C. Kamphuis, K. Dercksen, D. Hiemstra, T. Pieters, and A. P. de Vries. Exploring task-based query expansion at the TREC-COVID track. *CoRR*, abs/2010.12674, 2020. URL <https://arxiv.org/abs/2010.12674>
- C. Kamphuis. Graph Databases for Information Retrieval. In *Advances in Information Retrieval*, pages 608–612, Cham, 2020. Springer International Publishing. ISBN 978-3-030-45442-5
- D. Hiemstra, G. Hendriksen, C. Kamphuis, and A. P. de Vries. Challenges of Index Exchange for Search Engine Interoperability. In *Proceedings of the 5th International Open Search Symposium*, OSSYM23. Open Search Foundation, 2023

Although this work relates to the main subject of the thesis, it is not used as source material for the work described in this thesis.

Chapter 2

Background

Machines take me by surprise
with great frequency.

Alan Turing - 1950

2.1 Introduction

This chapter aims to provide the context in which this work is written. By introducing the related scientific fields, this chapter sketches a broader context wherein this research exists.

First, the field of information retrieval is introduced. Within the field of information retrieval, different approaches to *retrieving information* have been used throughout the decades. These approaches will be introduced such that the reader understands what kind of techniques are used to retrieve information. The work described in this thesis would also be considered information retrieval research.

Secondly, techniques from different research areas are also used for this research. Specifically, the research described in this thesis extensively uses techniques typically investigated by the *data management* community. This field will also be introduced, where especially the techniques used for the research described in this work are highlighted.

Then, the concept of *graphs* will be explained. Graphs themselves are just mathematical structures that model the relations between objects. The objects are modeled as *nodes* and relations between them as *edges*. Many kinds of graphs exist; they are helpful for modeling problems that can be

expressed formally through this framework. Different kinds of graphs will be explained, and examples of how to use them will be provided.

Finally, we will describe the idea of *reproducible science*. In general, scientific work must be reproducible. In our research, we spend additional effort on reproducible science. What reproducibility exactly entails will be introduced.

2.2 Information Retrieval

The research described in this thesis is *information retrieval* research. Colloquially, one can refer to information retrieval as the science of everything relating to search engines. This field encompasses all aspects: from user experiences of search systems to storage algorithms and fast retrieval of the information items users search. Baeza-Yates et al. [1999] introduce information retrieval as the following:

Information retrieval (IR) deals with the representation, storage, organization of, and access to information items.

Following this description, an information retrieval system, i.e., a search engine, is a system that allows users to access information items that they are looking for. How this works internally is usually not interesting for the user, they just want to find the information they need. Typically, and also in this thesis, the term *document* refers to information items, even though the item the user seeks is not necessarily a literal document.

Consider the following situation; someone wants to know whether it will rain in the coming 30 minutes. They might query a web search engine with the text *weather*. In this case, the user does not care how the search engine works but only about the result. In order to satisfy the user's information need, the search engine needs to (1) present the correct weather forecast and (2) do this quickly. The user will be dissatisfied if the search engine presents the incorrect weather forecast or cannot do this in milliseconds.

In order to answer this inquiry for information correctly, the information retrieval system needs more context. It can only correctly know the weather if it is known where the user resides. When accessing search engines through a mobile device or computer, this information is sent along with the query as metadata. This way, the search engine can correctly answer even though the user did not explicitly give this information.

When the user's location is known, the search engine needs to find a webpage that correctly provides the weather for that particular location. As there are billions of web pages in the search engine's index, correctly

identifying which one contains information about the weather at that particular location and time and then retrieving it in milliseconds is the next challenge.

2.2.1 Inverted Indexes

Search engines must organize the data because it is infeasible to iterate over all web pages to check if the correct information is available. The most common data structure used in search engines is the *inverted index*.

Inverted indexes have been used for decades in the field of information retrieval. They are designed to access documents quickly. If a system cannot provide the data quick enough, the user is likely to switch to a different system. Documents are typically formed through words that form sentences, which form paragraphs, and, eventually, stories. When considering a document in its standard form, it is not trivial to determine efficiently whether it contains a specific word. When searching for something, it is paramount that the search engine knows what documents contain the keywords in the query. To efficiently access this data, instead of storing documents as is (a mapping from documents to words), search engines store the inverted information (a mapping from words to documents). To provide an example, consider the following three short documents:

1. Cats and dogs are animals.
2. Cats are smart animals.
3. Dogs are great at tricks.

Then the inverted index looks something like this:

animal: [1 2], **cat:** [1 2], **dog:** [1 3], **great:** 3, **smart:** 2, **trick:** 3

A couple of interesting observations can be made here. Not all words appear in the inverted index. Typically words with little “meaning”, or so-called stop words, are removed. These words tend to contain no semantic information and can therefore be dropped. This process is called stopword removal or *stopping*. Then, words are sometimes shortened to a *stem*. Let us say we have a query only mentioning the word “dog”, so not plural, then we still want to be able to find the documents that mention the plural form. This shortening to the word’s stem is often called *stemming*. When words are reduced to their dictionary form, this process is referred to as *lemmatization*. In this case, the words are ordered alphabetically. When the number of documents in a collection increases, the number of unique

words also increases. By storing the words alphabetically, a computer can find the entries for that particular word more efficiently.

When someone queries the search engine with the following query: “dog tricks”, the search engine can directly find which documents contain these words using the inverted index. It can then automatically discard all documents that do not contain these words. For the keywords, the search engine retrieves the lists associated with them. These lists are called *posting lists* in information retrieval. Entries in such lists are referred to as *postings*. In this case, the postings only contain the document identifiers as data. Generally, however, information like how often a word appears in a document or its location is also stored in the posting. In practice, posting lists are compressed, and instead of storing the direct identifiers, the gaps between them are stored such that the index becomes smaller. In literature, these gaps are referred to as delta gaps [Moffat and Zobel, 1996].

In this case, the posting lists of the words “dog” and “trick” are retrieved: [1 3] and [3]. Then some scoring method can process these lists and assign a score to the documents present in these lists. As document 3 is the only document containing both words, it makes sense to consider this document the most relevant. However, assessing relevance ordering is more complicated when multiple documents contain both words with different frequencies. To create ranked lists of documents based on their estimated relevancy, different models for scoring documents have been proposed in the past. In the following sections, these models are presented.

2.2.2 Ranking Models

Boolean Retrieval

The Boolean retrieval model was used in the early days of information retrieval [Lancaster and Fayen, 1973, Lancaster, 1979, van Rijsbergen, 1979]. Boolean retrieval can be formulated as the following; let,

$$T = \{t_1, t_2, \dots, t_n\} \quad (2.1)$$

be the set of all index terms that appear in the collection. Let,

$$D = \{d_1, d_2, \dots, d_m\} \quad (2.2)$$

be the set of all documents, where every document is a subset of T . Specifically, the terms that appear in that document. Then a query can be any Boolean expression over T . All documents that adhere to the Boolean expression formed by the query are considered relevant. To give an example, consider the same three documents as before:

1. Cats and dogs are animals.
2. Cats are smart animals.
3. Dogs are great at tricks.

Then, D is formulated as $\{d_1, d_2, d_3\}$ with,

$$d_1 : \{cat, dog, animal\} \quad (2.3)$$

$$d_2 : \{cat, smart, animal\} \quad (2.4)$$

$$d_3 : \{dog, great, trick\} \quad (2.5)$$

One could be interested in finding all documents that mention dogs but not cats. Then we can formulate the following Boolean query:

$$Q = dog \wedge \neg cat \quad (2.6)$$

If we find all subsets that adhere to the individual terms in this conjunction, we get the following set expression:

$$Q = \{d_1, d_3\} \cap \{d_3\} = \{d_3\} \quad (2.7)$$

This shows that document d_3 is the only document that mentions dogs, but not cats. Although it is possible to express complicated expressions with Boolean logic, a document either satisfies the expression or does not. All documents that satisfy the expression are considered of equal importance; creating a ranking between them cannot be done with this approach alone.

Using Boolean retrieval, there is no apparent difference between data retrieval and information retrieval; there is only one correct solution: all documents that satisfy the restrictions imposed by the query.

Vector Space Models

In 1975, the vector space model was introduced by Salton et al. [1975]. The idea behind the vector space model is to represent documents and queries as vectors. Consider a collection that has t unique terms, then we can represent a document d in the collection as a t dimensional vector $\mathbf{d}$ ¹:

$$\mathbf{d} = (w_0, \dots, w_t) \quad (2.8)$$

¹We use bold text to distinguish between document d , and its vector representation $\mathbf{d}$

where w_i represents the weight associated with the i -th term in the collection for this document. These weights can be binary, or if one only wants to consider the importance of the term in the document, they can be the term frequency or any weight that considers the “general” term importance.

Then if we have a query q , we can represent it in the same way:

$$\mathbf{q} = (q_0, \dots, q_t) \quad (2.9)$$

where q_i represents the weight of the i -th term in the collection. Most values in these vectors will be zero.

Now it is possible to calculate a similarity score between the query and every document. A widely used similarity metric is the cosine similarity between two vectors, which between a document d and a query q is defined as; the dot-product of vectors $\mathbf{d}$ and $\mathbf{q}$ divided by the product of their ℓ_2 -norms:

$$\cos(d, q) = \frac{\mathbf{d} \cdot \mathbf{q}}{\|\mathbf{d}\| \|\mathbf{q}\|} \quad (2.10)$$

Calculating this for the documents in the collection gives an estimated value of relevance for all of them. Ordering them based on this estimated value then produces a ranked list where the highest ranked document is presumed to be the most interesting for the user. As mentioned before, the weights represented in the document vector do not necessarily have to be binary. Salton et al. showed that the product of the term frequency with a general measure of term importance is highly effective. The measure they used for term importance was the inverse document frequency proposed by Spärck Jones [1972] three years earlier.

Probabilistic Relevance Models

In 1976, Robertson and Spärck Jones developed the probabilistic ranking framework. Within this framework, it is assumed that a document has a certain probability of being relevant given a query. Then, a search system that implements models within this framework ranks documents with the highest probability of being relevant the highest.

Within this framework, making an assumption of binary independence, which is nicely explained by Robertson and Zaragoza [2009], the so-called Robertson-Spärck Jones weight can be derived. When assuming that no relevance labels are available, this leads to an approximation of the classical inverse document frequency.

By extending this function to include within-document term frequency information and document length normalization, the well-known BM25

function can be derived. Chapter 3 will focus on this model, and additional information explaining this model will be presented there.

Language Models

Later, language models were proposed [Song and Croft, 1999, Hiemstra, 2001, Zhai and Lafferty, 2002]. The idea behind language models is that documents are considered samples of language, and queries are modelled as the result of a generative process from these samples. This process generates terms in the query by randomly selecting one of the words from that sample. Let the probability of a document d , consisting of n different terms, generating a term t_i be:

$$P(t_i|d) = \frac{tf_{t_i,d}}{\sum_{j=0}^n tf_{t_j,d}} \quad (2.11)$$

with $tf_{t_i,d}$ being the term frequency of term t_i in document d .

Then, if we have a query Q , it is possible to calculate the likelihood that a document sample generated this query. We know for every term in the query what the probability is that one of the document samples generated that term, so with Bayes rule:

$$P(d|Q) = \frac{P(d) \times \prod_{q \in Q} P(q|d)}{P(q)} \quad (2.12)$$

Assuming an uniform relevance prior for the documents, and the fact that $p(q)$ is constant for all documents we can determine the document with maximum likelihood through:

$$P(d|Q) \propto \prod_{q \in Q} P(q|d) \quad (2.13)$$

There are some issues when we take the document with the maximum likelihood. First, if a query contains multiple terms, a document can only get a non-zero likelihood of generating that query if all query terms exist in that document. If this is not the case, the probability of generating that term would be zero. As we take the product of the individual probabilities of the query terms, if one of them would be zero, the product would also be. The second issue is that there is no term-independent specificity weighting. Terms with a higher degree of specificity should increase the probability of relevance more than “filler” words. The previously described inverse document frequency took care of this in the vector space model.

Zhai and Lafferty [2004] show that a simple smoothing technique addresses both issues at once. Instead of only considering the samples formed

by the documents, a collection-wide sample is being used, in a mixture model. This collection-wide sample generates terms according to the frequency in which the terms appear in the collection:

$$P(t_i|C) = \frac{df(t_i)}{\sum_{j=0}^N df(t_j)} \quad (2.14)$$

With N being the number of terms in the collection, and $df(t_i)$ the document frequency of term t_i ; i.e., the number of documents in the collection in which term t_i appears. Then, the probability estimated by the document sample is interpolated by the probability based on the collection sample:

$$P(d|Q) \propto \prod_{q \in Q} \omega \cdot P(q|d) + (1 - \omega) \cdot P(q|C) \quad (2.15)$$

with ω being the smoothing factor, i.e., how much is the collection sample weighted against the document sample?

This technique is also known as Jelinek-Mercer smoothing [Jelinek and Mercer, 1980]. Smoothing the probability estimates like this, solves both problems: As the collection sample contains all terms, it is no longer possible for probabilities generated by a specific term to be zero, so the resulting product will not be zero either. Words with a higher specificity also get a boost; words with a high specificity have a low probability of being generated by the collection sample, meaning that it is more important for a document to contain that word to achieve a high probability of relevance. Words with a low specificity have a smaller effect; the probability of them being generated by the collection sample is already relatively high.

Learning-to-Rank

With the increase in compute power, learning-to-rank became a popular approach to ranking. The methods described in the previous sections can all be considered unsupervised learning methods if we look at them from a machine learning perspective. Learning-to-rank models would be considered supervised or reinforcement learning methods from that perspective. The goal is to learn a model that, given a query, ranks documents relevant to that query higher than those not relevant to that query. In general, there are three ways to achieve this; for all these methods, relevance labels are assumed to be available from which the function can be learned. Learning-to-rank methods can be categorized into the following three categories [Liu, 2010]:

Pointwise Pointwise methods are most comparable to standard regression methods (e.g. linear regression). The idea is to learn a function that directly estimates a document's relevance label. Consider that we have n query-document pairs for which relevance labels are available, with m independent variables per pair. Such variables can be query dependent, e.g. the scores produced by the models described in the previous sections, or they can be query independent, e.g. a document importance measure like the PageRank score [Page et al., 1999]. The relevance labels for these pairs can then be expressed as a vector $\mathbf{y}$, and the independent variables as a matrix $\mathbf{X}$:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} \quad (2.16) \quad \mathbf{X} = \begin{bmatrix} x_{1,1} & \cdots & x_{1,m} \\ \vdots & \ddots & \vdots \\ x_{n,1} & \cdots & x_{n,m} \end{bmatrix} \quad (2.17)$$

The goal would be to learn a function $\hat{\mathbf{y}} = f(\mathbf{X})$ such that some distance measure between $\mathcal{L}(\hat{\mathbf{y}}, \mathbf{y})$ is minimized. Commonly used measures for this are mean square error or mean absolute error. The pointwise learning approach to ranking has some undesirable effects, however. Consider a situation where we have three documents; $[d_1, d_2, d_3]$, with relevance labels $[1, 0, 0]$, i.e. document d_1 is relevant and the other documents are not. Then, two different ranking functions, a and b , might produce the following relevance estimates: $a = [1, 0.9, 0.8]$ and $b = [0, 0.1, 0.2]$, which results in the following respective rankings: $[d_1, d_2, d_3]$ and $[d_3, d_2, d_1]$. The ranking produced by system a is objectively better. It ranks the relevant document as the most relevant. System b , however, ranks the relevant document as the least relevant. If we evaluate the two systems with a loss function like mean absolute error however, system b would seem the better one as the mean absolute error to $[1, 0, 0]$ is smaller.

Pairwise Pairwise ranking tries to deal with this problem by not directly learning from the relevance labels but by looking at the relative preferences between pairs of documents. The idea behind this is that it does not matter what the actual ranking score value is, as long as it is higher for relevant documents than for non-relevant ones.

Pairwise functions consider, given a query, two documents at a time. The function tries to determine which one of the two documents is more relevant than the other. If we consider the same set of relevance labels $\mathbf{y}$ and independent variables $\mathbf{X}$, then if we would take a pair of documents x_k and x_l , with their corresponding relevance labels y_k and y_l , assuming

$y_k \neq y_l$, then we will learn the following function:

$$f(x_k, x_l) = \begin{cases} 0 & \text{if } y_k < y_l \\ 1 & \text{if } y_k > y_l \end{cases} \quad (2.18)$$

This function can be learned by any binary classification method. After the method is learned, it can be used to determine an ordering of documents. Pairwise methods tend to be more expensive than pointwise methods, as the number of samples to train on for pointwise methods is n , with n being the number of documents for which relevance assessments are available per query. For pairwise methods, every document with a label can be compared with every other document, creating n^2 training samples per query.

A disadvantage of pairwise methods is that every pair is treated equally when training a model. There is no difference when comparing the two highest-ranked documents to comparing the two lowest-ranked documents, even though users typically only consider documents at the top of the ranking [Joachims et al., 2005].

Listwise Instead of looking at pairs of documents only, the listwise approach considers the whole ranking. This way, the top of the ranking can affect the function’s learning more than the bottom of the ranking. Generally, when learning a listwise learning-to-rank method, the goal is to optimize some ranking metric quality directly. For example, LambdaMart [Burges, 2010] calculates what the differences on ranking metrics are by virtually swapping documents, these differences are used for determining gradients. (The gradients can not be calculated directly as the evaluation metric is not continuous with respect to the model parameters.) Examples of ranking metrics will be introduced in Section 2.2.3.

Multistage Ranking In many cases, applying inference to all documents in the collection is prohibitively costly. To reduce computational costs, multi-stage ranking is often used [Broder et al., 2003]. A non-learning-to-rank approach is often used to create an initial ranking, under the assumption that the top- k documents contain all relevant documents. The learned method only reorders the top- k documents to present to the user.

In these cases, it can make sense to rank the top- k documents using an inverted index, while the reranking is done with another ranking approach.

Vector Space Models revisited

With the rise of large language models after the introduction of BERT by Devlin et al. [2019], the vector space model has become more popular

again due to what is called *dense* retrieval.

The basic idea is to estimate the relevance of a document d for a query q as follows:

$$\text{rel}(q, d) = \omega(\phi(q), \psi(d)) \quad (2.19)$$

Here, ψ and ϕ are learned functions that use large language models to map the query and the document to vectors with relatively low dimensionality [Lin, 2022].² These functions map the query and the document to the same space. When documents are relevant to the query, the resulting vectors should be similar, while if the document is not relevant, they should be dissimilar, as measured through a similarity function ω , often the dot-product or the cosine similarity as shown in Equation (2.10).

Finding the vectors closest to a query vector then corresponds directly to the nearest neighbor search problem. In the case of dense retrieval, dedicated dense retrieval systems (e.g. FAISS [Johnson et al., 2019]) tend to be used instead of an inverted index.

2.2.3 Evaluation

In the previous section, a variety of ranking approaches has been introduced. When comparing different ranking methods to each other, we need to be able to measure the quality of a result list for a given query. Many metrics have been introduced in the literature. Here, we focus on the most common evaluation metrics, including all the metrics used in this thesis.

The metrics are @ k metrics, these only evaluate the documents ranked up to rank k , i.e., if $k = 30$, we only consider the first 30 documents produced by the system. The evaluation metrics that we describe are either set metrics or ranking metrics, set metrics consider the first k documents, but not the ranking within the first k documents, while ranking metrics also consider the ranking of the first k documents retrieved.

Precision

A straightforward approach to evaluating a set of documents produced by a system is through precision. The precision metric counts the number of relevant documents found by the system, which is then divided by the total number of documents found by the system:

$$P@k = \frac{\sum_{i=1}^k \text{rel}(d_i)}{k} \quad (2.20)$$

²10² - 10³ instead of the 10⁸ - 10⁹ parameters of the LLM's.

In this case, $\text{rel}(\cdot)$ is a function that returns 1 if the document is relevant; otherwise, 0. If $k = 30$, then only the first 30 documents are considered. If ten of these documents are relevant, then $P@30 = \frac{10}{30} = \frac{1}{3}$. Precision can, however, be somewhat limited in its use. In this example, the quality of the found documents might not seem good, as only a third of the produced documents are relevant. However, when only ten relevant documents are present in the collection, the $P@30$ found is the maximum attainable.

Recall

Another approach to evaluate a set of documents would be to measure how many relevant documents were retrieved compared to the total number of relevant documents in the collection. This measure is called recall:

$$R@k = \frac{\sum_{i=1}^k \text{rel}(d_i)}{\sum_{i=1}^N \text{rel}(d_i)} \quad (2.21)$$

Here N is the number of documents in the collection. If, again, $k = 30$ and we found 10 relevant documents out of 15 known relevant documents in the collection, then $R@30 = \frac{10}{15} = \frac{2}{3}$. The problem with recall is that it does not distinguish well between systems if only few relevant documents exist. Also, it is far from straightforward to determine the total number of relevant documents for an arbitrary query.

Average Precision

Consider the situation where three documents are retrieved, and two systems might produce rankings with the following relevance labels: $[1, 0, 0]$ and $[0, 0, 1]$. Both rankings achieve the same $P@3$ and $R@3$ scores. The previously introduced metrics cannot distinguish between ranking quality in this situation. It is, however, clear that the order of documents in the first ranking should be preferred over the document-list in the second ranking. Instead of measuring the precision only at rank k , the average precision measure calculates it every time a relevant document is encountered. The sum of these values is divided by the number of relevant documents found up to k :

$$AP@k = \frac{1}{\sum_{i=1}^N \text{rel}(d_i)} \sum_{i=1}^k P(i) \cdot \text{rel}(d_i) \quad (2.22)$$

Looking back at the two rankings, $[1, 0, 0]$ and $[0, 0, 1]$, the rankings produce an $AP@3$ of 1 and $\frac{1}{3}$, respectively.

(Normalized) Discounted Cumulative Gain

Discounted cumulative gain is an evaluation metric developed after the observation that users often only look at the highest-ranked documents and that documents might not have binary relevance labels (when we refer to their labels as graded relevance assessments) [Järvelin and Kekäläinen, 2002]. If a document is ranked at, say, rank 5, it might be possible that it is not considered by the user anymore, because they would rather reformulate the query. Plus, the metrics described before would rank the following two rankings to be of equal quality: $[2, 1, 0]$ and $[1, 2, 0]$ Discounted cumulative gain was developed with the intention that there should be a discounted gain every time a new relevant document is found, while non-binary relevance labels can be taken into account. This discount is applied by dividing by the logarithm of the rank plus one:

$$DCG@k = \sum_{i=1}^k \frac{\text{rel}(d_i)}{\log_2(i+1)} \quad (2.23)$$

In this case, $\text{rel}(\cdot)$ does not return a binary value but the actual relevance assessment label value.

A disadvantage of DCG is that the scores produced by this metric can vary greatly between rankings. If many highly relevant documents exist for a query, the values DCG takes can be much higher than if only one moderately relevant label is available. Normalized discounted cumulative gain takes care of this by dividing by the highest possible DCG that can be achieved. So normalized discounted cumulative gain is calculated by dividing the DCG of the ranking, by the DCG of the ideal ranking:

$$NDCG@k = \frac{DCG@k}{IDCG@k} \quad (2.24)$$

with $IDCG@k$ being the $DCG@k$ for a perfect ranking.

Reciprocal Rank

The metrics described before score the whole ranking up to rank k . This might be unnecessary, for example, in a question answering scenario. If a document found is relevant, it may not matter anymore what the relevance assessment of the remaining documents. After all, the user would already have the answer, and likely stop looking for more information. The reciprocal rank therefore takes only the first found relevant document into account. Reciprocal rank [Voorhees and Tice, 2000] is the inverse of the index of the

first found document:

$$RR@k = \frac{1}{\text{rank}(d)} \quad (2.25)$$

where only the first k documents are considered. Here, $\text{rank}(d)$ refers to the position of the first found relevant document. If the index of the first relevant is greater than k , the resulting score is 0.

2.2.4 Significance Testing

Often, significance testing is used to evaluate if one approach is better compared to another. First, it is assumed that the output scores of two approaches, measured for example by one of the above metrics, to be sample distributions. Then with some assumptions, it is possible to calculate the probability that we would expect to see the difference observed to at least be as extreme if they were sampled from the same population distribution. If this probability is lower than, say for example 5%, one can assume that the population distributions are not the same, and that the higher scoring approach is better than the other; i.e. it is not likely that the observed difference can be explained due to a sampling error alone.

However, some of the measures described above can still be problematic when they are used to compare different ranking approaches. It is not always entirely clear how to do significance testing on some metrics. For example, Fuhr [2018] argues that in the case of reciprocal rank, the scores are not on an interval scale, but an ordinal one. It is not possible to calculate a mean score for an ordinal scale, which is however necessary for the many common tests for significance, e.g. a t-test. Fuhr describes this, and other mistakes made when evaluating IR systems, plus how to avoid them. There are however some contrary views [Sakai, 2021] that do not necessarily agree with these views. Specifically, Sakai argues that are cases where ordinal is often averaged, and tested with parametric tests; for example in the case of multi-point scale ratings in user studies. To Sakai it also not completely clear if Reciprocal Rank cannot be considered as a interval scale.

2.3 Data Retrieval

In the case of information retrieval, the concept of *relevance* is critical. When someone issues a query, it does not necessarily have to be the case that words in the query must be present for the document to be relevant. Information retrieval has to deal with the uncertainty inherent to the information seeking process. Also, if only some relevant documents are found, this does not

necessarily have to be a problem, as long as the information need of the user is satisfied. If we think about a search engine from a user's perspective, a document is only relevant if it contains the information they seek. In order to find this information, the user has to express their information need in a query, which the systems then have to interpret somehow in order to present the user (hopefully) the correct information. It is however possible that two different users issue the same query, e.g., *the weather*. Suppose one user is interested in today's weather, while the other is interested in tomorrow's weather. Different documents will be relevant for the user in these cases: relevancy is also tied to the user's expectations, that may not be known to the system.

In the field of *data retrieval*, the notion of relevancy does not exist. When we speak about data retrieval, we are interested in retrieving all data that fulfills some *predicate*. A predicate is a function that takes some input, i.e. the data, and it returns either true or false. In data retrieval we are interested in the data for which such a predicate resolves to true. It assumes that there is no ambiguity in the data and the request for data. An example query could be; to provide a list of all dog breeds. In this case, the system should return a list of all dog breeds known by the system. If it would accidentally forget to return one dog breed or includes an answer that is not correct, the system works incorrectly. The most commonly used data retrieval systems are relational database management systems (RDBMS) or relational databases. Relational databases represent tabular data as relations, a concept of from relational algebra; where a relation is defined as a set of tuples. In a relational database, a table corresponds to the relation, a row is one of the tuples, and a column can be considered an attribute in the tuple. Relational databases implement the operators as defined by relational algebra, which means users can select columns from a table (projection), filter rows that fulfill certain predicates (selection), and combine tables into new ones (join). There are differences between the theoretic framework of relational algebra and its wide-spread implementation in relational databases. For example, in practice, it is often possible to have duplicate rows in a relational database, which is inconsistent with a set.

In the 1970s, relational databases have standardized on the structured query language (SQL). SQL is a declarative language; expressions specify which data needs to be retrieved, but not how. Depending on what a user expects from the program, choosing different database management systems for different applications makes sense. In the case of storing employee records for a company, it makes sense to store this data in a database suited for online *transaction* processing (OLTP). Processing a transaction means

that we update the data; e.g. adding new data to a table. When data analysis for scientific studies is important, storing data in a database system focused on online *analytics* processing (OLAP) makes more sense. With analytics we mean queries that retrieve and process data, as opposed to queries that update the data.

2.3.1 Physical vs. Logical models

As SQL is a declarative language, data management systems’ physical and logical models are strongly separated. Although many dialects of SQL exist, basic SQL queries are system agnostic and can run on different types of SQL engines. The logical model of data management systems are the SQL queries; what should be done. The physical model is how this should be done. Depending on which system is used, things like join ordering, how the data is processed, or how it is represented internally, can be very different. The logical model does not “mind” how the query is resolved, as long as it is done correctly.

In the case of information retrieval, the physical and logical models are often more entangled. Model specification and query processing often co-exist. When this is the case, it can be harder to detect why systems differ when they implement the same models. Chapter 3 will highlight this in more detail.

Lin [2022] discusses how information retrieval might benefit from separating physical and logical models in the context of dense retrieval models. Chapter 4 will discuss his proposal in more detail.

2.3.2 Columns vs. Rows

Earlier database management systems were row oriented. Row-oriented means that when we look at how data is physically stored on the computer, all data within a tuple is represented near each other in memory. This orientation is especially well suited when a database needs to process many transactions. In the past, some IR systems have been built on top of such databases. However, these were much less efficient than when inverted indexes were used. Later, column-oriented databases were developed Copeland and Khoshafian [1985], Boncz [2002]. These databases are more suited for analytical queries, at the cost of more expensive transaction processing. As these databases are designed for efficient and scalable analytics, they are also more suited for information retrieval systems. This practice is used throughout this thesis.

2.3.3 NoSQL and Graphs

In recent years there has been interest in database systems that approach data retrieval differently than traditional relational database systems [Cattell, 2011]. These systems approach data retrieval without relying on the relational model. This could, for example, mean that joins are inefficient or not even possible to express within the system. This might seem limiting, but if certain operations are not supported, it might be possible to increase the efficiency of other operations. A very efficient key-value store like BigTable [Chang et al., 2006] is a good example to illustrate this trade-off.

In particular, graph database (e.g. Neo4J [Webber, 2012]) systems are a type of NoSQL database management systems that have been of interest in the database community in the last few years. The concept of graphs is explained in Section 2.4. The research presented in Chapter 4 explores the possible benefits of considering these databases from an IR point of view.

2.3.4 Embedded Databases

Database management systems usually run on dedicated database servers. It is, however, also possible to use so-called embedded database systems. These systems do not require a dedicated database server, but work within the application process itself. An advantage of embedded databases is that they share memory space with the application process. This makes it possible to instantly move data from the database to a usable format for the application process. For information retrieval applications, this might be beneficial, especially in the case of multistage ranking approaches as described in Section 2.2.2. However, embedded databases might not be the best solution in every case. When a dedicated database server is run, it is easier for multiple clients to read and write to the same database simultaneously.

2.3.5 Databases for Information Retrieval

As databases are designed for the purpose of data retrieval, they are especially suited for information retrieval in the case of Boolean retrieval. However, it is harder to express the models that capture relevancy using a purely data retrieval approach. When expressing models that represent uncertainty, i.e. model a likelihood measure of relevance, much data must be processed efficiently—inverted indexes are used to select critical data quickly. Older (row-based) database systems were not efficient enough to retrieve all necessary data to create rankings and model uncertainty simultaneously.

After columnar databases became prominent, it was possible to express simple ranking models that consider uncertainty, with query plans that can be executed in a reasonably efficient manner. In this thesis, we will use, column oriented databases for information retrieval problems. Specifically, MonetDB [Boncz, 2002], a non-embedded columnar database, is used for the work presented in Chapter 3, and DuckDB [Raasveldt and Mühleisen, 2019], an embedded column oriented database, is used for the work presented in Chapters 4 and 6.

2.4 Graphs

As mentioned in Section 2.3.3, in this thesis, we present work where approach information retrieval using graphs. In this section we describe what a graph is and which kinds of graphs exist. The graphs definitions follow the descriptions by Sakr et al. [2021] and Angles [2018].

Basic Graphs A *graph* is a structure consisting of a set of objects where pairs of these objects can be related. Typically, these objects are called nodes or vertices, and the relation between a pair of objects is referred to as an edge. In this context, a relation means something different than a relation in the context of relational algebra. In a RDBMS, a relation is implemented through a table, so tuples are “related” if they exist in the same table. In a graph, a relation describes a link between two objects. These objects do not necessarily have to be presented in the same table if we would use the relational model. We still use the word relation when talking about graphs, as done in literature.

More formally, a graph G is a pair of vertices V and edges E ;

$$G = (V, E) \tag{2.26}$$

where V is the set of vertices in the graph, and E a set of pairs (sets with two elements) of vertices:

$$E = \left\{ \{v_1, v_2\} \mid \{v_1, v_2\} \in V^2 \right\} \tag{2.27}$$

Graphs can be depicted graphically illustrated as a set of circles, where every circle represents a vertex. If an edge exists between two vertices, a line is drawn between the two circles. Figure 2.1a, for example illustrates a graph with $V = \{x_1, x_2, x_3\}$ and $E = \{\{x_1, x_2\}, \{x_2, x_3\}\}$. It does not necessarily have to be the case that every vertex is included in at least one edge.

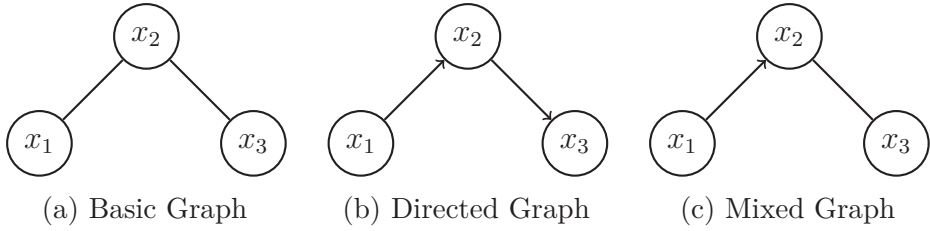


Figure 2.1: Example of (a) a simple graph, (b) a directed graph, and (c) a mixed graph.

A graph is called *bi-partite* if it is possible to separate the set of vertices in two disjoint subsets, such that all edges only connect vertices from the two subsets with each other (there are no edges between vertices in the same subset).

The graph's structure can be used to analyze how things relate. Considering geographical maps, we can model countries as vertices and create an edge between two countries if they are neighboring. Then, we can use the graph structure to analyze the possible paths from one country to another where only the fewest possible other countries are crossed, on a smaller scale this is useful for navigation systems. For many use cases, however, it is helpful to let graphs have less restrictive rules on how they might be formed.

Directed Edges The most basic extension of graphs is that they are allowed to have directed edges. This means an edge can exist from one node to another, but this connection does not automatically also exist in the opposite direction. In the case of countries and their neighbors, this does not make sense; if one country borders another, the opposite is also true. However, on the internet, one website might contain a hyperlink to another, while the opposite does not have to be the case. We need directed graphs to properly model the situation. Again, we can use this model to see how many paths exist from one website to another, with the restriction of visiting two other websites. In this case, E is a set of ordered pairs of vertices:

$$E = \{(v_1, v_2) \mid (v_1, v_2) \in V^2\} \quad (2.28)$$

Figure 2.1b shows an illustration of the following directed graph: $V = \{x_1, x_2, x_3\}$ and $E = \{(x_1, x_2), (x_2, x_3)\}$

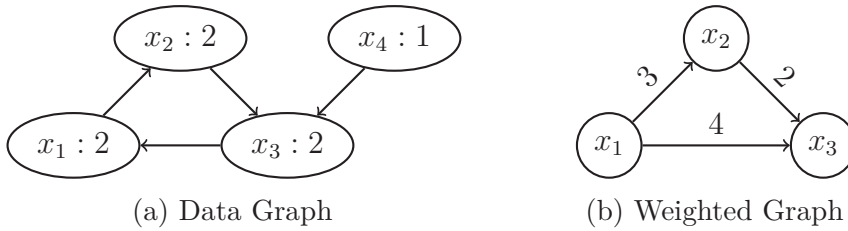


Figure 2.2: Example of (a) a data graph, and (b) a weighted graph.

Mixed Graphs Mixed graphs are graphs that accept both undirected and directed edges. In that case, a graph is defined with two sets of edges, where one set describes the undirected edges and the other the directed edges:

$$G = \{V, E_1, E_2\} \quad (2.29)$$

with E_1 and E_2 being defined as edges as presented in Equations (2.27) and (2.28) respectively. Figure 2.1c shows the illustration of a mixed graph: $V = \{x_1, x_2, x_3\}$, $E_1 = \{\{x_1, x_2\}\}$, and $E_2 = \{(x_2, x_3)\}$. Mixed graphs are sometimes used in applications where the direction of an edge is not known but it is clear that a relation exists. Types of these graphs appear, for example, in causality research.

2.4.1 Labeled Graphs

None of the types of graphs we described so far contain data themselves; we can only reason about the structure of the graphs. Even though this already has many practical applications, graphs are more useful for the applications we are interested in when they are allowed to contain data. There are two ways to add data to graphs; it is possible to add data to the vertices or the edges.

Data Graph A data graph is a type of graph where vertices have data. Figure 2.2a shows an example of a (directed) graph where every vertex is labeled with a number. These graphs are, for example, used in the PageRank algorithm [Page et al., 1999]. In this algorithm, every vertex represents a webpage, and a directed edge exists between a node a to b if a link on website a points to website b . Then by randomly walking over this graph or visiting a new website randomly, it is possible to count how often a node has been visited. The number of visits corresponds to the PageRank score. Normalized PageRank scores correspond to the probabilities found in a Markov process in equilibrium.

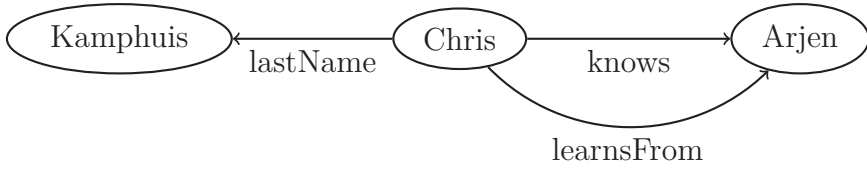


Figure 2.3: Example of a W3C RDF graph.

Weighted Graph A weighted graph is a type of graph where edges contain data. Figure 2.2b shows an example of a (directed) graph where every edge is labeled with a number. For example, these kinds of graphs are used to calculate the shortest distance between nodes. They are, for example, useful for navigation systems. There are two possible paths if we need to travel from location x_1 to location x_3 . The path from x_1 directly to x_3 is shorter than going from x_1 to x_3 via x_2 . When graphs are larger, algorithms like Dijkstra’s algorithm [Dijkstra, 1959] can efficiently find the shortest path over weighted graphs.

Pregel/Giraph Graph An obvious extension of the data graph is, of course, adding data to both the vertices and edges. In this context, we refer to them as Pregel/Giraph graphs [Malewicz et al., 2010, Apache Software Foundation, 2012], as these systems are widely used implementations of these graphs. Some algorithms need graphs to have both edge and vertex data. An example is the node2vec algorithm, which generates vector representations of the vertices of the graph. These graphs are used for various machine learning applications. Formally the data on the graph can be expressed by a labeling function:

$$\rho : (V \cup E) \rightarrow w \mid w \in \mathbb{R}^+ \quad (2.30)$$

In this case, we assume that the data are positive numbers, as often it is a constraint imposed by algorithms. There are however no theoretical restrictions from graph theory that impose this constraint. The data graph and weighted graph can be expressed with the same labeling function, with just the vertices or the edges as input for the labeling function.

2.4.2 RDF graphs

Research Description Framework (RDF) graphs are directed graphs where edges have labels and where multiple edges may exist between a pair of nodes. Graphs that allow multiple edges between a pair of nodes are also called *multi-graphs*. They are a standard by the World Wide Web Consortium

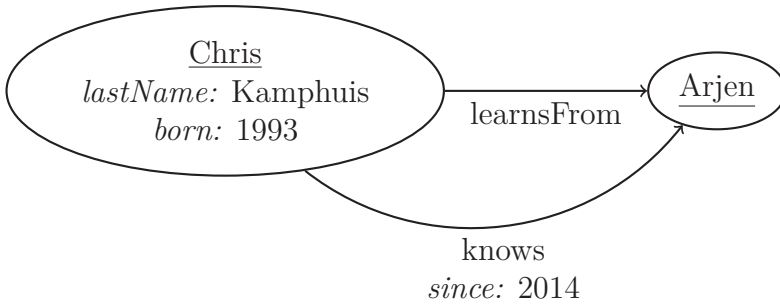


Figure 2.4: Example of a Property Graph.

(W3C). The idea behind RDF graphs is that they can be resolved by a collection of so-called triples. These triples contain a *subject*, *predicate*, and *object*. The predicate describes the relation (edge) between the subject (node) and the object (node); a collection of triples forms a graph. RDF graphs are widely used to represent knowledge graphs. These are widely used to store information about named *entities*; things in the real world like persons, locations, or organizations. As a rule of thumb, you could say that everything that could have a Wikipedia entry can be considered a named entity. Knowledge graphs have many applications in information retrieval. It is, for example, possible to generate entity cards [Hasibi et al., 2017a] on top of search engine result pages using knowledge graphs. Figure 2.3 shows an example of a small RDF graph.

2.4.3 Property Graphs

The final type of graph we discuss is the property graph. Property graphs allow everything we described before, while vertices can also have labels, and vertices and edges can have property-value pairs (multiple values per property are even allowed). A property-value represents some named data for either a vertex or edge. For example, if a node represents a person, the person’s age can be a property of the node. Depending on the type of property graph, the graph might also contain multiple vertex or edge labels. Figure 2.4 shows an example of a property graph. The node with label *Chris* has last name and birth year properties. The edge with label *knows* from the node *Chris* to the node *Arjen* has a timestamp as a property. Property graphs have high expressive power, and have become the preferred data model for graph databases.

Formally, to complement our previous definitions of vertices and edges, we assume the existence of three infinite sets; L , P , and A . L is an infinite

set containing vertex/edge labels, P is an infinite set containing property names, and A is an infinite set of atomic values. Given a set X , $\text{SET}^+(X)$ provides the set of all finite subsets of X . Using these definitions, we can define the following two functions that assign labels, and property-value pairs:

$$\lambda : (V \cup E) \rightarrow \text{SET}^+(L) \quad (2.31)$$

$$\sigma : (V \cup E) \times P \rightarrow \text{SET}^+(A) \quad (2.32)$$

In Chapter 4, we built on this representation to explore graph databases for IR.

2.5 Reproducible Science

In order to establish experimental results in science, it is essential that they can independently be verified. Many fields have had problems where the reproducibility of scientific studies fell short. In psychology, a large reproducibility study was carried out [Open Science Collaboration, 2015], based on a selection of 100 studies published in 2008. The goal was to reproduce the findings of all these studies, to determine the state of reproducibility of the field. Of the original 100 studies, 97 presented significant results. However, when trying to reproduce these results, only 35 out of 97 of these studies succeeded to reproduce significant results.

It is unclear why that many scientific studies were not reproducible. Nevertheless, many type II errors happened; failure to reject the null hypothesis while it is false. Scientific misconduct could explain some of these results, but it is unlikely that this is happening at such a large scale. All studies this project tried to reproduce were published in just three journals. These journals were peer-reviewed; only the scientific studies the reviewers thought were good enough were published. This process, of course, introduces a selection bias. Studies that show more impactful results have a higher probability of being published. This bias favors papers with type II errors, even when the scientific conduct was proper.

When a scientific study is published, people might interpret its results as fact. However, one should consider experimental results only as evidence of how the world might work. Then, whenever the results are reproduced, more evidence is gathered to support the findings. By reproducing science, our understanding of the world becomes more precise, and eventually, we can assume with high certainty a hypothesis to be true.

In the field of information retrieval, reproducibility has also been a topic of interest. Reproducibility is discussed several times in the dissertation,

and only some observations on the topic of reproducibility in our field are highlighted here. Before introducing these observations, it is important to clearly define what is meant by reproducibility. Within the Association for Computing Machinery (ACM), a distinction has been made between the concepts of repeatability, replicability, and reproducibility³:

- *Repeatability* is verification of results produced by the same research group, using the same resources. Confirming that when one reruns their experiments, the same results are produced. This might seem trivial, but when software is updated to some newer version, the results that the software produces could differ slightly.
- *Replicability* is the verification of results produced by a different research group, but using the same resources. This could, for example, be done by running the publicly available code for a research project and verifying its results. This is more tricky than it might look. Often, for example in scientific papers, the exact parameter settings of the software might be left out, making it hard to verify results correctly.
- *Reproducibility* is the verification of results produced by a different research group that uses its own (different) resources. Typically, in a reproduction study, the scientific study is done from scratch using the instructions presented in a paper. As many details might be left out, it is even harder to reproduce results exactly. However, if a reproduction study can confirm the results of previous work, this is a strong indication of the correctness of the results presented in that work.

In Chapter 3, we present empirical results about the BM25 algorithm. Many systems implement this algorithm, while the reported effectiveness results as measured through established metrics (a subset of metrics described in Section 2.2.3) vary substantially. How this can happen will be discussed. BM25, however, is a ranking method often used as a baseline method. When comparing newer algorithms to BM25, it makes quite a difference if its implementation reports a relatively low score, or a high one.

Armstrong et al. [2009] showed that many improvements in ranking algorithms presented throughout the years did not add up. Significant improvements were presented in many cases, where these were compared against weak baselines. Also, there was no upward trend of retrieval effectiveness, which would be expected if we find repeated improvements. More

³<https://www.acm.org/publications/policies/artifact-review-badging>, last accessed - September 2025

recently, Yang et al. [2019] showed new methods are still being compared against implementations of BM25 that have non-optimal hyperparameter settings. This leads to methods that are less effective than presented.

For this thesis, all software and data produced is publicly available and released on Zenodo⁴ in order to facilitate reproducible science, following the guidelines of the research data management policy of the Institute for Computing and Information Science of the Radboud University.⁵ Additionally, in some chapters, the importance of reproducibility will be highlighted and discussed in more detail.

⁴<https://zenodo.org>, last accessed - September 2025

⁵<https://www.ru.nl/rdm/>, last accessed - September 2025

Chapter 3

IR using Relational Databases

The Analytical Engine has no pretensions whatever to originate anything. It can do whatever we know how to order it to perform

Ada Lovelace - 1843

Abstract

Throughout the years, there have been many attempts to express information retrieval problems using relational databases. This chapter will highlight one of the latter attempts that revived the idea of expressing bag-of-words ranking functions using SQL. A prototype system that uses these expressions is presented, dubbed OldDog, after the work by Mühleisen et al. (*Old Dogs Are Great at New Tricks: Column Stores for IR Prototyping*). This system can be used for rapid IR prototyping and is especially helpful in the context of reproducible information retrieval research. Also, when researchers report that they used BM25, it is not always clear which variant they mean. Many tweaks to Robertson et al.’s original formulation have been proposed. Does this ambiguity “matter”? We attempt to answer this question with a large-scale reproducibility study of BM25,

considering eight variants implemented in the OldDog system. Experiments on three newswire collections show no significant effectiveness differences between them, not even for Lucene’s (often maligned) approximation of document length.

3.1 Introduction

Where information retrieval researchers commonly use inverted indexes as data structures, there is also a rich history of researchers using relational databases to represent the data in information retrieval systems. Different approaches in the literature present varying degrees of success. Given this context, we arrive at the first research question:

- *Research Question 1: What are the benefits of using relational databases for information retrieval?*

In order to answer this question, first, we review the history of using database systems for IR. Then, one of the latter attempts of using a relational database for information retrieval will be highlighted. Using this work, a prototype system is built, dubbed “OldDog”. This system will be used in a reproduction experiment, which compares several variants of BM25 with each other. This reproduction study confirms and rebuts previous findings from the literature and verifies that relational database systems are suited for running IR experiments.

3.2 Related work

3.2.1 Boolean retrieval

Perhaps the earliest work on using relational databases for information retrieval is the work by Schek and Pistor [1982]. In their work, the authors recognized that the relational data model is widely accepted as an interface to query structured data. However, it is inconvenient to use unstructured data like text. They proposed extending the relational model by allowing Non-First Normal-Form (NF²) relations. This extension allows for text queries to be more easily expressed. The systems that can be built in this language are still Boolean retrieval systems (as described in Section 2.2.2). At the time, that worked well, but scoring was not a feature considered. Similarly, Macleod [1991] compared the inverted index approach to using the relational model. Macleod showed how queries of the IBM STAIRS

system could be expressed using the relational model. These were, however, still Boolean queries, so scoring using uncertainty was not considered.

3.2.2 Probabilistic Relational Algebra

Fuhr [1996] recognized that where databases contain structured/formatted data, IR systems deal with unformatted data requiring uncertain inference. They propose to express this uncertainty using a probabilistic relational algebra (PRA) [Fuhr and Rölleke, 1997]. PRA can be considered an extension of standard relational algebra. The basic idea behind PRA is that tuples are assigned weights; the weight represents the probability that the tuple belongs to the relation. These probabilities give two advantages. Uncertainty of information can be expressed explicitly, and tuples representing answers to queries can be ordered by the weights representing this uncertainty. Tuples with high probabilities are ranked higher than those with lower ones. Although these extensions give advantages over Boolean retrieval, how to assign these probabilities to, for example, a document-term pair remains a question.

3.2.3 IR on top of a database cluster

Grabs et al. [2004] have proposed PowerDB-IR, developed to run IR applications on a scalable infrastructure. It should also be able to update the data quickly while retrieving up-to-date results. Grabs et al. achieve this by assigning every document to a category, e.g., sports or news, in their experiment. A dedicated node is created for every category, containing documents, inverted lists, and statistics tables. The system supports both single-category and multi-category searches. For a “single query” search, the following ranking score value is calculated:

$$\text{RSV}(d, q) = \sum_{t \in q} tf(t, d) \cdot idf(t)^2 \cdot tf(t, q) \quad (3.1)$$

Here $tf(t, d)$ is the term frequency of term t in document d , $idf(t)$ is the inverted document frequency of term t (which is squared in this formula), and $tf(t, q)$ is the term frequency of term t in the query text. Calculating this is straightforward: all statistics necessary are stored on the node decided to that category. However, when one wants to search on multiple (or all) categories, subscores need to be calculated for all relevant nodes before they can be aggregated to a final score. This approach is quite costly, as variances between executions times on the nodes are high, but this work may have proposed the first real IR in SQL approach.

3.2.4 Integrating DB + IR

Chaudhuri et al. [2005] also identified the need for systems that integrate database and IR functionalities. In their view, database systems need to be more flexible for scoring and ranking, while IR systems cannot adequately handle structured data and metadata. Chaudhuri et al. put together a list of seven requirements that a DB + IR system should be able to support, of which they identified the following three requirements as the most important:

1. *Flexible scoring and ranking.* It should be possible to customize the ranking function for different applications; a news search system probably needs different ranking functions and settings than a web search system.
2. *Optimizability.* Following standard database approaches, queries in a DB+IR system should have a query optimizer that considers the workload and the data characteristics. For example, when only one relevant result is sufficient, the system should be able to abort when a relevant document is found.
3. *Metadata and ontologies.* Other than metadata that describes data sources, metadata that is used for understanding information demands might be needed. This metadata could be, for example, an ontology or a lexicon used for more effective ranking strategies.¹

To build a system that can support these requirements, the authors identified four alternatives for designing a DB+IR system:

1. *On-top-of-SQL.* The IR functionalities are built on top of a SQL engine. The disadvantage of this approach is that it is challenging to customize efficient access for both IR and DB functionalities.
2. *Middleware.* In this approach, SQL and IR engines run simultaneously. The two disadvantages of using this approach are that the API needs to talk to two systems, which can have very different design philosophies, and the data needs to be shared between systems, incurring a large overhead and making it harder to combine both functionalities.
3. *IR-via-ADTs.* The third approach is building an IR system using abstract data types (ADT). The authors argue that this approach makes the system more customizable than the previous approaches.

¹Latent representations generated by large language models would have been a great example of this kind of metadata had their paper been written today.

However, the authors also note that query optimization in the presence of user defined functions is complicated. Also, when programmers need to work with such a system, it has the full complexity of SQL plus the complexity of working with ADTs.

4. *RISC*. The final approach is what the authors prefer; IR functionalities build on top of a relational *storage* engine, as described in an earlier work by them [Chaudhuri and Weikum, 2000]. The DB+IR systems should then be built on top of this engine. The storage-level core should be build “Reduced Instruction Set Computing” style (RISC).

Although the approaches described in this work are interesting, they do not provide prototypes to compare them. (The goal of this paper was to present a theoretical framework for tackling this problem.) Even though it is not the preferred option of Chaudhuri et al., this PhD thesis research has focused on the *On-top-of-SQL* approach, assuming that recent developments in database system engineering have overcome the shortcomings with respect to ranking efficiency.

3.2.5 Handwritten plans and Array Databases

Héman et al. [2006] participated in the TREC TeraByte track using the relational engine MonetDB/X100 [Boncz et al., 2005]. They were able to express ranking functions efficiently and effectively in this system. In their participation, they used BM25 as a scoring function. In order to reduce the amount of computation necessary for every document-term pair, the BM25 score was precalculated. A shortcoming in this approach is that the query plans were not generated from SQL, but handwritten (essentially the RISC approach of above). Having to handwrite queries makes the system challenging to use for IR researchers. Also, because all BM25 scores were precalculated (albeit with some compression), more storage was needed than when only the term frequencies would be saved.

The same research group [Cornacchia et al., 2008] also ran experiments on the TREC TeraByte track using the array database SRAM (Sparse Relational Array Mapping). SRAM automatically translates BM25 queries to run them on a relational engine (specifically X100). However, SRAM is quite an esoteric query language, only used by the researchers themselves, and no publicly available (prototype) system offers support for this approach.

Table 3.1: Results presented by Mühleisen et al. [2014]. MAP and P@5 on the ClueWeb12 collection are reported for five different systems that each claim to rank their documents using BM25. The table shows that only the two database systems (MonetDB and Vectorwise) achieve the same effectiveness score.

System	MAP	P@5
Indri [Strohman et al., 2005]	0.246	0.304
MonetDB [Boncz, 2002]	0.225	0.276
Vectorwise [Zukowski et al., 2012]	0.225	0.276
Lucene [Apache Software Foundation, 2013]	0.216	0.265
Terrier [Ounis et al., 2005]	0.215	0.272

3.2.6 Retrieval models only using SQL

In more recent work, Mühleisen et al. [2014] showed that the commonly used BM25 ranking function could be expressed easily using SQL. This was done similarly to Grabs et al. [2004]. In this work, the MonetDB [Boncz, 2002] and Vectorwise [Zukowski et al., 2012] systems were used, making the runtime much faster than Grabs et al.’s original results. Mühleisen et al. specifically focused on the retrieval efficiency of systems and compared the efficiency of inverted indexes with systems built on top of relational engines. They argued that instead of using a custom build IR system using an inverted index, researchers could store their data representations in a column-oriented relational database and formulate the ranking functions using SQL. They showed that their implementation of BM25 in SQL is on par in efficiency and effectiveness compared to systems that use an inverted index.²

There was an interesting observation in the paper to highlight: All the systems evaluated in this paper implement a ranking function they refer to as BM25. However, there was a substantial difference between the effectiveness scores produced by these systems, as shown in Table 3.1. The only two systems that achieved the same effectiveness score were the two database systems (MonetDB and Vectorwise).³

These results were surprising, as the authors took specific care to keep document preprocessing identical for all systems. Still, the observed differ-

²This was especially the case for the Vectorwise system.

³Although the same research group developed these two systems, they were completely separate projects.

Table 3.2: Results from the RIGOR workshop [Arguello et al., 2016]. MAP@1000 on the .GOV2 collection is reported for four different systems that run BM25. The table shows that all four implementations report a different effectiveness score.

System	MAP@1000
ATIRE [Trotman et al., 2012]	0.290
Lucene [Apache Software Foundation, 2013]	0.303
MG4J [Boldi and Vigna, 2005]	0.299
Terrier [Ounis et al., 2005]	0.270

Table 3.3: Results from the OSIRRC workshop [Clancy et al., 2019b]. AP, P@30, and NDCG@20 on the robust04 collection are reported for seven different systems that run BM25. As shown in the table, all implementations report (again) a different effectiveness score.

System	AP	P@30	NDCG@20
Anserini [Clancy et al., 2019a]	0.253	0.310	0.424
ATIRE [Trotman et al., 2012]	0.218	0.320	0.421
ielab [Scells and Zuccon, 2019]	0.183	0.261	0.348
Indri [Hauff, 2019]	0.239	0.300	0.404
OldDog [Kamphuis and de Vries, 2019b]	0.243	0.299	0.400
Pisa [Mallia et al., 2019]	0.253	0.312	0.422
Terrier [Câmara and Macdonald, 2019]	0.236	0.298	0.405

ence in MAP of 3% absolute was the largest deviation in the scores reported.

3.2.7 Reproducibility

The paper by Mühleisen et al. [2014] is not the only one that reports differences in effectiveness scores for BM25. In the SIGIR 2015 Workshop on Reproducibility, Inexplicability, and Generalizability of Results (RIGOR) [Arguello et al., 2016] and the Open-Source IR Replicability Challenge (OSIRRC) workshop [Clancy et al., 2019b] similar results can be observed. See Tables 3.2 and 3.3, respectively.

It is unclear why the results between these systems differ this much; many

explanations are possible. Examples include; differences in preprocessing,⁴ different hyperparameter settings, differences in the definition of the inverse document frequency (*idf*) used, or erroneous implementation of the ranking function. Using, for example, non-optimized hyperparameter settings can lead to considerable gaps in differences between effectiveness scores. Yang et al. [2019] showed that in many cases, new ranking methods had been proposed that compared the results of a newly proposed method to a non-fine-tuned version of BM25, making the results look better than they are. The choices for the BM25 hyperparameters are often left out of papers where it is the baseline compared against. As BM25 is often used as a baseline, it is crucial to understand why these differences exist and how they arise.

3.3 Prototype OldDog

As shown in Table 3.3, one of the workshop’s submissions was the prototype system developed for this PhD thesis research [Kamphuis and de Vries, 2019b]. This prototype is a software project to replicate and extend the database approach to information retrieval presented in Mühleisen et al. [2014]. The prototype was based on their work, so we dubbed it *OldDog*. OldDog uses column store database MonetDB [Boncz, 2002] for query processing. Mühleisen et al. produced the database tables to represent the data typically found in an inverted index, using a custom program run on Hadoop. Instead, we created a Lucene index using the Anserini tool suite developed by Yang et al. [2018a]. Anserini takes care of standard document preprocessing. From Anserini’s Lucene index, we extracted the data necessary to fill the tables. Three tables are constructed. One table contains the data that represents the documents, another one that represents the terms, and one that contains all data that relate terms to documents. To illustrate, the results for a document named “doc1” with the text “I put on my coat and green hat” shown in Table 3.4.

Using these tables, we can easily express bag-of-words ranking functions in SQL queries. The default ranking function implemented in our implementation of OldDog is the version of BM25 that had been proposed by Robertson et al. [1994].

⁴But this cannot explain the differences reported in Mühleisen et al., as they ensured preprocessing was the same for all systems.

Table 3.4: Example of tables representing the data in the OldDog system, the dict tables contains all term specific data, the terms table represents all the postings, and the docs table contains all document specific data.

(a) dict			(b) docs			(c) terms		
termid	term	df	docid	name	length	termid	docid	tf
1	put	1	1	doc1	8	1	1	1
2	coat	1				2	1	1
3	green	1				3	1	1
4	hat	1				4	1	1

3.3.1 Docker

For the submission to the OSIRRC workshop [Clancy et al., 2019b], we created a docker image of OldDog⁵ [Kamphuis and de Vries, 2019b]. Mühleisen et al. implemented a conjunctive variant of BM25 (all query terms have to be present in a document in order for a document to be considered relevant). When creating the submission for the workshop, we noticed that the effectiveness scores were substantially lower than those of other submissions. The retrieval effectiveness degraded more than we expected, considering the results in previous work. When removing the conjunctive constraint, the effectiveness of the system increased. So our prototype supports both conjunctive and disjunctive versions of BM25. Our entry in Table 3.3 presents the effectiveness scores of the disjunctive variant. The number of relevant documents per topic for this collection was likely low; decreasing the effectiveness of the system when imposing a conjunctive constraint.

3.3.2 Ease-of-Use

Having implemented BM25 in a database system enabled us to carry out some experiments quite easily that are more complicated when using an inverted index. Filtering out the terms with a large document frequency is easy, as all document frequencies are stored in one table. We updated the table removing the terms with large document frequency in only two lines

⁵<https://hub.docker.com/r/osirrc2019/olddog>, last accessed - September 2025

```
1 CREATE table dict AS SELECT * FROM odict WHERE df <= (  
2     SELECT 0.1 * COUNT(*) FROM docs  
3 );
```

Figure 3.1: This code updates the `docs` table such that all terms with a document frequency greater than a tenth of the collection size are removed.

of SQL, as shown in Figure 3.1. This approach could be an automatic way to remove stopwords from a collection. This filter was too strict to improve retrieval effectiveness but can easily be fine-tuned. For example, we could join against a table storing a traditional stopword list.

3.4 Variants of BM25

Having “OldDog” set up, we can run retrieval experiments using relational databases. As mentioned in the previous section, it is still unclear why the differences between implementations of BM25 of different systems were this big. Also, many different variants of BM25 have been proposed in the literature, that claim to be more effective than the original formulation. A study by Trotman et al. [2014] compared several variants and found that improvements presented in the literature cannot be reproduced. As we now have a system in which the BM25 formula is written directly in SQL, we can easily swap this version of BM25 with its proposed improvements: in the query we only need to change definition of the ranking formula. By using OldDog, we can ensure the data representation is the same when we compare these variants; the results will only reflect what the effects are of applying a different variant of BM25. This way, we can validate Trotman et al.’s findings.

In the following section we will first describe the original formulation of BM25, and then variants of BM25 proposed in the literature.

Robertson et al. [1994]

BM25 was developed within the probabilistic ranking framework 2.2.2. The original formulation consists of two parts. The first part is derived from the binary independence relevance model [Robertson and Zaragoza, 2009], which results in an approximation of the classical inverse document frequency (*idf*) for a query term t :

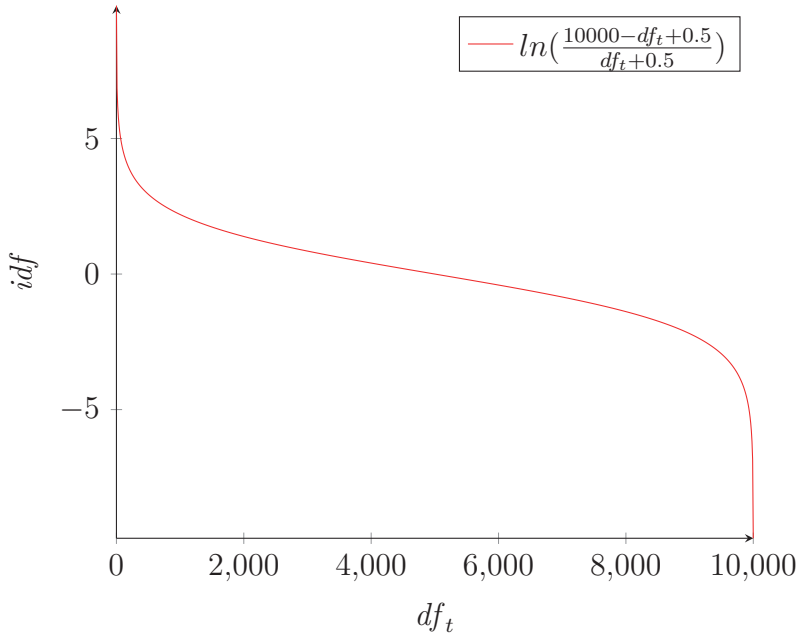


Figure 3.2: Inverse document frequency as used by Robertson et al. [1994]

$$w_t^{\text{IDF}} = \log \left(\frac{N - df_t + 0.5}{df_t + 0.5} \right) \quad (3.2)$$

where N is the collection size, and df_t are the number of collection documents containing query term t .

There is, however, a negative consequence of using this formula for weighting term importance. Assume a collection with 10,000 documents; then, it is possible to plot the idf for each term as shown in Figure 3.2. The figure shows that the idf score becomes negative when $df_t > \frac{N}{2}$. This happens for terms that appear in more than half of all documents, e.g.: “the” or “a”. Many systems do not consider these terms when searching by keeping a list of common words that can be ignored (stop words). However, when these words are considered, a negative idf would decrease the relevance scores of documents with the query term in the document – variations of BM25 have been proposed to deal with this anomaly, that are discussed in the following sections.

The second part of BM25 can be considered as a term frequency weighting tf . The two parts are multiplied to get something like the traditional term frequency-inverse document frequency weighting $tf \times w_i^{\text{IDF}}$. However,

the tf in BM25 is not just a linear factor; every additional term occurrence increases the ranking score value less than the previous one. For example, a term being present twice in a document versus once provides more information than a term being present ten times versus nine. To achieve this effect of “diminishing returns” of additional occurrences, the following formula defines the tf component of the ranking:

$$\frac{tf}{k + tf} \text{ where } k > 0 \quad (3.3)$$

This approach ensures that the term frequency does not increase linearly. In the final formulation of BM25, k is written as k_1 , because earlier versions of this ranking formula also had k_2 and k_3 parameters.

Lastly, a second component is added to correct for variation in document length. It is, however, unclear how one should deal with documents being longer than others; the document’s author can be verbose, in which case additional term occurrences do not provide extra information. On the other hand, a document can be lengthy because more relevant information is provided, and the document is more relevant than its shorter counterpart. For these reasons, the following soft length normalization is introduced:

$$(1 - b) + b \times \left(\frac{L_d}{L_{avg}} \right) \text{ with } 0 \leq b \leq 1 \quad (3.4)$$

The value of b controls the degree of length normalization. Combining these parts, including the correction for term frequency and length normalization, we get BM25 as initially proposed by Robertson et al. [1994]:

$$\sum_{t \in q} \log \left(\frac{N - df_t + 0.5}{df_t + 0.5} \right) \cdot \frac{tf_{td}}{k_1 \cdot \left(1 - b + b \cdot \left(\frac{L_d}{L_{avg}} \right) \right) + tf_{td}} \quad (3.5)$$

Lucene (default)

Lucene is a widely used open source search engine software library [Apache Software Foundation, 2013]. The variant of BM25 implemented in Lucene (as of version 8) introduces two main differences. As mentioned, the idf component of Robertson et al. [1994] is negative when $df_t > \frac{N}{2}$. To avoid negative values in all possible cases, Lucene adds a constant value of one before calculating the $\log$ value. Second, the document length used in the scoring function is compressed (in a lossy manner) to a one-byte value,

denoted L_{lossy} .⁶ With only 256 distinct document lengths, Lucene pre-computes for each possible length the value of:

$$k_1 \cdot \left(1 - b + b \cdot \left(\frac{L_{\text{lossy}}}{L_{\text{avg}}} \right) \right) \quad (3.6)$$

This results in fewer computations at query time. Equation (3.7) describes BM25 as implemented in Lucene:

$$\sum_{t \in q} \log \left(1 + \frac{N - df_t + 0.5}{df_t + 0.5} \right) \cdot \frac{tf_{td}}{k_1 \cdot \left(1 - b + b \cdot \left(\frac{L_{\text{lossy}}}{L_{\text{avg}}} \right) \right) + tf_{td}} \quad (3.7)$$

Lucene (accurate)

Equation (3.8) represents our attempt to measure the impact of Lucene’s lossy document length encoding. We implemented a variant that uses exact document lengths but is otherwise identical to the Lucene default.

$$\sum_{t \in q} \log \left(1 + \frac{N - df_t + 0.5}{df_t + 0.5} \right) \cdot \frac{tf_{td}}{k_1 \cdot \left(1 - b + b \cdot \left(\frac{L_d}{L_{\text{avg}}} \right) \right) + tf_{td}} \quad (3.8)$$

ATIRE [Trotman et al., 2012]

Equation (3.9) shows BM25 as implemented by ATIRE; it implements the *idf* component of BM25 as $\log(N/df_t)$, which also avoids negative values. We will investigate the difference of this formulation with that of Lucene later in Section 3.6. The TF component is multiplied by $k_1 + 1$ to make it look more like the classic RSJ weight [Robertson and Spärck Jones, 1976]; this does not affect the resulting ranked list, as all scores are scaled linearly with this factor.

$$\sum_{t \in q} \log \left(\frac{N}{df_t} \right) \cdot \frac{(k_1 + 1) \cdot tf_{td}}{k_1 \cdot \left(1 - b + b \cdot \left(\frac{L_d}{L_{\text{avg}}} \right) \right) + tf_{td}} \quad (3.9)$$

BM25L [Lv and Zhai, 2011c]

BM25L builds on the observation that BM25 penalizes longer documents too much when compared to shorter ones. The *idf* component differs to

⁶The lossy encoding is described in this ticket: <https://issues.apache.org/jira/browse/LUCENE-7730>, last accessed - September 2025

avoid negative values. The TF component is reformulated as follows:

$$\frac{(k_1 + 1) \cdot c_{td}}{k_1 + c_{td}} \quad (3.10)$$

with

$$c_{td} = \frac{tf_{td}}{1 - b + b \cdot \left(\frac{L_d}{L_{avg}}\right)} \quad (3.11)$$

being a normalized tf using pivoted length normalization [Singhal et al., 1996]. The c_{td} component is further modified by adding a constant δ , boosting the score for longer documents. The authors report using $\delta = 0.5$ for the highest effectiveness. Equation (3.12) presents the final formulation of BM25L:

$$\sum_{t \in q} \log \left(\frac{N + 1}{df_t + 0.5} \right) \cdot \frac{(k_1 + 1) \cdot (c_{td} + \delta)}{k_1 + (c_{td} + \delta)} \quad (3.12)$$

BM25+ [Lv and Zhai, 2011a]

BM25+, as shown in Equation (3.13), encodes a general approach for dealing with the issue that ranking functions unfairly prefer shorter documents over longer ones. Lv and Zhai proposed adding a lower-bound bonus when a term appears at least once in a document. The difference with BM25L is a constant δ to the TF component. The idf component is again changed to a variant that disallows negative values.

$$\sum_{t \in q} \log \left(\frac{N + 1}{df_t} \right) \cdot \left(\frac{(k_1 + 1) \cdot tf_{td}}{k_1 \cdot \left((1 - b) + b \cdot \left(\frac{L_d}{L_{avg}} \right) \right)} + \delta \right) \quad (3.13)$$

BM25-adpt [Lv and Zhai, 2011b]

BM25-adpt is an approach that varies k_1 per term (i.e., uses term specific k_1 values). In the original formulation of BM25, k_1 can be considered a hyperparameter that regulates the increase of score for additional occurrences of a term; k_1 ensures that every additional occurrence gets discounted as it provides less information than its previous. However, Lv and Zhai argued that this does not necessarily have to be the case. If there are many fewer documents with $r + 1$ occurrences versus r , it should provide more information than when the number of documents is almost the same. In order

to find the optimal term-specific k_1 value, the authors want to maximize the information gain for that particular query term. First, they identify the probability of selecting a document randomly from the collection that contains the term q at least once as:

$$p(1|0, q) = \frac{df_r + 0.5}{N + 1} \quad (3.14)$$

The probability of a term occurring one more time is defined as:

$$p(r + 1|r, q) = \frac{df_{r+1} + 0.5}{df_r + 1} \quad (3.15)$$

In both these formulas, constants 1 and 0.5 are added for smoothing to avoid zero probabilities. Then the information gain from r to $r + 1$ occurrences is computed by subtracting the initial probability:

$$G_q^r = \log_2 \left(\frac{df_{r+1} + 0.5}{df_r + 1} \right) - \log_2 \left(\frac{df_r + 0.5}{N + 1} \right) \quad (3.16)$$

Here df_r is not defined as a standard document frequency but based on the length normalized term frequency:

$$df_r = \begin{cases} |D_t|_{c_{td} \geq r-0.5} & r > 1 \\ df_t & r = 1 \\ N & r = 0 \end{cases} \quad (3.17)$$

In this case df_t is the “normal” document frequency, and c_{td} is the same as in BM25L (using the pivoted method for length normalization introduced by Singhal et al. [1996]):

$$c_{td} = \frac{tf_{td}}{1 - b + b \cdot \left(\frac{L_d}{L_{avg}} \right)} \quad (3.18)$$

This implies that for $r = 0$ that df_r is equal to the number of documents in the collection, and for $r = 1$ it is equal to the “normal” document frequency. Otherwise, it will be the number of documents with at least r occurrences of the term (rounded up) using the pivoted method c_{td} .

Next, the information gain is calculated for $t \in \{0, \dots, T\}$, until $G_q^t > G_q^{t+1}$. This threshold is chosen as a heuristic: When t becomes large, the estimated information gain can be very noisy. So T is chosen as the smallest value that breaks the worst burstiness rule [Church and Gale, 1995]

(the information gain starts decreasing). The optimal value for k_1 is then determined by finding the value for k_1 that minimizes the following equation:

$$k'_1 = \arg \min_{k_1} \sum_{t=0}^T \left(\frac{G_q^t}{G_q^1} - \frac{(k_1 + 1) \cdot t}{k_1 + t} \right)^2 \quad (3.19)$$

Essentially, this gives a value for k_1 that maximizes information gain for that specific term; k_1 and G_q^1 are then plugged into the BM25-adpt formula:

$$\sum_{t \in q} G_q^1 \cdot \frac{(k'_1 + 1) \cdot tf_{td}}{k'_1 \cdot \left((1 - b) + b \cdot \left(\frac{L_d}{L_{avg}} \right) \right) + tf_{td}} \quad (3.20)$$

We found that the optimal value of k_1 is not defined for about 90% of the terms. A unique optimal value for k_1 only exists when $t > 1$ while calculating G_q^t . For many terms, especially those with a low df , $G_q^t > G_q^{t+1}$ occurs before $t > 1$. In these cases, picking different values for k_1 has virtually no effect on retrieval effectiveness. For undefined values, we set k_1 to 0.001, the same as Trotman et al. [2014].

TF $l \circ \delta \circ p \times \text{IDF}$ [Rousseau and Vazirgiannis, 2013]

Rousseau and Vazirgiannis observed that in different variants of BM25, term frequency is normalized following several heuristics (i.e. non-linear gain, document length normalization, lower-bounding). When composing a ranking function, the order in which these heuristics are applied will change the ranking function. Testing different compositions of these heuristics, they found one composition that significantly outperformed traditional BM25:

$\text{TF}l \circ \delta \circ p \times \text{IDF}$, as shown in equation 3.23, models the non-linear gain of a term occurring multiple times in a document as:

$$1 + \log(1 + \log(tf_{td})) \quad (3.21)$$

To ensure terms occurring at least once in a document get boosted, the approach adds a fixed component δ , following BM25+. These parts are combined into the TF component using the pivoted method for length normalization [Singhal et al., 1996]:

$$c_{td} = \frac{tf_{td}}{1 - b + b \cdot \left(\frac{L_d}{L_{avg}} \right)} \quad (3.22)$$

The same IDF component used as in BM25+, which gives us $\text{TF}l \circ \delta \circ p \times \text{IDF}$:

$$\sum_{t \in q} \log \left(\frac{N + 1}{df_t} \right) \cdot (1 + \log(1 + \log(c_{td} + \delta))) \quad (3.23)$$

3.5 Experiments

This section presents an empirical evaluation of the impact of the different choices of BM25 as described in Section 3.4. Our experiments were conducted using Anserini (v0.6.0) on Java 11 to create an initial index, and subsequently using relational databases for rapid prototyping, using “OldDog” [Kamphuis and de Vries, 2019b] after Mühleisen et al. [2014]; following that work, we use MonetDB as the column-oriented analytic database. Evaluations with Lucene (default) and Lucene (accurate) were performed directly in Anserini; the latter was based on previously-released code that we updated and incorporated into Anserini.⁷ The inverted index was exported from Lucene to OldDog, ensuring that all experiments share the same document processing pipeline (e.g., tokenization, stemming, stopwords removal). Specifically, Anserini uses the Lucene implementation of the porter stemmer [Porter, 1980], and Lucene’s default stopwords list for the English language. While exporting the inverted index, we precalculate all k_1 values for BM25-adpt as suggested by Lv and Zhai [2011b]. As an additional verification step, we implemented both Lucene (default) and Lucene (accurate) in OldDog and compared the results to the output from Anserini. We can confirm that the results are the same, setting aside unavoidable differences related to floating point precision. All BM25 variants are then implemented in OldDog as minor variations on the original SQL query provided in Mühleisen et al. [2014]. The term-specific parameter optimization for the *adpt* variant was already calculated during the index extraction stage, allowing us to upload the optimal (t, k) pairs and directly use the term-specific k values in the SQL query. The advantage of our experimental methodology is that we did not need to implement a single new ranking function from scratch.

The experiments use three TREC newswire test collections: TREC Disks 4 and 5, excluding Congressional Record, with topics and relevance judgments from the TREC 2004 Robust Track (Robust04) (topics 301-450 & 601-700); the New York Times Annotated Corpus, with topics and relevance judgments from the TREC 2017 Common Core Track (Core17); the TREC Washington Post Corpus, with topics and relevance judgments from the TREC 2018 Common Core Track (Core18). Following standard experimental practice, we assess ranked list output in terms of average precision (AP) and precision at rank 30 (P@30). The parameters shared by all models are set to $k_1 = 0.9$ and $b = 0.4$, Anserini’s defaults, derived from Trotman et al. [2012]. The parameter δ is set to the value reported as

⁷<http://searchivarius.org/blog/accurate-bm25-similarity-lucene>, last accessed - September 2025

Table 3.5: Effectiveness scores different BM25 variants. All were implemented as SQL queries, so the underlying data representations are the same.

	Robust04		Core17		Core18	
	AP	P@30	AP	P@30	AP	P@30
Robertson et al.	.2526	.3086	.2094	.4327	.2465	.3647
Lucene (default)	.2531	.3102	.2087	.4293	.2495	.3567
Lucene (accurate)	.2533	.3104	.2094	.4327	.2495	.3593
ATIRE	.2533	.3104	.2094	.4327	.2495	.3593
BM25L	.2542	.3092	.1975	.4253	.2501	.3607
BM25+	.2526	.3071	.1931	.4260	.2447	.3513
BM25-adpt	.2571	.3135	.2112	.4133	.2480	.3533
$TF_{\log op} \times IDF$	.2516	.3084	.1932	.4340	.2465	.3647

best in the corresponding source publication.

All experiments were run on a Linux desktop (Fedora 30, Kernel 5.2.18, SELinux enabled) with four cores (Intel Xeon CPU E3-1226 v3 @ 3.30 GHz) and 16 GB of main memory; the MonetDB 11.33.11 server was compiled from source using the `--enable-optimize` flag.

3.6 Results

Table 3.5 shows the effectiveness scores of the different BM25 variants. The observed differences in effectiveness are small and can be entirely attributed to variations in the scoring function; our methodology fixes all other parts of the indexing pipeline (e.g., tag cleanup, tokenization, and stopwords). Both an ANOVA and Tukey’s HSD show no significant differences between any variant on all test collections. These results confirm the findings of Trotman et al. [2014]: effectiveness differences are unlikely an effect of the choice of the BM25 variant. Across the IR literature, we find differences due to seemingly less impactful settings (such as the choice of stopwords [Dolamic and Savoy, 2010]) tend to be larger than the differences we observe here. Although we find no significant improvements over the original [Robertson et al., 1994] formulation, using a variant of BM25 that avoids negative ranking scores might still be worthwhile.

You might have caught that the effectiveness scores of ATIRE and Lucene (accurate) are the same. This is not a mistake. As explained, the

Table 3.6: Average retrieval time per query in ms: Anserini (top) and OldDog (bottom)

	Robust04	Core17	Core18
Lucene (default)	52	111	120
Lucene (accurate)	55	115	123
Robertson et al.	158 ± 25	703 ± 162	331 ± 96
Lucene (default)	157 ± 24	699 ± 154	326 ± 90
Lucene (accurate)	157 ± 24	701 ± 156	324 ± 88
ATIRE	157 ± 24	698 ± 159	331 ± 94
BM25L	158 ± 25	697 ± 160	333 ± 96
BM25+	158 ± 25	700 ± 160	334 ± 96
BM25-adpt	158 ± 24	700 ± 157	330 ± 92
TF _{l_od_op} × IDF	158 ± 24	698 ± 158	331 ± 96

$k_1 + 1$ in ATIRE scales the scores linearly and does not affect the ranking. So the only difference that can change the effectiveness scores is the different *idf* functions. However, these are practically the same, especially when a collection has a large number of documents (N):

$$\log \left(\frac{N}{df_t} \right) = \log \left(\frac{N - df_t + df_t}{df_t} \right) \quad (3.24)$$

$$= \log \left(\frac{N - df_t}{df_t} + \frac{df_t}{df_t} \right) \quad (3.25)$$

$$= \log \left(\frac{N - df_t}{df_t} + 1 \right) \quad (3.26)$$

$$\approx \log \left(\frac{N - df_t + 0.5}{df_t + 0.5} + 1 \right) \quad (3.27)$$

Switching our attention from effectiveness to efficiency, Table 3.6 presents the average retrieval time per query in milliseconds (without standard deviation for Anserini, which does not report time per query). MonetDB uses all cores for inter- and intra-query parallelism, while Anserini is single-threaded.

Comparing Lucene (default) and Lucene (accurate), we find negligible differences in effectiveness. However, the differences in retrieval time are also negligible, which calls into question the motivation behind the original length approximation. Currently, the similarity function and, thus, the

document length encoding are defined at index time. Storing exact document lengths would allow for different ranking functions to be swapped at query time more effortlessly, as no information would be discarded at index time. Accurate document lengths might additionally benefit downstream modules that depend on Lucene. We suggest that Lucene might benefit from storing exact document lengths.

3.7 Conclusion

In summary, the previous sections describe a double reproducibility study. The study methodologically validated the usefulness of databases for IR prototyping and performed a large-scale study of BM25 to confirm the findings of Trotman et al. [2014]. It does not seem to matter which of the multitude of BM25 variants is used. Furthermore, to return to our research question, we can conclude that using relational databases for information retrieval is beneficial. Because data processing and storage are separated in relational databases, comparing different ranking functions in a relational system is much easier compared to a system that uses an inverted index. The work by Mühleisen et al. [2014] also confirmed that relational databases could be as efficient as inverted indexes in retrieval tasks. In short, databases have use cases in which they are easier to work with, while it remains possible to obtain efficient retrieval systems.

Chapter 4

From Tables to Graphs

Day by day, however, the machines are gaining ground upon us; day by day we are becoming more subservient to them

Samuel Butler - 1863

Abstract

This chapter introduces GeeseDB. GeeseDB is a Python toolkit for solving information retrieval research problems, that leverages graphs as data structures. It aims to simplify information retrieval research by allowing researchers to formulate graph queries through a graph query language. GeeseDB is built on top of DuckDB, an embedded column-store relational database for analytical workloads. GeeseDB is available as an easy-to-install Python package. In only a few lines of code, users can create a first-stage retrieval ranking using BM25. Queries read and write Numpy arrays and Pandas dataframes at negligible data transformation cost. Therefore, the results of a first-stage ranker expressed in GeeseDB can be used in various stages in the ranking process, enabling all the power of Python machine learning libraries with minimal overhead.

4.1 Introduction

In recent years there has been a lot of exciting new information retrieval research that uses non-text data to improve the efficacy of search applications. Examples of these are learning-to-rank or ranking with the use of knowledge graphs. All these research directions have improved search systems' effectiveness by using more diverse data. Although more diverse data sources are considered, these systems are often implemented through a coupled architecture: first-stage retrieval is often carried out with different software than that used in later retrieval stages, where novel reranking techniques tend to be used. In our view, researchers could benefit from a system where retrieval stages are more tightly coupled, which facilitates the exploration of how to use non-content data for ranking and serves the data in a format suitable for reranking with, e.g., transformers [Gao et al., 2021, Luan et al., 2021, Lin et al., 2020b], tree-based methods [Burgess, 2010] or graph-based methods [Page et al., 1999, Hasibi et al., 2016, Balog, 2018, Dalton et al., 2014].

The previous chapter demonstrates how relational databases can be used for information retrieval problems. Integrating alternative data sources into search systems is easier when using a relational database instead of an inverted index. In the case of information retrieval, graph data can often be used to improve search effectiveness. One of the most famous examples where graphs help information retrieval is the PageRank algorithm [Page et al., 1999]. Although it might be possible to express ranking over graph data in relational databases, they are not designed with graph structures in mind.

The data management community has shown much interest in graph databases in recent years. Graph databases are different from relational databases in that graphs are the main representation for managing the data, instead of the usual relations. Graphs can be considered a better abstraction for real-world data than relations. Graphs focus much more on concepts (nodes in a graph) and how they relate to each other (edges in a graph), while in a relational database this information is implicit, requiring knowledge of the schema. From the types of graphs discussed in Chapter 2, we focus on using *property graphs* in this chapter. This type of graph contains nodes and directed edges, which can be labeled; nodes and edges can also have associated key-value pairs.

The research goal in this chapter is two-fold; firstly, we develop a system that can search efficiently in information represented as graphs. In the previous chapter, we showed that relational databases could search efficiently and help reproducible research. However, these systems do not support

graph data structures directly in their interface to the user (e.g., the query language). Systems built on inverted indexes use coupled architectures, which can introduce unnecessary overhead.

Secondly, we want to create a system where both first and second-stage retrieval are directly supported. In IR research, multi-stage retrieval systems contain components that are not run in the same ecosystem; introducing overhead in retrieval efficiency. So in this chapter, the prototype system GeeseDB (Graph Engine for Exploration and Search) is introduced; it tries to leverage the same techniques for search as relational databases do while also naturally being able to support graphs. This leads to the main research question for this chapter:

- *Research Question 2: Can we extend the benefits from using relational databases for information retrieval to using graph databases while being able to express graph-related problems easier?*

In order to answer this research question, we will use our own prototype system GeeseDB. Before introducing GeeseDB, we first look into work that combine two systems for two-stage retrieval experiments.

4.2 Related work

In modern IR it is not uncommon to use coupled architectures. In many situations a ranking model is used that is efficient, but only reasonably effective. BM25 can be considered a model that can be calculated cheaply providing a good score for a more advanced (but less efficient) retrieval model that could achieve higher effectiveness. In such cases, BM25 can be used to calculate an initial top- k ranking that contains (presumably) all relevant documents. Then, a more expensive model only has to calculate the final ranking scores over the top- k documents. k is orders of magnitude smaller than the collection size.

In this section we start discussing methods that implement multi-stage approaches, then systems that implement them are highlighted. In particular there will be a focus on the coupled-architectures aspect of these systems. Finally, some proposed methods will be discussed to deal with cons arising from coupled architectures.

4.2.1 Learning To Rank

In learning-to-rank (LTR), multi-stage retrieval approaches are widely used; LTR models are often not efficient enough to compute a relevance score

value for all documents in the collection. In these cases, for example, a decision tree based ensemble like LambdaMART [Burges, 2010] is used to rerank only the top- k documents retrieved by BM25. Also, these models take input that is usually not stored in an inverted index. So, in order to use these models, the output from the inverted index needs to be combined with other data.

Deveaud et al. [2014] showed that for the TREC Contextual Suggestion track [Dean-Hall et al., 2014], reranking using user-based features significantly improved retrieval effectiveness over a baseline language modeling approach. These features show that data not present in the document’s text can help the retrieval effectiveness of systems; metadata helps to rank. Macdonald et al. [2012] showed examples of features that have been effective for learning to rank, many of these features are non-textual. Some influential features are: click count, click entropy, and the number of displayed results in a session. Agichtein et al. [2006] showed that for web search, incorporating user behavior features significantly improves retrieval effectiveness.

4.2.2 Dense Retrieval

Where traditionally, search was carried out using inverted indexes, neural retrieval has become more prevalent in recent years. When neural retrieval [Lin et al., 2022] is implemented as a vector search approach, it is referred to as dense retrieval. Consider, for example, the work by Gao et al. [2021]; in their work, they proposed `clear`, a model that aims to complement lexical models with a dense retrieval model. In their study, they use BM25 as the lexical model. The BM25 scores are calculated using an inverted index (Anserini [Yang et al., 2018a]), while the dense retrieval model is a fine-tuned BERT bi-encoder. This dense model calculates a score using a different maximum inner-product search (MIPS) system (FAISS [Johnson et al., 2019]) to calculate a similarity score for the query document pairs. Then scores are weighted and summed to calculate a new ranking score value.

In a similar study, Luan et al. [2021] compared several methods on the MS MARCO document and passage datasets. In their work, one of the two best models on the passage dataset was to combine a lexical model with a BERT-based bi-encoder. The scores were combined by retrieving the top- k documents for both the lexical model and the bi-encoder. In order to do this, an inverted index (Anserini) is used to retrieve the documents for the lexical model, while the bi-encoder uses a MIPS system (ScaNN [Guo et al., 2020]).

A third example, is the study of Lin et al. [2020b] that experiments with

distilled dense representations. They do an extensive comparison study, and in their work, the best results are reported by either multi-stage or hybrid dense + sparse retrieval methods. The multi-stage retrieval method used BM25 in combination with BERT-large. Although effective, it is a method with a high latency and high energy cost because of the cross-encoder used. A TCT-ColBERT (Tightly-Coupled Teacher - ColBERT) bi-encoder with doc2query-T5 sparse retrieval was the best hybrid method. This sparse retrieval method is a BM25 retrieval method that expands the documents with T5 language model questions [Raffel et al., 2020]. These questions were generated by letting the T5 language model generate questions that the passage might answer. Sparse retrieval is carried out by Anserini, and dense retrieval by FAISS.

4.2.3 Knowledge Graphs and Entities

Knowledge graphs have been used a lot as well in IR research. Hasibi et al. [2016] proposed the Entity Linking incorporated Retrieval (ELR) approach. The idea behind this concept is that when entity annotations are available for queries and documents, they can be incorporated into the ranking model. Their publication only focuses on entity retrieval, but the method can generalize to ad-hoc document and passage retrieval. The entities are incorporated by creating a dedicated index for entities on which entity scores can be calculated. These are then merged by the document scores calculated using the regular index. By incorporating entity scores, ELR has been shown to increase effectiveness scores for multiple baseline methods.

Dalton et al. [2014] used entity features for query expansion. One of the features was calculated by tagging queries with an entity-linking system. When entities were found, information from the knowledge base linked to them was included in query expansion. Alternative names of the entity were also included in the query expansion. Then, this expanded query provided a score for all documents. Another feature they calculate is obtained by ranking the entities in the knowledge base using the original query text. Then the data (e.g., words in the descriptions) included in the knowledge base can be included as terms for query expansion. Multiple other features using entities were also calculated,

For additional background related to entities in search, we refer to Balog [2018], an excellent book on entity-oriented search, that surveys many methods that use entities for search.

Basically, in order to include entities, a knowledge graph needs to be queried. This tends to be done by a different system than the system used

for regular retrieval. If we could integrate IR in graph databases, a single system would solve the complete process.

4.2.4 Current approaches

Recently, retrieval methods using neural networks have become state-of-the-art, mainly using large language models. In this section the two most commonly used research systems are described, highlighting how they with neural approaches to IR. PyTerrier [Macdonald et al., 2021] and Pyserini [Lin et al., 2021] are modern research systems that are widely used for information retrieval research. These systems are built as inverted index systems, with a separate component to handle neural retrieval models.

PyTerrier

PyTerrier by Macdonald et al. [2021], is a Python extension of Terrier [Ounis et al., 2005]. Terrier is an open-source search engine system written in Java. It implements state-of-the-art indexing and retrieval techniques on top of an inverted index. It is a system suited for rapidly developing retrieval experiments concerning document collections with many documents. The PyTerrier extension was developed to express complex IR pipelines in Python. Using overloaded Python operators (e.g. » for passing the output of one retrieval building block to the input of another) directly, it is possible to set up an IR pipeline using PyTerrier in only a few lines of code.

PyTerrier was expanded to support state-of-the-art large language model reranking approaches and dense retrieval methods. As the PyTerrier framework is developed for the Python ecosystem, other (re-)ranking methods that are implemented in Python can readily work with PyTerrier. The PyTerrier system is ideal for modern information retrieval research, that often concerns reranking through learned methods, as they are often developed in Python.

Note, however, that PyTerrier uses a coupled architecture. First-stage retrieval is carried out using Terrier in Java (e.g., a BM25 run), and then the resulting top documents are reranked in Python. For this, some overhead is introduced because of data transfer between processes.

Pyserini

Pyserini was developed by Lin et al. [2021], which is a Python extension of Anserini [Yang et al., 2018a]. Pyserini and Anserini have a lot in common with PyTerrier and Terrier. Just like Terrier, Anserini is developed in Java, but on top of Apache Lucene. Anserini is developed as a system with solid

support for reproducible research. Anserini provides reproducible baselines for many retrieval benchmarks that can be run with minimal investment of resources.

Pyserini is developed as a Python wrapper around the Anserini retrieval system. Within the Python ecosystem, it is possible to access the Anserini internals and replicate the same experiments implemented in Anserini. Similarly to PyTerrier, Pyserini also contains features that are not available directly in Anserini; ranking methods that make use of large language models, such as dense retrievers and BERT-based cross encoders. As these methods are generally developed in Python, it is much easier to support them with Pyserini compared to Anserini.

Anserini is included in Pyserini as a JAR file, so when using Pyserini for first-stage retrieval, data needs to be retrieved from JAVA before it can be materialized in Python. The setup is similar to that of Pyterrier.

4.2.5 Proposed approaches

Separation of the Logical and Physical Model

As shown in the previous sections, dense and sparse retrieval methods are often implemented using different systems, which introduces the coupled architecture. Lin [2022] recognizes that sparse retrieval models tend to be implemented using inverted indexes and dense retrieval methods with approximate nearest neighbor systems. Although these systems are different, the input and the output are the same: the input is a query, and the output is a ranked list of the top- k documents the ranking methods deem the most relevant.

In order to be able to formalize these similarities and distinctions more clearly, Lin proposes a distinction between the logical and physical retrieval models. Logical models specify encoders that map documents and queries to a representation, and a comparison function that defines how a ranking score value can be computed from comparing these representations. Physical models define how to create a top- k ranking from a large corpus given a query. The most straightforward physical approach would simply apply the logical model to every query-document pair, but through optimizations it might not necessary to compare all documents to the query. For example, dense retrieval methods often make use of the HNSW datastructure for nearest neighbor search [Johnson et al., 2019] which uses a greedy search method that might miss some relevant documents.

A Graph Query Language for IR

It would be ideal for a system to encapture these different ways of approaching IR research natively; work with multiple stages in the ranking process, but also reason about entities and knowledge graphs. Kamphuis and de Vries [2019a] proposed to use a graph query language for information retrieval problems. In the context of separating logical and physical models, a graph query language can be interpreted as a logical model, while the engine that implements the graph query language can be considered the physical model.

If it is possible to express both first stage and second stage retrieval in such a query language, both stages would be computed by the same physical model, automatically ensuring a coupled architecture is not necessary. Having a graph query language available, ensures that more complicated retrieval strategies can be expressed compared to when a relational query language is used.

So ideally we develop a system with the following two capabilities; the system supports first retrieval stage directly, and graph queries over the output of this first stage retrieval can be processed without a coupled architecture. In order to fulfill these needs, we developed GeeseDB in a follow up study.

4.3 GeeseDB

GeeseDB¹ is a prototype Python toolkit for information retrieval that leverages graphs as data structures, allowing metadata and graph-oriented data to be easily included in the ranking pipeline. The toolkit is designed to quickly set up first-stage retrieval and make it easy for researchers to explore new ranking models.

In short, GeeseDB aims to provide the following functionalities:

- GeeseDB is an easy-to-install, self-contained Python package available through PyPI with as few as possible dependencies. It contains topics and relevance judgments for several standard IR collections out-of-the-box, allowing researchers to develop new ranking models quickly.
- First stage (sparse) retrieval is directly supported. In only a few lines of code, it is possible to load documents and create BM25 rankings.

¹<https://github.com/informagi/geesedb>, last accessed September 2025

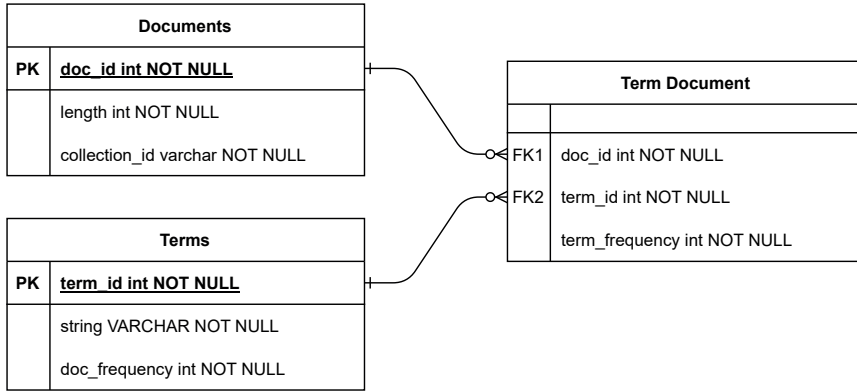


Figure 4.1: Database schema by Mühleisen et al. for full text search in relational databases

- Data is served in a usable format for later retrieval stages. GeeseDB allows directly running queries on Pandas data frames for efficient data transfer to sequential reranking algorithms.
- Easy data exploration is supported through querying data with SQL, but more interestingly, using a graph query language (based on the Cypher query language), making exploring new research avenues easier. This prototype supports a subset of the graph query language Cypher, similar to the property graph database model query language as described by Angles [2018].

4.3.1 Design

At the core of GeeseDB lies the full-text search design presented by Mühleisen et al. [2014] as discussed already in Chapter 3. This work proposes a column-store database for IR prototyping, which uses the database schema described in Figure 4.1, consisting of three database tables. Using these three tables, the authors show that BM25 can be easily expressed as a SQL query with latencies on par with custom-built IR engines. In GeeseDB, we use the same relational schema for full-text search. Instead of seeing the document data and term data as tables that relate to each other through a many-to-many join table, it is also possible to consider this schema as a bipartite graph. In this graph, documents and terms are considered nodes connected through edges. If a term occurs in a document, an edge exists between that term and the document. GeeseDB uses the data model of property graphs labeled

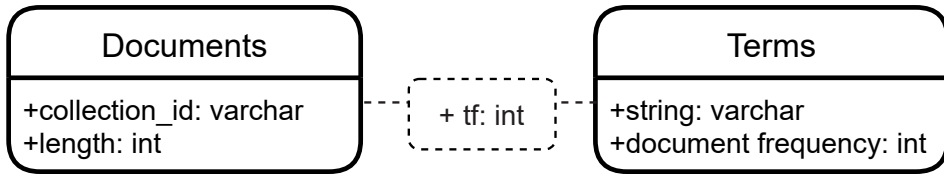


Figure 4.2: Graph schema representing bipartite document-term graph

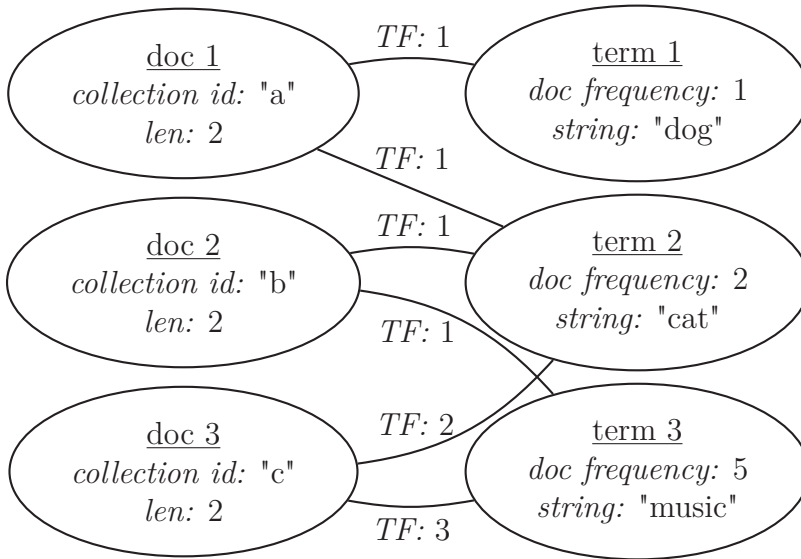


Figure 4.3: Example term-document graph that maps to relational database schema

multigraphs where both edges and nodes can have property-value pairs. The database schema, as described in Figure 4.1, would then translate to the property graph schema shown in Figure 4.2. A small example of a graph represented by this schema is shown in Figure 4.3. Document nodes contain document-specific information (i.e., document length and the collection identifier), term nodes contain information relevant to the term (i.e., the term string and the term's document frequency), and the edges between document and term nodes contain term frequency information (i.e., how often is the term mentioned in the document represented the respective nodes it connects).

If one wants to, for example, also store position data, this graph can easily be extended to a graph where the edges store term positions. If a term appears multiple times in a document, the property graph model will allow

multiple edges between two nodes. The graph schema that we described by Figure 4.2 maps one-to-one to the relational database schema described by Figure 4.1, when we represent nodes by standard relational tables that represent specific data units (terms, documents), while many-to-many join tables represent the edges. So, even though we think of the data as graphs, they are still represented as relational tables in the backend. When using GeeseDB for search, we at least expect the document-term graph to be present. New node types can be introduced to explore new search strategies. In GeeseDB we assume one database per collection.

Backend

GeeseDB is built on top of DuckDB [Raasveldt and Mühleisen, 2019], an in-process SQL OLAP (analytics optimized) database management system. DuckDB is designed to support analytical query workloads. It aims explicitly to process complex, long-running queries where a significant portion of the data is accessed, conditions matching the case of IR research. DuckDB has a client Python API which can be installed using `pip`.² Afterward, it can be used directly. DuckDB has a separate API built around NumPy³ and Pandas,⁴ providing NumPy/Pandas views over the same underlying data representation without incurring data transfer (usually referred to as “zero-copy” reading). Pandas DataFrames can be registered as virtual tables, allowing to query the data in Pandas DataFrames directly. GeeseDB inherits all these functionalities from DuckDB.

As DuckDB is a database management system, we can execute analytical SQL queries on tables containing our data, including the BM25 rankings described by Mühleisen et al. [2014]. By default, the BM25 implementation provided with GeeseDB implements the disjunctive variant of BM25 instead of the conjunctive variant they used. Although the conjunctive variant of BM25 can be calculated more quickly, the differences between effectiveness scores are noticeable on smaller collections. Also, disjunctive would mostly likely better match the mental model of our IR researchers as GeeseDB users. We only support the original formulation of BM25 by Robertson et al. [1994]. However, supporting other versions of BM25 as described in the previous chapter is trivial.

²<https://duckdb.org/docs/installation/?version=stable&environment=python>, last accessed - September 2025

³<https://numpy.org/>, last accessed - September 2025

⁴<https://pandas.pydata.org/>, last accessed - September 2025

```

1 MATCH (d:docs)-[]-(a:authors)-[]-(d2:docs)
2 WHERE d.collection_id = "96ab542e"
3 RETURN DISTINCT d2.collection_id

```

Figure 4.4: An example cypher query that finds all documents that were written by the same author that wrote the document with the `collection_id` “96ab542e”

```

1 SELECT DISTINCT d2.collection_id
2 FROM docs AS d2
3 JOIN doc_author AS da2 ON (d2.collection_id = da2.doc)
4 JOIN authors AS a2 ON (da2.author = a2.author)
5 JOIN doc_author AS da3 ON (a2.author = da3.author)
6 JOIN docs AS d ON (d.collection_id = da3.doc)
7 WHERE d.collection_id = '96ab542e'

```

Figure 4.5: SQL query that corresponds to the graph query described in Figure 4.4.

4.3.2 Graph Query Language

GeeseDB distinguishes itself from alternatives, database-backed systems (OldDog) [Kamphuis and de Vries, 2019b], or native systems (Anserini [Yang et al., 2018a], Terrier [Ounis et al., 2005]) by offering a graph query language, based on Cypher [Francis et al., 2018]. Where SQL queries start with what will be returned, a Cypher queries ends with it in the return clause. In Cypher the Match pattern is used to find specific graph patterns in the data. This can, for example, be a specific node of a certain type, or a graph pattern that describes a subgraph. GeeseDB implements Cypher’s basic graph pattern-matching queries for retrieving data. An example of a graph query supported by GeeseDB is presented in Figure 4.4. This co-authorship query finds all documents written by the same authors as those who wrote document “96ab542e”. For comparison, Figure 4.5 illustrates the same query represented in SQL; much more complex than the Cypher version, due to the join conditions that have to be made explicit. In order to connect the “docs”-table with the “authors”-table two joins are needed; reconnecting the “docs” table again introduces two more joins.

GeeseDB supports the following Cypher keywords: `MATCH`, `RETURN`, `WHERE`, `AND`, `DISTINCT`, `ORDER BY`, `SKIP`, and `LIMIT`. Instead of using `WHERE`

```

1 MATCH (d:docs {d.collection_id: "96ab542e"})
2 RETURN d.len

```

Figure 4.6: Graph query where the length of document with `collection_id` is returned.

to filter data, it is also possible to use graph matching patterns that include filters; as shown in Figure 4.6; the query returns the length of document “96ab542e”. Here the filter is defined between the curly braces. Everything that is not directly supported yet by our implementation can, of course, still be expressed in SQL, which is fully supported.⁵ In order to know how to join nodes to each other if no edge information has been provided, GeeseDB stores information on the graph schema. This way, GeeseDB knows how nodes relate to each other through which edges.

4.3.3 Usage

GeeseDB comes as an easy-to-install Python package that can be installed using pip, the standard package installer for Python:

```
$ pip install geesedb
```

After installing GeeseDB, we can immediately start using it. It is also possible to install the latest commit by installing the latest version directly from GitHub. As an example, we will show how to use GeeseDB for the background linking task of the TREC News Track [Soboroff et al., 2019]. The goal of this task is: *Given a news story, find other news articles that can provide meaningful context or background information.* These articles can then be recommended to the reader to help them understand the context in which these news articles occur. The collection used for this task is the Washington Post V3 collection⁶ released for the 2020 edition of TREC. It contains 671,945 news articles published by the Washington Post between 2012 and 2020, and 50 topics with relevance assessments (topics correspond to collection identifiers of documents for which relevant data has to be found). The articles in this collection contain valuable metadata; in particular, we will investigate how to best use the article authorship information. We extracted 25,703 unique article authors, where it is possible that multiple

⁵GeeseDB supports the graph queries by translating them to their corresponding SQL queries. After all, both nodes and edges are just tables in the backend.

⁶<https://trec.nist.gov/data/wapost/>, last accessed - September 2025

authors co-wrote a news article. We also annotate documents with entity information which was obtained by using the Radboud Entity Linker [van Hulst et al., 2020]. In total 31,622,419 references to 541,729 unique entities were found. The links also contain mention and location information, as well as the entity type (`ner_tag`) found by the linker’s entity recognition module (the entity type is part of a link, as the entity linker can assign different tags to the same entity.) Figure 4.7 illustrates the data schema that we use for the background linking task.

4.3.4 Indexing and Search

In order to start, a database containing at least the document and term information needs to be created. Figure 4.8 shows how the data can be easily loaded using CSV files.

Instead of loading the data from CSV files, it is also possible to load the text data directly using the CIFF format for information retrieval data exchange [Lin et al., 2020a]. GeeseDB also has functionalities to create CSV files from the CIFF format. Authorship information and entity links are loaded similarly. After loading the data, we can easily create a BM25 ranking for ad hoc search in the Washington Post collection, as shown in Figure 4.9.

For the background linking task, however, we do not have regular topics; we only have the document identifiers of the documents for which we need to find relevant background info. In order to search for relevant background reading, queries that represent our information need have to be constructed. A common approach uses the top- k TF-IDF terms of the source article [Yang and Lin, 2019]. These can easily be found using the Cypher statement in Figure 4.10. Instead of using Cypher, it would also be possible to use SQL, as shown in Figure 4.11; however, this example shows again that the Cypher query is more elegant than SQL, and easier to comprehend.

Processing Cypher queries depends on the schema information that must be loaded, before queries can be issued. Graph schema info is external to DuckDB. For prototyping you handle it in a support class, that is written in Python. The schema data used in this chapter are available via GitHub. Using the terms found with Cypher, we can construct queries to pass to the searcher and create a BM25 ranking. The code that generates the rankings for all topics is presented in Figure 4.12. In only a few lines of Python code, it is easy to create rankings. From this point, writing the content of `hits` to a runfile and evaluating their effectiveness using `trec_eval` is trivial.

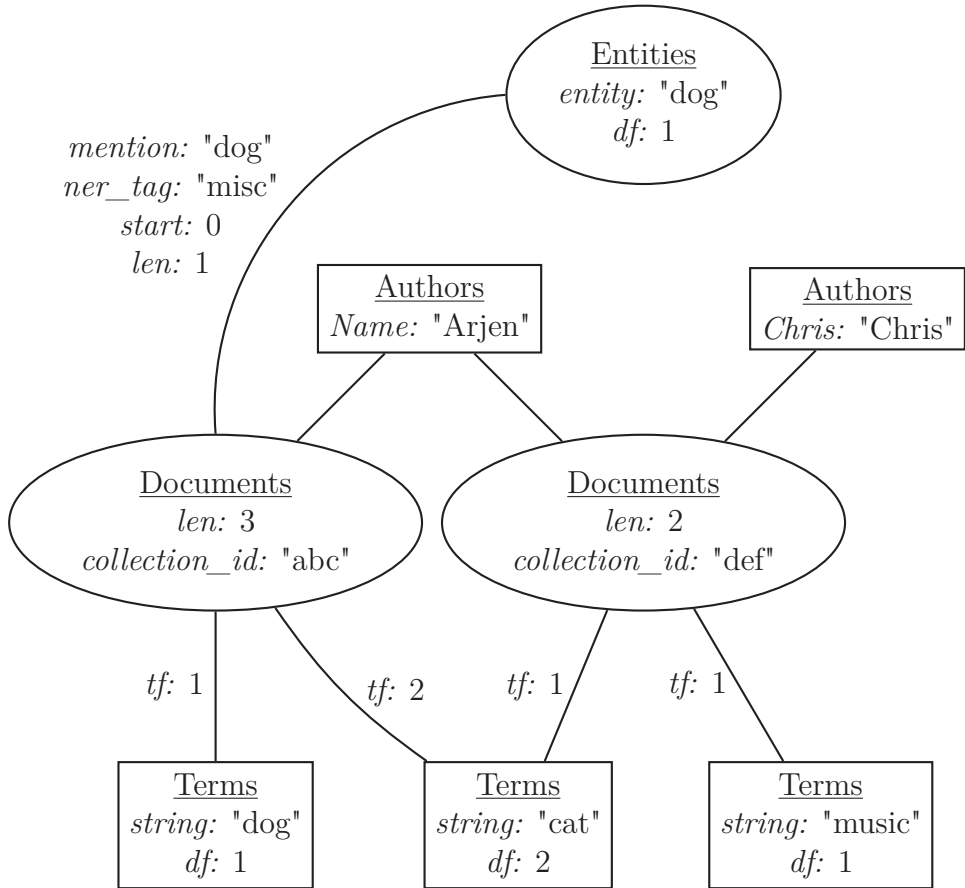


Figure 4.7: Example property graph for the TREC News Track's background linking task. The node types are authors, entities, terms, and documents. Edges connect document nodes to other types of nodes. Both edges and nodes can have properties (following the property graph model). Multiple edges may exist between one entity node and one document node, as one entity can be linked multiple times to one document. Note that some nodes here are squares instead of circles, this was done for readability.

```

1 from geesedb.index import FullTextFromCSV
2
3 index = FullTextFromCSV(
4     database='/path/to/database',
5     docs_file='/path/to/docs.csv',
6     term_dict_file='/path/to/term_dict.csv',
7     term_doc_file='/path/to/term_doc.csv'
8 )
9 index.load_data()

```

Figure 4.8: Load text data from the WashingtonPost collection formatted as CSV files in the format as described by Mühleisen et al. [2014]

```

1 from geesedb.search import Searcher
2
3 searcher = Searcher(
4     database='/path/to/database',
5     n=10
6 )
7 hits = searcher.search_topic('fast dogs')

```

Figure 4.9: Example on how to create a BM25 ranking for the query “fast dogs” that returns the top 10 documents. The Searcher class also takes the optional parameters k_1 and b , in this example it uses the default values 0.9 and 0.4 respectively [Trotman et al., 2012].

```

1 MATCH (d:docs {collection_id: ?})-[]-(t:term_dict)
2 RETURN string
3 ORDER BY tf*log(671945/df)
4 DESC
5 LIMIT 5

```

Figure 4.10: Prepared⁷ Cypher statement that finds the top-5 TF-IDF terms in a given document. The number 671945 on line 3 corresponds to the collection size, this could be computed in a sub-query.

```
1 SELECT term_dict.string
2 FROM term_dict
3 JOIN term_doc ON (term_dict.term_id = term_doc.term_id)
4 JOIN docs ON (docs.doc_id = term_doc.doc_id)
5 WHERE docs.collection_id = ?
6 ORDER BY term_doc.tf * log(671945/term_dict.df)
7 DESC
8 LIMIT 5;
```

Figure 4.11: Prepared SQL statement that finds the top-5 TF-IDF terms in a document.

Instead of “just” ranking with BM25, expressing strategies that would use metadata to adapt the ranking is straightforward. In the case of background linking, it makes sense to consider authorship information when recommending articles that might be suitable as background reading. As journalists often specialize in specific news topics (e.g., politics, foreign affairs, tech), the stories they write often share context. Also, when journalists collaborate on stories, they write together on topics they specialize in. As authorship information is represented in the data graph, we can use the information whether an article is written by the authors of the topic article or by someone they have collaborated with in the past. The graph query that finds the articles that this group of people has written is shown in Figure 4.13.

Depending on the number of documents this query identifies, different rescoring strategies can be decided upon. If the set of documents written by the authors or their co-authors is large, it is possible only to consider these documents, but if the set is small, a score boost might be more appropriate. Figure 4.14 shows an example of how only to consider documents found with the query in Figure 4.13. In this case, we ensure that at least 2000 documents are found before filtering.

For another example, expressiveness of the graph query language is also valuable when considering the occurrences of entities in news articles. Journalists write news articles that relate to events concerning, e.g., people, organizations, or countries. In other words, entities are often the subject of news. So, instead of using the most informative terms in a news article, it is worthwhile to consider the entities identified in the article instead. Important entities tend to be mentioned at the beginning of a news article for this collection [Kamphuis et al., 2019]; Figure 4.15 shows the Cypher

```

1  from geesedb.search import Searcher
2  from geesedb.connection import get_connection
3  from geesedb.resources import
   ↪ get_topics_backgroundlinking
4  from geesedb.interpreter import Translator
5
6  db_path = '/path/to/database'
7  searcher = Searcher(
8      database=db_path,
9      n=1000
10 )
11
12 translator = Translator(db_path)
13 c_query = """cypher TFIDF query"""
14
15 query = translator.translate(c_query)
16 cursor = get_connection(db_path).cursor
17 topics = get_topics_backgroundlinking(
18     '/path/to/topics'
19 )
20 for topic_no, collection_id in topics:
21     cursor.execute(query, [collection_id])
22     topic = ' '.join(cursor.fetchall()[0])
23     hits = searcher.search_topic(topic)

```

Figure 4.12: Create a BM25 ranking for all background linking topics using the top-5 TFIDF terms. Note that a processed topic file was used where only the topic identifier and article id are available. The topic file in this format is provided on our GitHub.

```

1  MATCH (d:docs)-[]-(a:authors)-[]-(d2:docs)-[]-(a2:authors)-
   ↪ []-(d2:docs {collection_id:
   ↪ ?})
2  RETURN DISTINCT d.collection_id

```

Figure 4.13: Cypher query to find documents written by co-authors of the authors of the topic article.

```

1  # import and first lines the same as previous example
2
3  author_c_query = """cypher authorship query"""
4  author_query = t.translate(author_c_query)
5
6  cursor = get_connection(db_path).cursor
7  topics = get_topics_backgroundlinking(
8      '/path/to/topics'
9  )
10 for topic_no, collection_id in topics:
11     cursor.execute(query, [collection_id])
12     topic = ' '.join(cursor.fetchall()[0])
13     hits = searcher.search_topic(topic)
14
15     cursor.execute(author_query, [collection_id])
16     docs_authors = {
17         e[0] for e in cursor.fetchall()
18     }
19     if len(docs_authors) > 2000:
20         hits =
            ↪ hits[hits.collection_id.isin(docs_authors)]

```

Figure 4.14: Find documents written by all authors that collaborated with the authors of the topic article. When there are more than 2000 documents found, we only consider these documents as background reading candidates.

```

1 MATCH (d:docs {collection_id: ?})-[]-(e:entities)
2 RETURN mention
3 ORDER BY start
4 LIMIT 5

```

Figure 4.15: Retrieve the first five entities mentioned in the topic article, and return the terms used to mention the entity.

query to retrieve the text mentions of the first five mentioned entities.

Before it is possible to search using the text describing the first five entity mentions, the text needs to be processed by an entity linker. The term data loaded in GeeseDB was already processed, as it was data loaded from CSV files built from a CIFF file created from an Anserini [Yang et al., 2018a] (Lucene) index. Pyserini [Lin et al., 2021] can be used to tokenize the text in the same way the documents were tokenized. Figure 4.16 shows the Python code where we extract the mentions, process them such that they become a usable query for GeeseDB, and then a BM25 ranking is created with this query.

4.4 Experiments

The previous sections give examples on how non textual content can be used for retrieval experiments. For example, authorship information can be used as a boost for finding related background reading articles. In the example the result set is filtered if sufficient articles are written by the same author, it might however be more interesting to use authorship information as a feature. Let us define a boosting value p for a document when it is written by the same author. Then we can score documents with, e.g., BM25 and adapt the score with p :

$$\text{score} = \begin{cases} p \cdot \text{BM25} & \text{if same author} \\ \text{BM25} & \text{otherwise} \end{cases} \quad (4.1)$$

We choose to multiply the scores by p instead of adding a constant values, as the distribution of BM25 scores can differ quite a bit depending on the topic.⁸ We run a series of retrieval experiments for different values of p and measure the effect on retrieval effectiveness. We run this experiment

⁸Topic terms with a high IDF will increase the BM25 scores for that topic.

```

1  from geesedb.search import Searcher
2  from geesedb.connection import get_connection
3  from geesedb.resources import
   ↪ get_topics_backgroundlinking
4  from geesedb.interpreter import Translator
5  from pyserini.analysis import Analyzer,
   ↪ get_lucene_analyzer
6
7  db_path = '/path/to/database'
8  searcher = Searcher(
9      database=db_path,
10     n=1000
11 )
12
13 analyzer = Analyzer(get_lucene_analyzer())
14
15 translator = Translator(db_path)
16 c_query = """cypher entity query"""
17 query = translator.translate(c_query)
18
19 cursor = get_connection(db_path).cursor
20 topics = get_topics_backgroundlinking(
21     '/path/to/topics'
22 )
23
24 for topic_no, collection_id in topics:
25     cursor.execute(query, [collection_id])
26     topic = ' '.join([e[0] for e in cursor.fetchall()])
27     topic = ' '.join(analyzer.analyze(topic))
28     hits = searcher.search_topic(topic)

```

Figure 4.16: Create a BM25 ranking for all background linking topics using the mention text of the first five linked entities in the source article. The Cypher query is shown in Figure 4.15

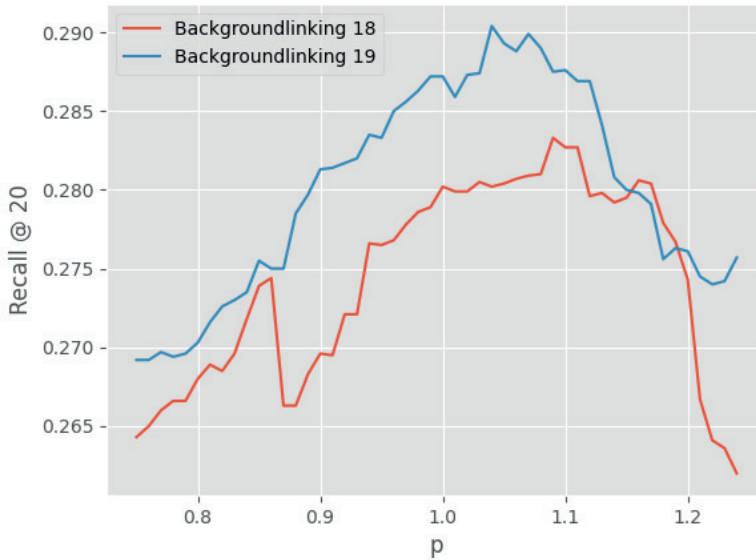


Figure 4.17: Recall @ 20 for varying levels of p

on the background linking tasks for TREC news 2018 and 2019. Both topic sets contain 60 topics. Varying for different values of p , we see the effects on retrieval effectiveness. We evaluate using recall, as BM25 is quite a cheap ranking method whose ranking should be the input for a more sophisticated second stage ranking system to reorder the top documents by their content.

Figures 4.17 to 4.19 show the effectiveness scores for values of p from 0.75 to 1.25 for recall at 20, 50 and 100.⁹ When $p = 1$, documents written by the same author get the same score. Values for p lower than 1 would decrease the scores for articles written by the same author. We see that for both collections, a slight boost of documents that are written by the same author, represented by the values $1 < p < 1.1$, does increase recall at different depths for both document collections. However when p becomes larger than 1.1 the recall tends to decrease, especially for recall at depth 20. This experiment shows that it is quite easy to use GeeseDB for retrieval experiments and adding metadata; showing that including metadata might benefit first stage ranking methods.

⁹Expensive large language model based rerankers do not rerank many documents.

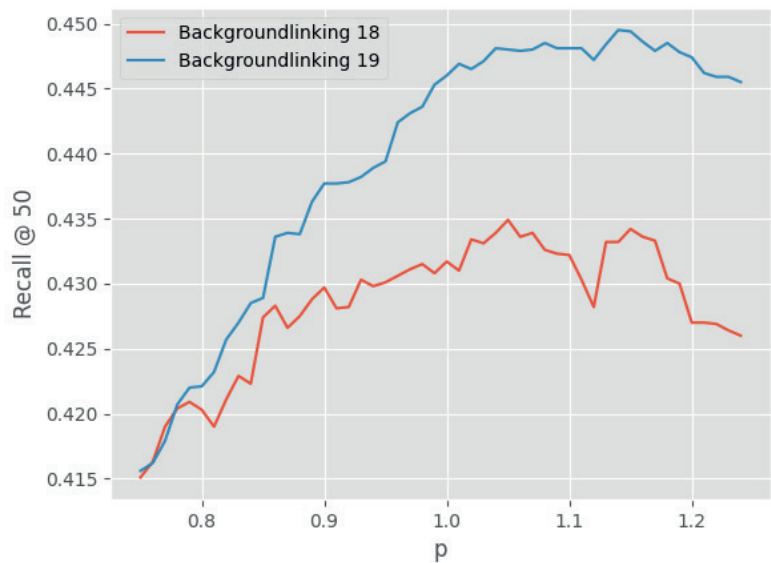


Figure 4.18: Recall @ 50 for varying levels of p

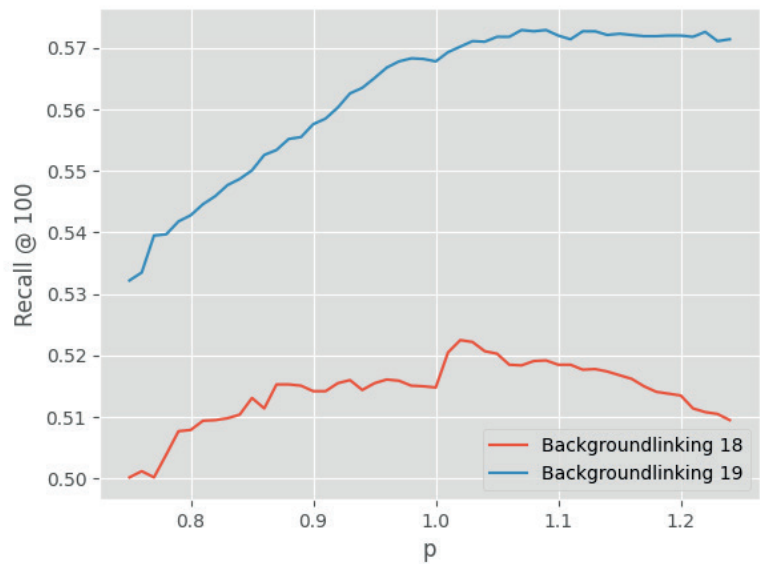


Figure 4.19: Recall @ 100 for varying levels of p

4.5 Conclusion

This chapter described the prototype implementation of GeeseDB, and how we envision graph databases be used for information retrieval research. The GeeseDB system can be considered the answer on our research question: *Can we extend the benefits from using relational databases for information retrieval to using graph databases, while being able to express graph-related problems easier?* As GeeseDB is built on top of a relational engine it automatically inherits the benefits of using relational databases for IR. As GeeseDB can process graph queries, we also are able to express graph related problems.

GeeseDB is still however a prototype system, and more functionalities need to be implemented. In particular, although the architecture is not coupled, it is not (yet) as efficient as traditional methods that do use coupled architectures.

Not all graph operators have been supported in the current implementation. In order to make GeeseDB a usable system for IR researchers, more operators need to be implemented and the system needs to be developed to be more robust.

Chapter 5

Creation of the Entity Graph

“Is this new question a worthy one to investigate?” This latter question we investigate without further ado, thereby cutting short an infinite regress.

Alan Turing - 1950

Abstract

REBL is an extension of the Radboud Entity Linker (REL) for Batch Entity Linking. REBL was developed after encountering unforeseen issues when trying to link the large MS MARCO v2 document collection with REL. In this chapter we discuss the issues we ran into and our solutions to mitigate them. REBL makes it easier to isolate the GPU heavy operations from the CPU heavy operations, by separating the mention detection stage from the candidate selection and entity disambiguation stages. By improving the entity disambiguation module we were able to lower the time needed for linking documents by an order of magnitude.

5.1 Introduction

In the previous chapter we have introduced GeeseDB, a system that can query IR graphs and create rankings. We have used the system to express queries over the Washington Post news article collection. The examples we showed made use of entity annotations that were created by the Radboud Entity Linker (REL) [van Hulst et al., 2020]. Although, the Washington Post collection has been used for retrieval experiments, it is mostly used as a benchmark for news retrieval. In recent years, the MS MARCO ranking collections have become the de-facto benchmark for information retrieval experiments that employ deep learning. These collections are considerably larger than the Washington Post collection; the Washington Post collection consists of 728,626 news articles and blog posts, while the MS MARCO v2 document collections consists of almost 12 million web documents. The MS MARCO v2 document collections has almost 20 times as many documents, which are also longer on average. When using REL, the Washington Post collection can be linked in a relatively short time. Using one GPU¹ the linking time was amounted to approximately one day (24 hours). As the documents in the MS MARCO v2 document collection are longer than those in the Washington Post collection, we expected it would take about a month of time to link this collection using the same system setup. As it is possible to divide the workload over multiple machines, the time needed to link all documents can be decreased easily by deploying more servers.

When starting to link the first segment of this collection however, the time it took to link was considerably larger than expected. In fact, after 72 hours there was a time-out after only linking 20 percent of the segment (200k documents), a processing job that should only take about 12 hours (following a rough estimation). Why entity linking took so much more time than estimated by extrapolation was unclear.

The goal of the research described in this chapter concerns our approach to understand and mitigate these efficiency issues, such that we could link the MS MARCO v2 document collection in an acceptable runtime. In databases runtime efficiency is often improved through set-at-a-time operations instead of tuple-at-a-time operations. This means that instead of doing all computations for one data entry at a time, one computation is calculated for all data entries, before starting the following computation.

We investigate whether the set-at-a-time execution model from database research might be applicable in the context of entity linking as well. Specifically, we developed a batch extension for REL dubbed REBL. REBL splits

¹NVIDIA GeForce RTX 3090

different stages in entity linking such that partial computations for the whole dataset can be calculated before starting the next computation. REBL improves REL efficiency by almost an order of magnitude, decreasing the processing time per document (excluding mention detection) on a sample of 5000 MS MARCO documents from 1.23 seconds to 0.13 seconds. Given modest computational resources, we demonstrate that REBL enables the annotation of a large corpus like MS MARCO v2. We discuss potential improvements that can be made to further improve batch entity linking efficiency. The REBL code and toolkit are available publicly at <https://github.com/informagi/REBL>.

Our third research question;

- *Research Question 3: When does information retrieval research benefit from graph data?*

is not directly answered by the research described in this chapter. The research in this chapter was needed in order to obtain a dataset that we want to approach the third question. So, it forms the basis for the research described in the next chapter that will try to answers this question. In order to give proper context for the research described in this chapter, first entity linking will be described in more depth, before explaining the Radboud Entity Linking system that was used for this research.

5.2 Related Work

Entity linking concerns identifying entity mentions in text and linking them to their corresponding entities in a knowledge graph. It fulfills a key role in the knowledge-grounded understanding of documents. It has been shown effective for diverse tasks in information retrieval [Gerritse et al., 2020, 2022, Xiong et al., 2017, Hasibi et al., 2016, Balog et al., 2013, Reinanda et al., 2015, Chatterjee and Dietz, 2022], natural language processing [Lin et al., 2012, Ferrucci, 2012], and recommendation [Yang et al., 2018b]. Utilizing entity annotations in these downstream tasks depends upon the annotation of text corpora with a method for entity linking. Due to the complexity of entity linking systems, this process is often performed by reliance on a third-party entity linking toolkit, where examples include DBpedia Spotlight [Mendes et al., 2011], TAGME [Ferragina and Scaiella, 2010], Nordlys [Hasibi et al., 2017b], GENRE [De Cao et al., 2021], Blink [Wu et al., 2020], and REL [van Hulst et al., 2020].

A caveat in existing entity linking toolkits is that they are not designed for batch processing large numbers of documents. Existing entity linking

toolkits are primarily optimized to annotate individual documents, one at a time. This severely restricts utilizing state-of-the-art entity linking tools, such as REL, Blink, and GENRE, that employ neural approaches and require GPUs for fast operation. Annotating millions of documents incurs significant computational overhead, to the extent that annotation of a large text corpus becomes practically infeasible using modest computational power resources, especially when tagging documents one by one. Batch entity linking is, however, necessary to build today’s data-hungry machine learning models, considering large text corpora like the MS MARCO v2 (12M Web documents) [Bajaj et al., 2016], or the newer ClueWeb22 collection [Overwijk et al., 2022].

5.3 REL

This research specifically concerns the Radboud Entity Linking (REL) toolkit, in the context of processing large corpora. REL is a state-of-the-art entity linking system as shown empirically and independently by Bast et al. [2023]. Its system architecture is similar to that of competing systems, including BLINK [Wu et al., 2020]. REL annotates individual documents efficiently, requiring only modest computational resources while performing competitively compared to the state-of-the-art methods on effectiveness. It considers entity linking as a modular problem consisting of three stages:

- *Mention Detection.* This step aims to identify all possible text spans in a document that might refer to an entity. If text spans that refer to entities are not appropriately identified, the system cannot correctly link the entity in later stages. REL is built to be a modular system, making it possible to use different named entity recognition (NER) systems for this stage. For REL the default system used is FLAIR [Akbik et al., 2019], which is a state-of-the-art NER system that uses contextual word embeddings.
- *Candidate Selection.* For every detected mention, REL considers up to $k_1 + k_2 (= 7)$ candidate entities. $k_1 (= 4)$ candidate entities are selected based on their a priori occurrence probability $p(e|m)$ (for entity e given mention m). These priors are pre-calculated by summing Wikipedia hyperlinks and priors from the CrossWiki corpus [Spitkovsky and Chang, 2012].² The other $k_2 (= 3)$ entities are chosen based on the

²So the conditional is based on occurrences in different collections, these are then used as a prior in the candidate selection stage.

similarity of their entity embedding to the word embeddings of the context. In this stage, a context of a maximum of 200-word tokens is considered.

- *Entity Disambiguation.* The final step aims to map the mention to the correct entity in a knowledge base. The candidate entities for each mention are obtained from the previous stage. REL implements the Ment-norm method proposed by Le and Titov [2018].

After all entities have been found, REL produces a JSON object that contains the entities and their respective locations in the source text using standoff annotation. Because of REL’s modular architecture, it is an ideal system to deploy set-at-a-time.

5.4 From REL to REBL

The MS MARCO v2 document collection contains 11,959,635 documents split into 60 compressed files, totaling roughly 33GB. Uncompressed, these files are in JSON line format (where every line represents a document, described by a JSON object). These JSON objects have five fields: *url*, *title*, *headings*, *body*, and *docid*. We wanted to link the documents’ titles, headings, and bodies for our experiments. We link to the 2019-07 Wikipedia dump, one of the two dumps also used in the initial development of REL.

In order to ease linking this size of data, we separated the GPU-heavy mention detection stage from the CPU-heavy candidate selection and entity disambiguation stages; the modified code can be found on GitHub.³ For REBL, the input for mention detection are the compressed MS MARCO v2 document files, and its output consists of the mentions found and their location in the document, in Apache Parquet format.⁴ These files and the source text are the input for the subsequent phases (candidate selection and entity disambiguation). The final output consists of Parquet files containing spans of text and their linked entities. In the following section, we discuss what is changed for mention detection, candidate selection, and entity disambiguation steps to make REL more suited to annotate large collections with entity links, in a batch processing manner.

³<https://github.com/informagi/REBL>, last accessed September 2025

⁴<https://github.com/apache/parquet-format>, last accessed September 2025

5.4.1 Mention Detection

REL [van Hulst et al., 2020] uses Flair [Akbik et al., 2019] as the default method for mention detection, a state-of-the-art named entity recognition system. In this section, we focus on inefficiencies that arise when interfacing between REL and flair. Flair uses the `segtok`⁵ package to segment an (Indo-European) document in sentences, internally represented as `Sentence` objects. These sentences are split into words/symbols represented as `Token` objects. When creating these representations, however, it is not possible to recreate the source text properly, as Flair adjusts the representations in its internal processing of text data. In order to have full control of these representations we decided to construct the underlying data structures for REBL ourselves. To do this, we used the `syntok`⁶ package, a follow-up version of `segtok`. Both packages were developed by the same author, who claims that the `syntok` package segments sentences better than `segtok`. Using this new representation we fixed two problems:

1. Flair removes white space characters when multiple occur after each other, REL accounts for this by counting the number of white spaces characters removed, a somewhat inefficient process prone to introduce inaccuracies in the mapping between source text and the annotated result. Using our own representation we know the start of every token, and we do not need an additional system that tracks the number of removed white spaces. So we change the interface between Flair and REL, as we create the Flair objects that REL uses manually.
2. Various zero width Unicode characters are removed by Flair from the source text before creating a token: zero width space (U+200B), zero width non-joiner (U+200C), variation selector-16 (U+FE0F), and zero width no-break space (U+FEFF). These characters occur rarely, but in a collection as large and diverse as MS MARCO v2, these characters do occur. When encountering these characters using REBL, the token objects were constructed such that the span and offset of the token still referred to that of the source text: For the case of the zero width space, we updated the `syntok` package. While, according to the Unicode standard, zero width space is not a whitespace character, it should be considered a character that separates two words. For the other Unicode characters removed by Flair, we manually update the span in the `Token` objects created by Flair such that they refer

⁵<https://github.com/fnl/segtok>, last accessed September 2025

⁶<https://github.com/fnl/syntok>, last accessed September 2025

correctly to the positions in the source text. Now, when Flair identifies a series of tokens as a possible mention, we can directly identify the location in the source text from the `Token` objects.

To efficiently use GPU resources, it is important to create batches of data to decrease the number of I/O operations. Every time a GPU calculation is called, data needs to be transferred from and to the GPU hardware. This way it is possible to parallelise the mention detection stage. Flair does in fact support named entity recognition in batches; multiple pieces of text can be sent to the GPU, to achieve an overall faster inference time (as fewer I/O operations are needed). Because REL was designed to tag one document at a time, it did not utilize this functionality. REBL exploits this feature, allowing users to specify the number of documents to be tagged simultaneously, increasing linking efficiency.

5.4.2 Candidate Selection and Entity Disambiguation

For candidate selection, REL makes use of a $p(e|m)$ prior, where e is an entity, and m is a mention. These priors are saved in an (SQLite) database, and up to 100 priors per mention are considered. However, data conversion between the client and the representation stored in the database incurred a high serialization cost. We updated this to a format that is faster to load, with the additional benefit of a considerably decreased database size.⁷ We experimented with data storage in the DuckDB column-oriented database as an alternative. However, we found that SQLite was (still) more efficient as a key-value store, at least in DuckDB’s state of development when we ran the experiments.

The entity disambiguation stage took much longer than reported in the original REL paper. This difference can partially be explained by the length of the documents to be linked. The documents evaluated by van Hulst et al. [2020] consisted of, on average, 323 tokens, with an average of 42 mentions to consider. The average number of tokens in an MS MARCO v2 document is about 1,800, with 84 possible mentions per document.⁸ Per mention, 100 tokens to the left and the right (so 200 total) are considered as context for the disambiguation model. The longer documents result in a higher memory consumption per context and document, with higher processing costs as the result.

⁷The table that represents the priors shrank from 9.6GB to 2.2GB.

⁸These figures are calculated over the body field; we also tagged the shorter title and headers fields.

We improve the efficiency of the entity disambiguation step such that it can be run in a manageable time. REL recreated a database cursor every time candidates were being generated. We rewrote the REL database code to create a single database cursor for the entity disambiguation module. When a mention occurs multiple times within a document, the exact same queries were issued to the database multiple times within a document. By caching the output of these queries, we could considerably lower the number of database calls needed. We ended up caching all database calls for a segment, as we ran the process for every segment separately.

The default setting of REL is to keep embeddings on the GPU after they are loaded, also ones that are not needed for disambiguation for following batches. This design decision, however, slowed down disambiguation when many documents were being processed consecutively, because operations like normalization were carried out over all embeddings on the GPU. A considerable speed-up has been achieved by clearing these embeddings as soon as a document is processed.

Finally, after retrieving the embeddings from the database, REL puts them in a Python list. We rewrote the REL code such that the binary data is directly loaded from NumPy, a data format that Pytorch can use directly.

5.5 Effects on Execution

In the mention detection stage, we improved tokenization and applied batching. The MS MARCO v1 collection does not contain characters that cause the problem in tokenization; the documents in that version of the collection were sanitized before it was published. In the MS MARCO v2 collection, 411,906 documents have tokens that would be removed automatically by Flair, which corresponds to 3.4% of all documents. Batching documents in the mention detection stage decreased the average time for finding all named entities. We used batches of size 10, as the documents are relatively large. The optimal batch size will depend on the available GPU memory, and the length of the documents that need to be processed.

Some documents in the MS MARCO v2 collection cannot be linked. This happens only in extraordinary cases where linking with entities did not make sense in the first place, an example being a document consisting of numbers only. Here, the `syntok` package created one long `Sentence` object from this file that could not fit in GPU memory.

Table 5.1 shows our improvements to the candidate selection and entity disambiguation step and describes how much time is saved in REBL. The code improvements to create the database cursor only once and to load

the data directly from NumPy had no noticeable effect on the overall run time of entity disambiguation and are not reported in this table. This is likely because the costs of these operations are quite low, relatively to the efficiency improvements made by the other changes. Note that the large standard deviations are primarily due to the differences in processing costs between long and short documents.

5.6 Conclusion and Discussion

This chapter described REBL, an extension for the Radboud Entity Linker. We utilized REL’s modular design to separate the GPU-heavy mention detection stage from the CPU-heavy candidate selection and entity disambiguation stages. The mention detection module is now more robust and reliable, using a better segmenter and preserving location metadata correctly. The candidate selection and entity disambiguation steps were updated to improve their runtime. Although it is now possible to run REL. [van Hulst et al., 2020] on MS MARCO v2 [Bajaj et al., 2016] in a (for us) somewhat reasonable time, we identified further improvements to implement.

In the candidate selection step, found mentions are compared to all other mentions. The complexity of this step is $O(n^2)$, with n being the number of mentions found in a document, which is especially problematic for longer documents. As we are only interested in similar mentions, it may be worthwhile to implement a locality-sensitive hashing algorithm to decrease the number of comparisons needed at this stage.

Per mention, all context tokens are considered. This context has to be constructed from the source document. As a result, we load the source data a second time during candidate selection. Alternatively, we could output the mention context in the mention detection stage, which could speed up the candidate selection stage as we do not have to reconstruct the context for a second time. However, this would significantly increase the size of the mention detection output. More experiments are needed to strike the right balance here. Having a streaming approach would probably mitigate this issue. Another way a streaming approach might benefit REBL is that now, because a two-step approach is being used, intermediate results are written to the file system in parquet format.

Overall, it has become clear that a data processing-oriented perspective on entity linking is necessary for efficient solutions. A set-at-a-time approach should be preferred over a tuple-at-a-time approach when large amounts of data have to be processed.

Table 5.1: Efficiency improvements for Candidate Selection and Entity Disambiguation. Improvements are calculated over a sample of 5000 documents using a machine with an Intel Xeon Silver 4214 CPU @ 2.20GHz using two cores with 187GB RAM and a GeForce RTX 2080 Ti (11GB) GPU. Improvements are cumulative; the times shown include the previous improvement as well.

Improvement	Seconds	Explanation
Old Candidate Selection + Entity Disambiguation	1.23 ± 2.09	Average time it takes to select candidates and disambiguate per document
No embedding reset	0.26 ± 1.60	The default setting of REL was to keep embeddings in GPU memory after they were loaded by clearing them from GPU memory after every document a speed up was achieved.
Cache database calls	0.15 ± 1.31	When an entity occurs within a document, there is a high probability of it occurring multiple times. By caching the calls, we increase memory usage but can lower the time needed for candidate selection and entity disambiguation.
Representation change candidates	0.13 ± 1.19	By representing the candidates better in the database, we were able to save on conversion time, lowering the time needed for candidate selection.

Chapter 6

Using the Entity Graph

The reader will have anticipated that I have no very convincing arguments of a positive nature to support my views. If I had I should not have taken such pains to point out the fallacies in contrary views.

Alan Turing - 1950

Abstract

MMEAD, or MS MARCO Entity Annotations and Disambiguations, is a resource of entity links for the MS MARCO datasets. We specify a format to store and share links for both document and passage collections of MS MARCO. Following this specification, we release entity links to Wikipedia for documents and passages in both MS MARCO collections (v1 and v2). Entity links have been produced by the REL and BLINK systems. MMEAD is an easy-to-install Python package, allowing users to load the link data and entity embeddings effortlessly. Using MMEAD takes only a few lines of code. Finally, we show how MMEAD can be used for IR research that uses entity information. We show how to improve recall@1000 and MRR@10 on more complex queries on the MS MARCO v1 passage dataset by

using this resource. We also demonstrate how entity expansions can be used for interactive search applications.

6.1 Introduction

The MS MARCO datasets [Bajaj et al., 2016] have become the *de facto* benchmark for evaluating deep learning methods for Information Retrieval (IR). The TREC deep learning track [Craswell et al., 2021], which has run since 2019, derives its datasets from the MS MARCO passage and document collections. The collections have been used in zero- and few-shot scenarios for diverse retrieval tasks and domains [Thakur et al., 2021, 2023, Xu et al., 2022]. They also serve as primary resources for training deep learning models for downstream IR tasks such as conversational search [Dalton et al., 2021] and search with knowledge graphs [Gerritse et al., 2022] to achieve state-of-the-art results.

Purely text-based neural IR models, trained using MS MARCO collections, are generally unable to reason over complex concepts in the social and physical world [Bosselut et al., 2021, Sciavolino et al., 2021]. In response, recently proposed neuro-symbolic methods aim to combine neural models and symbolic AI approaches, e.g., by using knowledge graphs, which map concepts to symbols and relations. An essential step in developing *neuro-symbolic* models is connecting text to entities that represent the world’s concepts formally. This step is mainly done using *Entity linking*, an intermediary step between text and knowledge graphs, which detects entity mentions in the text and links them to the corresponding entries in a knowledge graph.

Despite the proven effectiveness of neuro-symbolic AI – and for IR models in particular [Tran and Yates, 2022, Gerritse et al., 2022, Chatterjee and Dietz, 2022] – the IR community has made limited efforts to develop such models. We argue that the annotation of large-scale collections with entities is a primary hindrance; entity linking methods are computationally expensive. Running them over a large text corpus (e.g., MS MARCO v2 with 12M documents and 140M passages) requires extensive resources. This work aims to fill this gap by making entity annotations of the MS MARCO ranking collections readily available and easy to use. In the context of green AI / IR, computing the data only once, and then sharing it with many will reduce the amount of computation needed for people to carry out research on entity annotated web documents.

This chapter continues where the previous chapter ended. In the previous chapter the batch extension for the Radboud Entity Linking system, REBL,

is introduced. In this chapter MMEAD is presented, a resource that provides entity links for the MS MARCO document and passage ranking collections. Two state-of-the-art entity linking tools, namely REL [van Hulst et al., 2020], with efficiency improvements made in REBL, [Kamphuis et al., 2022] and BLINK [Wu et al., 2020], are utilized for annotating the corpora. The annotations are stored in a DuckDB database, enabling efficient analytical operations and fast access to the entities. The resource is available as a Python package and can be installed from PyPI effortlessly. The resource also includes a sample demo, enabling queries with complex compositional structures about entities.

The primary objective of the work in this chapter is to simplify the study of neuro-symbolic IR research, enabling further steps in neural IR. In our experiments, we show significant improvements on recall for neural re-ranking IR models when using MMEAD annotations as bag-of-word expansions for queries and passages. Our experiments reveal that the difference in effectiveness is even greater (in terms of both recall and MRR) for complex queries that require further reasoning over entities.

To show the usefulness of our resource beyond plain ranking, we also present how to enrich interactive search applications. Specifically, we demonstrate how to obtain entities’ geographical locations by relating the entities found in passages to their Wikidata entries. Plotting these entities on the world map shows that the MS MARCO passages can be geo-located all over the world. We can also move from location to web text by retrieving all passages associated with a geographical location.

To return to our last research question;

- *Research Question 3: When does information retrieval research benefit from graph data?*

This chapter attempts to demonstrate that graph data, in this case a graph of entities, can help retrieval by considering metadata that is included through graph operations. Although we have not used GeeseDB, the work described in Chapter 4, it would be a natural application of GeeseDB to carry out the query expansion methods described.

In summary, this chapter makes the following contributions:

- We annotate the documents of the MS MARCO passage and document collections and share these annotations. By sharing these annotations, we ease future research in neuro-symbolic retrieval, which extensively uses entity information. We also provide useful metadata such as Wikipedia2Vec [Yamada et al., 2016] entity embeddings.

- We provide a Python library that makes our data easy to use. All data is stored in DuckDB tables, which can be loaded and queried quickly. The library is easy to install through PyPI, and the entity annotations are available with only a few lines of code.
- We experimentally show that retrieval effectiveness measured by recall significantly increases when using MMEAD. The improvement is even greater for complex queries, where we observe low retrieval effectiveness using text-only IR models.
- We demonstrate how the data can be used in geographical applications. For example, we can plot on a static map all entities found in the MS MARCO v2 passage collection for which geographical data is available. Additionally, one can retrieve all passages associated with a geographical location.

MMEAD is publicly available at <https://github.com/informagi/mmead>. In this chapter we will first describe the systems that we used to create the MMEAD dataset. Then we describe how the dataset was designed, and how it can be used. With some experiments we show that using MMEAD can lead to higher retrieval effectiveness on the MS MARCO v1 dataset. We also demonstrate that this dataset can be used for geographical related tasks.

6.2 Background

In this section, we describe systems that are used for creating entity annotations on the MS MARCO collections for MMEAD. Some information presented in this section overlaps with that of the previous chapter, but the full context is included here to enable reading this chapter as standalone work.

6.2.1 REL

REL (Radboud Entity Linker) [van Hulst et al., 2020] is a state-of-the-art open-source entity linking tool designed for high throughput and precision. REL links entities to a knowledge graph (Wikipedia) using a three-stage approach: (1) mention detection, (2) candidate selection, and (3) entity disambiguation. We briefly explain these three steps:

1. *Mention Detection.* REL starts the entity linking process by first identifying all text spans that might refer to an entity. In this stage, it

is essential that all possible entities in the text are identified, as only the output of this stage can be considered an entity by REL. These spans are identified using a named entity recognition (NER) model based on contextual word embeddings. For our experiments, we use the NER model based on Flair [Akbik et al., 2019] embeddings.

2. *Candidate Selection.* Up to seven candidate entities are considered for every mention found by Flair. Part of these entities are selected according to the prior probability $P(e|m)$ of the mention m being linked to the entity e . This probability is estimated from existing dictionaries YAGO [Hoffart et al., 2011], Wikipedia, and CrossWikis [Spitkovsky and Chang, 2012]. Precisely, the top-4 ranked entities are selected based on $P(e|m) = \min(1, P_{Wiki}(e|m) + P_{YAGO}(e|m))$, where $P_{YAGO}(e|m)$ is a uniform probability from the YAGO dictionary and $P_{Wiki}(e|m)$ is computed based on the summation of hyper-link counts in Wikipedia and the CrossWikis corpus. The remaining three candidate entities are determined according to the similarity of an entity and the context of a mention (entities provided in the previous step are not considered). For the top-ranked candidates based on $P(e|m)$ probabilities, the context similarity is calculated by $\mathbf{e}^T \sum_{w \in c} \mathbf{w}$. Here $\mathbf{e}$ is the entity embedding for entity e , and $\mathbf{w}$ are the word embeddings in context c , with a maximum length of 200-word tokens. The entity and word embeddings are jointly learned using Wikipedia2Vec [Yamada et al., 2016].
3. *Entity Disambiguation.* The final stage tries to select the correct entity from the candidate entities and maps it to the corresponding entry in a knowledge graph (Wikipedia). For this, REL assumes a latent relation between entities in the text and utilizes the Ment-norm method proposed by Le and Titov [2018].

6.2.2 BLINK

BLINK [Wu et al., 2020] is a BERT-based [Devlin et al., 2019] model for candidate selection and entity disambiguation, which assumes that entity mentions are already given. When utilized in an end-to-end entity linking setup, BLINK achieves similar effectiveness scores as REL. Below we describe the three steps of mention detection, candidate selection, and entity disambiguation for end-to-end entity linking using BLINK.

1. *Mention Detection.* The mention detection stage can be done using an NER model. Like REL, we use Flair [Akbik et al., 2019] for mention

detection.

2. *Candidate Selection.* BLINK considers ten candidates for each mention. The candidates are selected through a bi-encoder (similar to the bi-encoder described by Humeau et al. [2019]) that embeds mention contexts and entity descriptions. The mention and the entity are encoded into two separate vectors. The similarity score is then calculated using the dot-product of the two vectors representing the mention context and the entity.
3. *Entity Disambiguation.* For entity disambiguation, BLINK employs a cross-encoder to re-rank the top 10 candidates selected by the candidate selection stage. The cross-encoder usage is similar to cross encoder described by Humeau et al. [2019], which employs a cross-attention mechanism between the mention context and entity descriptions. The input is the concatenation of the mention text and the candidate entity description.

6.2.3 DuckDB

DuckDB [Raasveldt and Mühleisen, 2019] is an in-process column-oriented database management system. It is designed with requirements that are beneficial for the MMEAD resource:

1. *Efficient analytics.* DuckDB is designed for analytical (OLAP) workloads, while many other database systems are optimized for transactional queries (OLTP). DuckDB is especially suitable for cases where analytics are more important than transactions. As we release a resource, transactions (after loading the data) are unnecessary, making an analytics database more useful than a transactional-focused one.
2. *In-process.* DuckDB runs in-process, which means no database server is necessary, and all data processing happens in-process. This allows the database to be installed from PyPI without any additional steps.
3. *Efficient data transfer.* Because DuckDB runs in-process, it can transfer data from and to the database more easily, as the address space is shared. In particular, DuckDB uses an API built around NumPy and Pandas, which makes data (almost) immediately available for further data analysis within Python.

DuckDB also supports the JSON and parquet file formats, making data loading especially fast when data is provided in such formats.

6.3 MMEAD

MMEAD provides links for MS MARCO collections v1 and v2 created by the REL entity linker, and links for the MS MARCO v1 passage collection by the BLINK entity linker. For REL, we use its batch entity linking extension, REBL [Kamphuis et al., 2022]. The knowledge graphs used for the REL and BLINK entity linkers are Wikipedia dumps from 2019-07 and 2019-08, respectively. Both dumps are publicly available from the linking systems’ Github pages.

6.3.1 Goals

The design criteria for MMEAD are based on the following goals:

- *Easy-to-use.* It should be easy to load and use the linked entities in experiments. With only a few lines of code, it should be possible to load entities and use them for analysis. Additional information should also be readily available, like where entities appear in the text and their latent representations.
- *High-quality entity links.* We wish to release high-quality entity links for the MS MARCO collections, so that applying neuro-symbolic models and reasoning over entities becomes feasible.
- *Extensibility.* It should be easy to link the collections with a different entity linking system and publish them in the same format as MMEAD. This way, we can integrate links produced by other entity linking systems and make them automatically available through the MMEAD framework.
- *Useful metadata.* Additional data that can help with experiments should be provided; this includes mapping entities to their respective identifiers and latent representations.

6.3.2 Design

Easy-to-use. To create an easy-to-use package, we make the MMEAD data publicly available as JSONL files, which is the same format as the MS MARCO v2 collections. Each line of JSON contains entity links for one of the documents or passages in the collections; see Figure 6.1. The corresponding document can be identified through the JSON field that represents the document/passage identifier: `docid` for documents and `pid`

for passages. Then, for every section of a document, a separate JSON field is available to access the entities in that section. For passages, there is only one section containing the entity annotations of the passage, while for MS MARCO v2 documents, we link not only the `body` of the document but also the `header` and the `title`.

All essential information about the entity mentions and linked entities is stored in the JSON objects. Specifically, the following metadata is made available: `entity_id`, `start_pos`, `end_pos`, `entity`, and `details`. The field `entity_id` stores the identifier that refers to the entry in the knowledge graph (Wikipedia, in our case). The `start_pos` and `end_pos` fields store the start and end positions of the text span that refers to the linked entity (i.e., as a standoff annotation of the entity mention). The positions are UTF-8 indices into the text, ready to be used in Python to extract the relevant parts of the document. The field `entity` stores the text representation of the entity from the knowledge graph. We chose to store this field for convenience and human readability. The `details` field is a JSON object that stores linker-specific information; examples include the entity type according to the NER module and the confidence of the identified mention.

High-quality entity links. MMEAD provides entity links produced by state-of-the-art entity linking systems. Both REL and BLINK are high precision entity linkers, ensuring that identified mentions and their corresponding entities are likely correct. The knowledge graphs used by the entity linkers are the same as those used in the original studies; this way, extensive research has been done to confirm the precision of the linking systems.

Extensibility. We ensure extensibility by clearly describing the format in which the entity links are provided. If another system shares its links in the same format, the MMEAD Python library can work with the data directly. The `details` field per entity annotation enables inclusion of linker-specific information. REL provides specific instructions on updating the system to newer versions of Wikipedia in its documentation, making it possible to easily release links to newer versions of Wikipedia.

Useful metadata. Alongside the entity links, we also provide additional useful metadata. Specifically, we release Wikipedia2Vec [Yamada et al., 2016] embeddings (300d and 500d feature vectors). REL uses the 300d Wikipedia2Vec feature vectors internally for candidate selection. These

Table 6.1: Number of entities linked by REL; we show the total number of entities found and how many entities there are per passage/document on average.

	passages	docs
MS MARCO v1	18,561,221 (2.10)	145,725,732 (45.34)
MS MARCO v2	233,254,024 (1.69)	661,183,287 (55.28)

feature vectors consist of word embeddings *and* entity embeddings mapped into the same high-dimensional feature space. These embeddings can be used directly for information retrieval research [Gerritse et al., 2020, 2022]. We also release a mapping of entities to their identifiers. The entity descriptions can change in different versions of Wikipedia, but their identifiers remain constant. The identifier can also be used to find the corresponding entity in other knowledge graphs such as Wikidata.

6.3.3 An Example

A passage from the MS MARCO v1 passage ranking collection is shown below.

A: While Pablo Picasso may be best known as one of the pioneers, along with George Braques, of the Cubism style of art, the Spanish painter and sculptor also painted in many other styles. Before Cubism,...

Several text spans in this text can be considered as entities. REL identifies three of these entities: *Pable Picasso*, *Georges Braque*, and *Spain*. It does however miss the concept *Cubism* as an entity. The output of the system is converted to our JSON specification, which results in the JSON object presented in Figure 6.1.

Table 6.1 shows the number of entities found in the collections by the REL system. Blink found 21,968,356 entity links for the v1 passage collection. For 11,177,904 entities, the two linking systems produced exactly the same output.

6.4 How To Use

MMEAD comes with easy-to-use Python code, allowing users to work with the resource effortlessly. To start, MMEAD can be installed from PyPI

```
{
  "passage": [
    {
      "entity_id": 24176,
      "start_pos": 9,
      "end_pos": 22,
      "entity": "Pablo Picasso"
    },
    {
      "entity_id": 12317,
      "start_pos": 76,
      "end_pos": 90,
      "entity": "Georges Braque"
    },
    {
      "entity_id": 26667,
      "start_pos": 124,
      "end_pos": 131,
      "entity": "Spain"
    }
  ],
  "pid": "3263"
}
```

Figure 6.1: Example of MMEAD annotations for a MS MARCO passage in JSON format.

```
1 from mmead import get_links
2 links = get_links('v1', 'passage', linker='rel')
```

Figure 6.2: Example of how to load MMEAD entity links for the MS MARCO v1 passage collection.

```
1 >>> links.load_links_from_docid(123)
2 {"passage":[{"entity_id":"7954681", ... }]
```

Figure 6.3: Example of how to load the entity links for a document. For formatting reasons, we do not show the full output.

using pip:

```
$ pip install mmead
```

After installation of MMEAD, the entity links can be loaded into a DuckDB database [Raasveldt and Mühleisen, 2019], as shown in Figure 6.2, with only a couple of lines of code.

When running this code for the first time, initialization will take some time, as all the data need to be downloaded and ingested into the DuckDB database. After loading the data for the first time, it is automatically stored on disk. Loading the persisted data for later usage will only take seconds.

Once the data is loaded, it is ready to use. We provide a simple interface to access the data. The code shown in Figure 6.3 loads the entity links available for a document in the MS MARCO v1 passage ranking collection. When using this function, the data is provided in JSON format, making it easy to access the annotations.

We also provide word embeddings and entity embeddings generated by Wikipedia2Vec [Yamada et al., 2016] based on the 2019-07 Wikipedia dump. These embeddings are stored in DuckDB tables and are available as Numpy arrays after loading. Figure 6.4 shows how embeddings are loaded using MMEAD. The example demonstrates that the entity embedding of *Montreal* and the word embedding of “Montreal” are closer to each other than the word embeddings of the two words “Montreal” and “green” based on dot-product as a similarity function. The dimensionality of the embedding vectors (300 or 500) can be specified in the code.

The mapping between the official Wikipedia identifiers and entity text representations is extracted from the 2019-07 Wikipedia dump. If entity

```
1  >>> from mmead import get_embeddings
2  >>> e = get_embeddings(300, verbose=False)
3  >>> montreal_word = e.load_word_embedding("Montreal")
4  >>> montreal_entity = e.load_entity_embedding("Montreal")
5  >>> green_word = e.load_word_embedding("green")
6
7  >>> montreal_word @ montreal_entity
8  31.83191792
9  >>> montreal_word @ green_word
10 5.55568354
11
12 >>> toronto_word = e.load_word_embedding("Toronto")
13 >>> toronto_word
14 array([-1.497e-01, -7.765e-01, -1.000e-02, ...])
15 >>> montreal_word @ toronto_word
16 21.62585146
```

Figure 6.4: Example code for loading word and entity embeddings. It shows that the dot-product between “Montreal” word and entity embeddings is greater than the dot-product of embedding vectors for the word “Montreal” and “green” (a word chosen at random). The word embeddings of Montreal and Toronto, two cities in Canada, are also more similar.

```
1 >>> from mmead import get_mappings
2 >>> m = get_mappings(verbose=False)
3 >>> m.get_id_from_entity('Montreal')
4 7954681
5 >>> m.get_entity_from_id(7954681)
6 'Montreal'
```

Figure 6.5: Entity names and identifiers are accessible in MMEAD. Given an entity text, we can directly find its corresponding identifier and vice versa.

annotations from another version of Wikipedia are available, the MMEAD mappings can be used to match entities between the dumps. Needless to say, emerging entities in newer versions of Wikipedia cannot be mapped to the version that is available in MMEAD (as their pages did not yet exist when making this dataset). However, existing entities in MMEAD can be mapped to newer versions of Wikipedia in a straightforward manner. Figure 6.5 shows how entity identifiers can be matched to their text and the other way around.

As DuckDB is used as a database engine for MMEAD, it is possible to directly access the underlying tables and issue structured queries in an efficient manner. Figure 6.6 shows an example, where a connection to the database is created, and the identifiers of passages containing the entity *Nijmegen* are retrieved.

All data can be downloaded directly as well, and links to the data are provided on our Github page.¹

6.5 Entity Expansion with MMEAD

To demonstrate the usefulness of MMEAD for (neural) retrieval models, we have conducted experiments that extend existing models with MMEAD annotations. These experiments serve a demonstrative purpose only, and the full potential of this resource is to be further explored in (neuro-)symbolic IR models [Gerritse et al., 2022, Tran and Yates, 2022].

¹<https://github.com/informagi/mmead>, last accessed September 2025

```

1  >>> from mmead import load_links
2  >>> cursor = load_links(
3  ...     'msmarco_v1_passage_links',
4  ...     verbose=False
5  ... )
6  >>> cursor.execute("""
7  ...     SELECT pid
8  ...     FROM msmarco_v1_passage_links_rel
9  ...     WHERE entity='Nijmegen'
10 ... """)
11 [(771129,), (1273612,), (1418035,), ... ]

```

Figure 6.6: All data is stored in DuckDB tables, and thus it is possible to directly access the tables and issue queries. In this example, we extract the identifiers of passages that contain the city of Nijmegen.

6.5.1 Methods

BM25 expansion. We experimented with three retrieval methods to show the benefits of entity annotation for passage ranking: one baseline method and two methods that use query entity expansion [Shehata, 2022] using REL:

- a BM25 – No Expansion.** As a baseline method, we used BM25 as implemented in Anserini [Yang et al., 2018a] using hyper-parameters $k_1 = 0.82$ and $b = 0.68$, tuned to be optimal for the MS MARCO dataset. MS MARCO was indexed normally, and no expansion was considered for the queries or the passages.
- b BM25 – Entity Text Expansion.** In this method, passages and queries are concatenated with the text representation of their annotated entities (from REL). Once the passages and queries have been expanded with entities, we run BM25 with the same hyper-parameter settings as described in **a**.
- c BM25 – Entity Hash Expansion.** Instead of using the text representation of entities as an expansion, we concatenate the passages and queries by the MD5 hash of the entity text (from REL). The use of MD5 hashing provides a consistent representation of multi-word terms and avoids partial or incorrect matching between queries and

- a. did sacajawea cross the pacific ocean with lewis and clark
- b. The same text as shown in a. + *Sacagawea Clark Pacific Ocean C. S. Lewis*
- c. The same text as shown in a. + 860324 a97fed 3e3b0e3fe907

Figure 6.7: Examples of queries for the three different experiments; (a) the non-expanded query, (b) the query with entity text expansion, and (c) the query with entity hash expansion. Text expansions are shown in italics. The MD5 hashes shown in (c) are shortened for formatting.

non-relevant passages; e.g., passages that contain the word “united”, do not benefit if the query contains “United States” as an entity. Again, after expansion, we run BM25 with the same hyper-parameter settings described in a.

In these experiments, the identified entities are deduplicated, i.e. entities are only added to queries and passages once. As a demonstration of the proposed text expansion methods, Figure 6.7 shows how the query expansion is performed using explicit and hashed forms. The added entities provide more precise context and help eliminate ambiguous terms. Figure 6.8 shows the expansion methods on the relevant passage for this query. The relevant passage can be found through our expansion technique. The linking system recognizes that both the query and the passage contain a reference to the entity *Sacagawea*, even though they are spelled differently in the query and the passage.

Reciprocal Rank Fusion. As a second series of experiments, we applied Reciprocal Rank Fusion (RRF) [Cormack et al., 2009] to the runs described above. RRF is a fusion technique that can combine rankings produced by different systems. RRF creates a new ranking by only considering the rank of a document in the input. Given a set of documents D and a set of rankings R , RRF can be computed as:

$$RRF(d \in D) = \sum_{r \in R} \frac{1}{k + r(d)} \quad (6.1)$$

- a.** Introduction. Sacagawea, as everyone knows, was the young Indian woman who, along with her baby, traveled with Lewis and Clark to the Pacific Ocean and back. She was a great help to the expedition and many organizations are preparing celebrations to commemorate the 200-year anniversary of the endeavor. y: M. R. Hansen. Sacagawea, as everyone knows, was the young Indian woman who, along with her baby, traveled with Lewis and Clark to the Pacific Ocean and back.
- b.** The same text as shown in a. + *Indian Ocean James Hansen Sacagawea India Oceania William Clark Meriwether Lewis Pacific Ocean*
- c.** The same text as shown in a. + fe6fc8 860324 aa84e6 7847ef 3e3b0e 7d31e0 2d8836 e58bef

Figure 6.8: The relevant passage for the query presented in Figure 6.7; (a) the non-expanded passage, (b) the passage with entity text expansion, and (c) the passage with entity hash expansion. Text expansions are in italics. The MD5 hashes shown in (c) are shortened for formatting.

Here k is a hyperparameter that can be optimized, but we simply used a default value of $k = 60$ for all settings.

This provides us with four new rankings; the RRF of the pairwise combinations of the three rankings described above and the RRF of all three of these runs:

- d. RRF – No Expansion + Entity Text.** RRF fusion of runs **a** and **b**. The run with no expansions and the run with entity text expansions are considered.
- e. RRF – No Expansion + Entity Hash.** RRF fusion of runs **a** and **c**. The run with no expansions and the run with entity hash expansions are considered.
- f. RRF – Entity Text + Entity Hash.** RRF fusion of runs **b** and **c**. The run with entity text expansions and the run with entity hash expansions are considered.
- g. RRF – No Expansion + Entity Text + Entity Hash.** RRF fusion of runs **a**, **b**, and **c**. All three runs are considered.

6.5.2 Experimental Setup

In our experiments, we use MMEAD as a resource to expand queries and passages with entities. The experiments are performed using the MS MARCO v1 passage ranking collection, where only queries containing at least one entity annotation are used. We do not expect meaningful differences for queries without any linked entities, as the expanded query is identical to the original query in that case.

As we expect the linked entities to provide additional semantic information about the queries and passages, we conduct further testing on the obstinate query sets of the MS MARCO Chameleons [Arabzadeh et al., 2021], which consist of challenging queries from the original MS MARCO passage dataset. In general, ranking methods show poor effectiveness in finding relevant matches for these queries. Our testing focuses on the bottom 50% of the worst-performing queries from the subsets of Veiled Chameleon (Hard), Pygmy Chameleon (Harder), and Lesser Chameleon (Hardest), which represent increasing levels of difficulty.

This gives us four query sets on which we evaluate; (1) all queries that contain entity annotations (*dev* – 1984 queries), (2) all queries in the hard subset that contain entity annotations (*hard* – 680 queries), (3) all queries in the harder subset that contain entity annotations (*harder* – 493 queries),

and lastly, (4) all queries in the hardest subset that have entity annotations (*hardest* – 322 queries).

The experiments are evaluated using Mean Reciprocal Rank (MRR) at rank ten and Recall (R) at rank one thousand. MRR@10 is the official metric for the MS MARCO passage ranking task, while R@1000 gives an upper limit on how well re-ranking systems could perform. The Anserini [Yang et al., 2018a] toolkit is used to generate our experiments.

6.5.3 Results

Table 6.2 presents the results of our experiments. If we first look at lines **a-c** in the results table, we can examine the effects of our expansion methods compared to the baseline run. Looking at R@1000, we can see that more relevant passages are found using entity expansion for the *dev* collection and its harder subsets. We do not find additional relevant documents/passages on the *dev* set when we use the entity hashes, and entity text seems to be the better approach. There is, however, no increase in MRR@10 when using this expansion method. Entity expansions help when evaluating using R@1000, especially when the queries are more complex. The difference in recall effectiveness becomes larger the more complex the queries get. MRR@10 only improves when using entity text expansion.

The reciprocal rank fusion methods are presented in lines **d-g**. When using these methods, the R@1000 increases more. Again, the subsets that contain more complex queries tend to benefit more. Regarding R@1000 effectiveness, the best RRF method uses a ranking from the normal, not expanded index, with the index that has been expanded with the entity text. Again, entity text expansion helps recall more than using hash expansion. Although the RRF methods improve recall, MRR@10 does not benefit from RRF when compared to using only one of the expansion techniques.

6.6 Beyond Quantitative Results

In the previous section, we demonstrated the potential value of MMEAD using quantitative evaluations, where we leverage entities to improve retrieval effectiveness in standard benchmark datasets. Beyond these quantitative results, MMEAD can also help enrich interactive search applications in various ways. This section describes a few such examples.

Entity links to Wikidata provide an entrée into the broader world of open-linked data, which enables integration with other existing resources. This allows us to build interesting “mashups” or support search beyond

Table 6.2: Results on the MS MARCO v1 passage collection, using only the queries that have entity annotations. Bolded numbers are the highest achieved effectiveness. Scores with a dagger (†) are significantly better compared to BM25 with no expansion (run *a*), following a paired t-test with Bonferroni correction. For MRR, we have not calculated significance scores due to its ordinal scale [Fuhr, 2018].

		R@1000				MRR@10			
		dev	hard	harder	hardest	dev	hard	harder	hardest
a. BM25	No Expansion	0.9111	0.7855	0.7444	0.6677	0.2413	0.0373	0.0137	0.0000
b. BM25	Entity Text	0.9183	0.8240†	0.7951†	0.7298†	0.2202	0.0385	0.0173	0.0057
c. BM25	Entity Hash	0.9105	0.7980	0.7576	0.6848	0.2199	0.0383	0.0175	0.0052
d. RRF	No Expansion + Entity Text	0.9338†	0.8436†	0.8124†	0.7500†	0.2372	0.0385	0.0163	0.0019
e. RRF	No Expansion + Entity Hash	0.9250†	0.8260†	0.7921†	0.7205†	0.2378	0.0367	0.0152	0.0034
f. RRF	Entity Text & Hash	0.9231	0.8260†	0.7982†	0.7314†	0.2218	0.0375	0.0161	0.0053
g. RRF	No Expansion + Entity Text & Hash	0.9313†	0.8370†	0.8043†	0.7376†	0.2358	0.0391	0.0156	0.0035

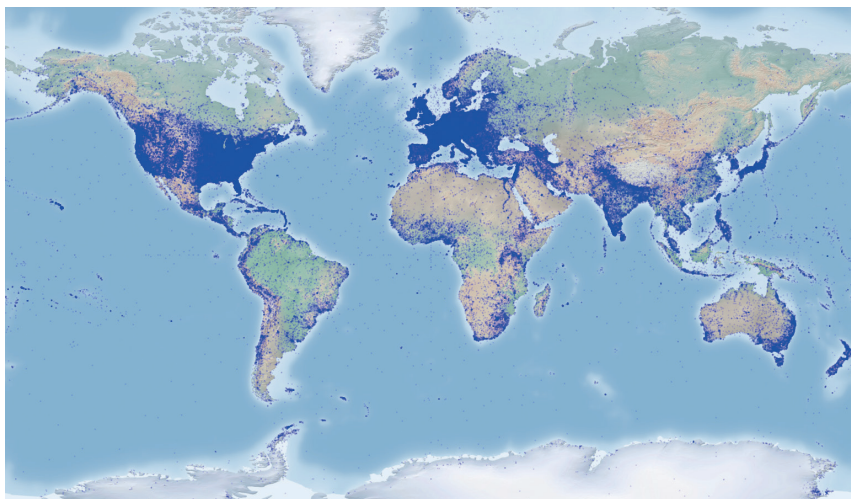


Figure 6.9: Locations of entities found in the MS MARCO v2 passage collection.

simple keyword queries. As a simple example, we can take the entities referenced in MS MARCO, look up the coordinates for geographic entities, and plot them on a map. Figure 6.9 shows a world map with all entities found in the MS MARCO v2 passage collection mapped onto it (each shown with a transparent blue dot). The results are as expected, where the blue dots’ density largely mirrors worldwide population density, although (also as expected) we observe more representation from entities in North America, Europe, and other parts of the world with better access to internet.

Figure 6.9 is a static visualization, but we can take the same underlying data and principles to create interesting interactive demonstrations. Geo-based search is an obvious idea, where users can specify a geographic region – either by dragging a box in an interactive interface to encompass a region of interest, or specifying a geographic entity. For example, the user might ask “Show me content about tourist sites in Paris” and receive passages about the Eiffel Tower in which Paris is not mentioned explicitly. Simple reasoning based on geographic containment relationships on open-linked data resources would be sufficient for answering this query. While it is possible that pretrained transformers might implicitly contain this information, they can never offer the same degree of fine-grained control provided by explicit entity linking.

As a simple demonstration, we have taken MMEAD, reformatted the entity links into RDF, and ingested the results into the QLever SPARQL

```

1 PREFIX rdf: <http://www.w3.org/1999/02/22-rdf-syntax-ns#>
2 PREFIX ex: <http://example.org/>
3 PREFIX schema: <https://schema.org/>
4 PREFIX rdfs: <http://www.w3.org/2000/01/rdf-schema#>
5 PREFIX passage: <http://example.org/passage>
6 PREFIX geo: <http://www.opengis.net/ont/geosparql#>
7 PREFIX wd: <http://www.wikidata.org/entity/>
8 PREFIX wdt: <http://www.wikidata.org/prop/direct/>
9 SELECT ?pid ?content ?entity ?label ?coord
10 WHERE {
11     ?pid rdf:type passage: .
12     ?pid schema:description ?content .
13     ?pid passage:has ?entity .
14     FILTER (regex(?entity, "wikidata", "i"))
15     ?entity rdfs:label ?label .
16     ?entity wdt:P625 ?coord .
17     ?entity wdt:P17 wd:Q142 .
18     FILTER (LANG(?label) = "en")
19 }

```

Figure 6.10: SPARQL query that produces all entities in the passages of the MS MARCO v2 collection that are related to the country of France.

engine [Bast and Buchhold, 2017].² By combining MMEAD with RDF data from Wikidata and OpenStreetMap, we can issue SPARQL queries such as “Show me all passages in MS MARCO related to France”.

The query is shown in Figure 6.10, which gives us 122,316 entities found in the collection that have a connection with France (most of them are located in France). We can automatically show the entities on a map, as presented in Figure 6.11 (showing the first 1000 entities found).

Not all linked entities are located in France, however. For example, some entities are related to France (entities for which France is mentioned in their Wikidata) but located elsewhere in the world. One of the blue dots in Germany is the source of the river *Moselle*. This river starts in Germany by splitting off from the *Rhine*, and then goes through France. Instead of querying for France, we can also query for different countries. Table 6.3

²<https://github.com/ad-freiburg/qllever>, last accessed September 2025

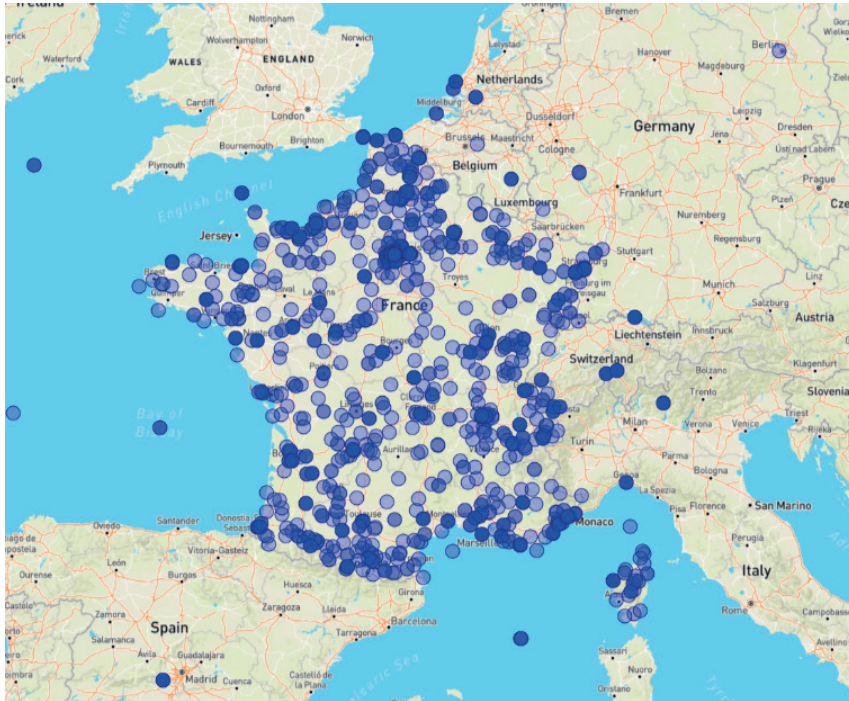


Figure 6.11: First 1000 entities found that are connected to France. Entities are represented with a blue dot on the map.

shows the number of entities found for a sample of countries.

Although we have used SPARQL for this experiment, it would also be possible to express this query using Cypher in GeeseDB. Depending on the exact schema, this would look like the query shown in Figure 6.12.

6.7 Conclusion and Future Work

This chapter introduces a new Information Retrieval resource called MMEAD, or MS MARCO Entity Annotations and Disambiguations. MMEAD contains entity annotations for the passages and documents in MS MARCO v1 and v2. These annotations simplify entity-oriented research on the MS MARCO collections. Links have been provided using the REL and BLINK entity linking systems. Using DuckDB, the data can quickly be queried, making the resource easy to use.

To return to our third research question:

RQ3: When does information retrieval research benefit from graph data?

Table 6.3: Number of entities found per country, for example countries where the entity has an English language label.

Country	WikiData ID	# Entities
United States	Q30	3,429,889
Canada	Q16	170,833
France	Q146	122,316
New Zealand	Q664	19,094
Peru	Q419	16,448
Iran	Q794	13,633
Ecuador	Q736	10,588
South Korea	Q884	9,718
Monaco	Q235	8,546
Singapore	Q334	6,597

```

1 MATCH (pid:Passage)-[:has]->(entity)
2 WHERE entity.uri CONTAINS "wikidata"
3 MATCH (pid)-[:description]->(content)
4 MATCH (entity)-[:label {lang: 'en'}]->(label)
5 MATCH (entity)-[:P625]->(coord)
6 MATCH (entity)-[:P17]->(:Country)
7 RETURN pid, content, entity, label, coord

```

Figure 6.12: Cypher query with similar behaviour as the SPARQL query shown in 6.10.

With MMEAD, we support information retrieval research that combines deep learning and entity graph information. Entity links can be considered as a bi-partite graph where entities and documents are nodes, that are connected when entities appear in documents. Using these entities for query and passage expansion, we showed experimentally a significant increase in recall effectiveness. When using reciprocal rank fusion, the effectiveness difference becomes even more prominent and recall increases significantly with the new relevant passages found. The question remains open whether these passages previously missed, can be ranked higher by new retrieval models.

We also demonstrated that our resource can enrich interactive search applications. In particular, we present a demo where all entities related to geographical locations can be positioned on a map. Specifically we queried the data as a RDF triple graph.

Using the MMEAD format, releasing entity links for collections beyond MS MARCO is also possible. We already showed that using entity links improves recall when using the linked entities for query expansion. What the effects are when training, e.g., dense passage retrieval methods that include the entity links, is yet to be investigated – an exciting research opportunity that lies ahead.

Chapter 7

Conclusion

We can only see a short distance ahead, but we can see plenty there that needs to be done.

Alan Turing - 1950

This thesis will be concluded using the research questions as a guide. We will look once more at the work presented in the individual chapters, and tie it together with the research questions. Finally, I want to reflect on the work described in this thesis and present research opportunities that follow this work.

7.1 Contributions

Recall the main research question and its subquestions:

- **Research Question: How can information retrieval benefit from graph databases and graph query languages?**
- *Research Question 1: What are the benefits of using relational databases for information retrieval?*
- *Research Question 2: Can we extend the benefits from using relational databases for information retrieval to using graph databases while being able to express graph-related problems easier?*
- *Research Question 3: When does information retrieval research benefit from graph data?*

Chapter 3 introduces the history of using relational databases for information retrieval. One of the latter attempts, using columnar databases for retrieval, is highlighted. This approach is re-implemented as a prototype system: “OldDog”. OldDog is used for a reproduction experiment, where many expressions of BM25 are compared against each other. Because OldDog was built using a relational database, we could express the logical model of the ranking functions separately from their physical models. This way, we could find the differences between models while keeping the data the same. This work demonstrates the usefulness of using relational databases for reproduction experiments in information retrieval. Looking at research question 1: *What are the benefits of using relational databases for information retrieval?*, we have demonstrated that relational databases provide a framework for easily comparing different ranking methods, as shown in the reproduction study. The system is reasonably efficient, making it easy to set up for new experiments.

Chapter 4 takes the data model for information retrieval using relational databases, as presented in Chapter 3, and extends this model to a graph data model. This model is implemented in GeeseDB, a prototype graph database for information retrieval. GeeseDB is built on top of DuckDB and uses its column-oriented tables and the fact that it can be run in-process. With a robust backend, we can express graph queries for information retrieval and run them on GeeseDB. GeeseDB supports all functionalities that OldDog already supports, plus it makes the expression of more complicated problems through the graph framework easier. Considering research question 2: *Can we extend the benefits from using relational databases for information retrieval to using graph databases, while being able to express graph-related problems easier?*, we find that, with GeeseDB, it is possible to take all the benefits from using relational databases for information retrieval, while making it possible to express more complex problems through a graph query language.

Chapter 5 concerns improvements on the Radboud Entity Linking (REL) system. Some unforeseen issues were encountered when trying to deploy this system to annotate a large document collection. The time it took to tag the whole collection was much larger than expected, making it unfeasible to annotate the whole corpus in a reasonable time. By improving the efficiency of several components of the software and including a batch extension, it was possible to speed up parts of the software by a factor of ten. After deploying these improvements, it was possible to use the REL software for large corpora.

Chapter 6 continues with the system created in Chapter 5 and deploys

it to annotate the large MS MARCO document and passage collections. A specification on how to share these annotations, independent of the linking system, is proposed. Following this specification, we make the annotations created by the REL system publicly available. We also publish data created by the BLINK entity linking system for the MS MARCO v1 passage collection. Using the annotations by REL, we show on the MS MARCO v1 collection that query expansion with entity annotations significantly improves recall for complicated queries. Although this is just a simple experiment, we show that these annotations contain valuable information for finding relevant documents, and they could be beneficial for more sophisticated methods. Also, through a demonstration that combines entity annotations with geographical information, we show that entity annotations benefit interactive search applications—considering research question 3: *When does information retrieval research benefit from graph data?*. We demonstrate that information retrieval can benefit from graph data, such as an entity graph. Employing it for retrieval techniques leads to an increase in retrieval effectiveness. Also, the graph model is ideal for interactive search systems.

Finally, to address the main research question; **How can information retrieval benefit from graph databases and graph query languages?**. This thesis demonstrates the usefulness of graph databases for information retrieval by creating a prototype system, GeeseDB, that is directly useful for information retrieval. With GeeseDB, it is possible to set up retrieval experiments using the graph model quickly. It inherits all benefits of using relational databases for information retrieval. We show that we can increase retrieval effectiveness by using a graph of entities. This graph of entities can improve interactive search applications as well.

7.2 Future Work

Following the work presented in this thesis, interesting research opportunities lie ahead. In particular, I would like to highlight three directions that seem promising to improve retrieval:

Firstly, an obvious follow-up on the work presented in this thesis is a study that combines the work of Chapters 4 and 6. In Chapter 4, we introduced GeeseDB, a graph-based system for information retrieval. In Chapter 6, we use graph data produced by the REL system. It would be interesting to see if we can reproduce the results of Chapter 6 using GeeseDB. We used an SPARQL engine for the demonstration in Chapter 6. While convenient for the purpose of this demonstration, it would also be

interesting to see how well GeeseDB would hold up in terms of efficiency compared to the SPARQL engine.

The examples shown in Chapter 4 are relatively straightforward. Additional studies should be carried out that consider more complex graph structures. For example, can we employ a graph structure and utilize paths in the graph? These paths could propagate probabilities used to calculate relevancy scores, a strategy that is naturally expressed using graphs.

Lastly, it would be interesting to see if the concepts of graph databases for information retrieval and deep learning for information retrieval can be married together. In the last couple of years, deep learning for information retrieval has become more prominent, especially with the rise of large language models. The representations found by these models can be expressed in graph databases, and both techniques can benefit from each other. Additional studies need to be conducted to determine how this can be done well.

Bibliography

- E. Agichtein, E. Brill, and S. Dumais. Improving Web Search Ranking by Incorporating User Behavior Information. In *Proceedings of the 29th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '06, page 19–26, New York, NY, USA, 2006. Association for Computing Machinery. ISBN 1595933697. URL <https://doi.org/10.1145/1148170.1148177>.
- A. Akbik, T. Bergmann, D. Blythe, K. Rasul, S. Schweter, and R. Vollgraf. FLAIR: An Easy-to-Use Framework for State-of-the-Art NLP. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (Demonstrations)*, NAACL '19, pages 54–59, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-4010. URL <https://aclanthology.org/N19-4010>.
- R. Angles. The Property Graph Database Model. In *Proceedings of the 12th Alberto Mendelzon International Workshop on Foundations of Data Management*, AMW '18, Aachen, 2018. CEUR-WS.org.
- Apache Software Foundation. Giraph, 2012. URL <https://giraph.apache.org/>.
- Apache Software Foundation. Lucene, 2013. URL <https://lucene.apache.org/>.
- N. Arabzadeh, B. Mitra, and E. Bagheri. MS MARCO Chameleons: Challenging the MS MARCO Leaderboard with Extremely Obstinate Queries. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, CIKM '21, page 4426–4435, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450384469. URL <https://doi.org/10.1145/3459637.3482011>.
- J. Arguello, M. Crane, F. Diaz, J. Lin, and A. Trotman. Report on the SIGIR 2015 Workshop on Reproducibility, Inexplicability, and Generalizability of Results (RIGOR). *SIGIR Forum*, 49 (2):107–116, jan 2016. ISSN 0163-5840. doi: 10.1145/2888422.2888439.
- T. G. Armstrong, A. Moffat, W. Webber, and J. Zobel. Improvements That Don't Add up: Ad-Hoc Retrieval Results since 1998. In *Proceedings of the 18th ACM Conference on Information and Knowledge Management*, CIKM '09, page 601–610, New York, NY, USA, 2009. Association for Computing Machinery. ISBN 9781605585123. URL <https://doi.org/10.1145/1645953.1646031>.
- R. Baeza-Yates, B. Ribeiro-Neto, et al. *Modern information retrieval*, volume 463. ACM press New York, 1999.
- P. Bajaj, D. Campos, N. Craswell, L. Deng, J. Gao, X. Liu, R. Majumder, B. M. Andrew McNamara, T. Nguyen, M. Rosenberg, X. Song, A. Stoica, S. Tiwary, and T. Wang. MS MARCO: A Human Generated MACHine Reading COMprehension Dataset. In *InCoCo@NIPS*, 2016.

- K. Balog. *Entity-Oriented Search*. Springer Nature, Gewerbestrasse 11, 6330 Cham, Switzerland, 2018.
- K. Balog, H. Ramampiaro, N. Takhirov, and K. Nørvåg. Multi-Step Classification Approaches to Cumulative Citation Recommendation. In *Proceedings of the 10th Conference on Open Research Areas in Information Retrieval*, OAIR '13, page 121–128, Paris, France, 2013. Le centre de hautes études internationales d’informatique documentaire. ISBN 9782905450098.
- H. Bast and B. Buchhold. QLever: A Query Engine for Efficient SPARQL+Text Search. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, CIKM '17, page 647–656, New York, NY, USA, 2017. Association for Computing Machinery. ISBN 9781450349185.
- H. Bast, M. Hertel, and N. Prange. A Fair and In-Depth Evaluation of Existing End-to-End Entity Linking Systems. In H. Bouamor, J. Pino, and K. Bali, editors, *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 6659–6672, Singapore, Dec. 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.411. URL <https://aclanthology.org/2023.emnlp-main.411>.
- P. Boers, C. Kamphuis, and A. P. de Vries. Radboud University at TREC 2020. In *NIST Special Publication 1266: The Twenty-Ninth Text REtrieval Conference Proceedings (TREC 2020)*, TREC'20, Gaithersburg, Maryland, 2020. [SI]: NIST. URL <https://trec.nist.gov/pubs/trec29/papers/RUIR.N.pdf>.
- P. Boldi and S. Vigna. MG4J at TREC 2005. In *The Fourteenth Text REtrieval Conference (TREC 2005) Proceedings*, number SP 500-266 in Special Papers. NIST, 2005. URL <http://mg4j.di.unimi.it/>.
- P. A. Boncz. *A Next-Generation DBMS Kernel For Query-Intensive Applications*. PhD thesis, University of Amsterdam, May 2002.
- P. A. Boncz, M. Zukowski, and N. Nes. MonetDB/X100: Hyper-Pipelining Query Execution. In *Proceedings of the Second Biennial Conference on Innovative Data Systems Research*, CIDR'05, pages 225–237. www.cidrdb.org, 2005.
- A. Bosselut, R. Le Bras, and Y. Choi. Dynamic Neuro-Symbolic Knowledge Graph Construction for Zero-shot Commonsense Question Answering. In *Proceedings of The Thirty-Fifth AAAI Conference on Artificial Intelligence*, volume 35 of AAAI '20, pages 4923–4931, 2021.
- A. Z. Broder, D. Carmel, M. Herscovici, A. Soffer, and J. Zien. Efficient query evaluation using a two-level retrieval process. In *Proceedings of the Twelfth International Conference on Information and Knowledge Management*, CIKM '03, page 426–434, New York, NY, USA, 2003. Association for Computing Machinery. ISBN 1581137230. URL <https://doi.org/10.1145/956863.956944>.
- C. J. C. Burges. From RankNet to LambdaRank to LambdaMART: An Overview. Technical report, Microsoft Research, 2010. URL http://research.microsoft.com/en-us/um/people/cburges/tech_reports/MSR-TR-2010-82.pdf.
- A. Câmara and C. Macdonald. Dockerising Terrier for The Open-Source IR Replicability Challenge (OSIRRC 2019). In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 26–30, Aachen, 2019. CEUR-WS.org. URL <https://ceur-ws.org/Vol-2409/docker02.pdf>.
- R. Cattell. Scalable SQL and NoSQL data stores. *SIGMOD Rec.*, 39(4):12–27, May 2011. ISSN 0163-5808. URL <https://doi.org/10.1145/1978915.1978919>.

- F. Chang, J. Dean, S. Ghemawat, W. C. Hsieh, D. A. Wallach, M. Burrows, T. Chandra, A. Fikes, and R. E. Gruber. Bigtable: A Distributed Storage System for Structured Data. In *7th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, pages 205–218, 2006.
- S. Chatterjee and L. Dietz. BERT-ER: Query-Specific BERT Entity Representations for Entity Ranking. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '22, page 1466–1477, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450387323. doi: 10.1145/3477495.3531944.
- S. Chaudhuri and G. Weikum. Rethinking Database System Architecture: Towards a Self-Tuning RISC-Style Database System. In *Proceedings of the 26th International Conference on Very Large Data Bases*, VLDB '00, page 1–10, San Francisco, CA, USA, 2000. Morgan Kaufmann Publishers Inc. ISBN 1558607153.
- S. Chaudhuri, R. Ramakrishnan, and G. Weikum. Integrating DB and IR Technologies: What is the Sound of One Hand Clapping? In *Proceedings of the Second Biennial Conference on Innovative Data Systems Research*, CIDR'05, 2005.
- K. W. Church and W. A. Gale. Poisson Mixtures. *Natural Language Engineering*, 1(2):163–190, 1995. doi: 10.1017/S1351324900000139.
- R. Clancy, Z. Akkalyoncu Yilmaz, Z. Z. Wu, and J. Lin. University of Waterloo Docker Images for OSIRRC at SIGIR 2019. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, page 36, Aachen, 2019a. CEUR-WS.org. URL <https://ceur-ws.org/Vol-2409/docker04.pdf>.
- R. Clancy, N. Ferro, C. Hauff, J. Lin, T. Sakai, and Z. Z. Wu. The SIGIR 2019 Open-Source IR Replicability Challenge (OSIRRC 2019). In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR'19, page 1432–1434, New York, NY, USA, 2019b. Association for Computing Machinery. ISBN 9781450361729. doi: 10.1145/3331184.3331647.
- G. P. Copeland and S. N. Khoshafian. A decomposition storage model. In *Proceedings of the 1985 ACM SIGMOD International Conference on Management of Data*, SIGMOD '85, page 268–279, New York, NY, USA, 1985. Association for Computing Machinery. ISBN 0897911601. doi: 10.1145/318898.318923. URL <https://doi.org/10.1145/318898.318923>.
- G. V. Cormack, C. L. A. Clarke, and S. Buettcher. Reciprocal Rank Fusion Outperforms Condorcet and Individual Rank Learning Methods. In *Proceedings of the 32nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '09, page 758–759, New York, NY, USA, 2009. Association for Computing Machinery. ISBN 9781605584836. URL <https://doi.org/10.1145/1571941.1572114>.
- R. Cornacchia, S. Héman, M. Zukowski, A. P. Vries, and P. Boncz. Flexible and Efficient IR Using Array Databases. *The VLDB Journal*, 17(1):151–168, jan 2008. ISSN 1066-8888. doi: 10.1007/s00778-007-0071-0.
- N. Craswell, B. Mitra, E. Yilmaz, D. Campos, E. M. Voorhees, and I. Soboroff. TREC Deep Learning Track: Reusable Test Collections in the Large Data Regime. In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '21, page 2369–2375, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450380379. URL <https://doi.org/10.1145/3404835.3463249>.
- J. Dalton, L. Dietz, and J. Allan. Entity Query Feature Expansion Using Knowledge Base Links. In *Proceedings of the 37th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '14, page 365–374, New York, NY, USA, 2014. Association for Computing Machinery. ISBN 9781450322577. doi: 10.1145/2600428.2609628.

- J. Dalton, C. Xiong, and J. Callan. CAsT 2020: The Conversational Assistance Track Overview. In *The Twenty-Ninth Text REtrieval Conference (TREC 2020) Proceedings*, Gaithersburg, Maryland, USA, 2021. NIST.
- N. De Cao, G. Izacard, S. Riedel, and F. Petroni. Autoregressive Entity Retrieval. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net, 2021. URL <https://openreview.net/forum?id=5k8F6UU39V>.
- A. Dean-Hall, C. L.A. Clarke, J. Kamps, P. Thomas, N. Simone, and E. Voorhees. Overview of the TREC 2013 Contextual Suggestion Track. In *Proceedings of The Twenty-Second Text REtrieval Conference (TREC 2013) Proceedings*, TREC '13, Gaithersburg, Maryland, USA, 2014. National Institute for Standards and Technology (NIST).
- R. Deveaud, M.-D. Albakour, C. Macdonald, and I. Ounis. On the Importance of Venue-Dependent Features for Learning to Rank Contextual Suggestions. In *Proceedings of the 23rd ACM International Conference on Conference on Information and Knowledge Management, CIKM '14*, page 1827–1830, New York, NY, USA, 2014. Association for Computing Machinery. ISBN 9781450325981. doi: 10.1145/2661829.2661956.
- J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, NAACL '19, pages 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1423. URL <https://aclanthology.org/N19-1423>.
- E. W. Dijkstra. A Note on Two Problems in Connexion with Graphs. *Numer. Math.*, 1(1): 269–271, Dec. 1959. ISSN 0029-599X. URL <https://doi.org/10.1007/BF01386390>.
- L. Dolamic and J. Savoy. When Stopword Lists Make the Difference. *J. Am. Soc. Inf. Sci. Technol.*, 61(1):200–203, Jan. 2010. ISSN 1532-2882.
- P. Ferragina and U. Scaiella. TAGME: On-the-Fly Annotation of Short Text Fragments (by Wikipedia Entities). In *Proceedings of the 19th ACM International Conference on Information and Knowledge Management, CIKM '10*, page 1625–1628, New York, NY, USA, 2010. Association for Computing Machinery. ISBN 9781450300995. doi: 10.1145/1871437.1871689.
- D. Ferrucci. Introduction to “This is Watson”. *IBM Journal of Research and Development*, 56: 1:1–1:15, 05 2012. doi: 10.1147/JRD.2012.2184356.
- N. Francis, A. Green, P. Guagliardo, L. Libkin, T. Lindaaker, V. Marsault, S. Plantikow, M. Rydberg, P. Selmer, and A. Taylor. Cypher: An Evolving Query Language for Property Graphs. In *Proceedings of the 2018 International Conference on Management of Data, SIGMOD '18*, page 1433–1445, New York, NY, USA, 2018. Association for Computing Machinery. ISBN 9781450347037. doi: 10.1145/3183713.3190657.
- N. Fuhr. Models for Integrated Information Retrieval and Database Systems. *IEEE Data Engineering Bulletin*, 19(1):3–13, 1996.
- N. Fuhr. Some Common Mistakes In IR Evaluation, And How They Can Be Avoided. *SIGIR Forum*, 51(3):32–41, 2018. ISSN 0163-5840. URL <https://doi.org/10.1145/3190580.3190586>.
- N. Fuhr and T. Rölleke. A Probabilistic Relational Algebra for the Integration of Information Retrieval and Database Systems. *ACM Trans. Inf. Syst.*, 15(1):32–66, jan 1997. ISSN 1046-8188. doi: 10.1145/239041.239045.

- L. Gao, Z. Dai, T. Chen, Z. Fan, B. Van Durme, and J. Callan. Complement Lexical Retrieval Model with Semantic Residual Embeddings. In *Advances in Information Retrieval, ECIR '21*, pages 146–160, Cham, 2021. Springer International Publishing. ISBN 978-3-030-72113-8.
- E. J. Gerritse, F. Hasibi, and A. P. de Vries. Graph-Embedding Empowered Entity Retrieval. In *Advances in Information Retrieval: 42nd European Conference on IR Research, ECIR 2020, Lisbon, Portugal, April 14–17, 2020, Proceedings, Part I*, page 97–110, Berlin, Heidelberg, 2020. Springer-Verlag. ISBN 978-3-030-45438-8. doi: 10.1007/978-3-030-45439-5_7.
- E. J. Gerritse, F. Hasibi, and A. P. de Vries. Entity-Aware Transformers for Entity Search. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '22*, page 1455–1465, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450387323. doi: 10.1145/3477495.3531971.
- T. Grabs, K. Böhm, and H.-J. Schek. PowerDB-IR – Scalable Information Retrieval and Storage with a Cluster of Databases. *Knowl. Inf. Syst.*, 6(4):465–505, jul 2004. ISSN 0219-1377.
- R. Guo, P. Sun, E. Lindgren, Q. Geng, D. Simcha, F. Chern, and S. Kumar. Accelerating Large-Scale Inference with Anisotropic Vector Quantization. In *International Conference on Machine Learning*, 2020. URL <https://arxiv.org/abs/1908.10396>.
- F. Hasibi, K. Balog, and S. E. Bratsberg. Exploiting Entity Linking in Queries for Entity Retrieval. In *Proceedings of the 2016 ACM International Conference on the Theory of Information Retrieval, ICTIR '16*, page 209–218, New York, NY, USA, 2016. Association for Computing Machinery. ISBN 9781450344975. doi: 10.1145/2970398.2970406.
- F. Hasibi, K. Balog, and S. E. Bratsberg. Dynamic Factual Summaries for Entity Cards. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '17*, page 773–782, New York, NY, USA, 2017a. Association for Computing Machinery. ISBN 9781450350228. URL <https://doi.org/10.1145/3077136.3080810>.
- F. Hasibi, K. Balog, D. Garigliotti, and S. Zhang. Nordlys: A Toolkit for Entity-Oriented and Semantic Search. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '17*, page 1289–1292, New York, NY, USA, 2017b. Association for Computing Machinery. ISBN 9781450350228. doi: 10.1145/3077136.3084149.
- C. Hauff. Dockerizing Indri for OSIRRC 2019. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 44–46, Aachen, 2019. CEUR-WS.org. URL <https://ceur-ws.org/Vol-2409/docker06.pdf>.
- S. Héman, M. Zukowski, A. P. de Vries, and P. Boncz. MonetDB/X100 at the 2006 TREC TeraByte Track. In *NIST Special Publication: SP 500-272. The Fifteenth Text REtrieval Conference (TREC 2006) Proceedings*, TREC'06, Gaithersburg, Maryland, USA, 2006. [SI]: NIST. URL <https://trec.nist.gov/pubs/trec15/papers/cwi-heman.tera.final.pdf>.
- D. Hiemstra. *Using Language Models for Information Retrieval*. Phd thesis, Universiteit Twente, 2001.
- D. Hiemstra, G. Hendriksen, C. Kamphuis, and A. P. de Vries. Challenges of Index Exchange for Search Engine Interoperability. In *Proceedings of the 5th International Open Search Symposium, OSSYM23*. Open Search Foundation, 2023.

- J. Hoffart, M. A. Yosef, I. Bordino, H. Fürstenau, M. Pinkal, M. Spaniol, B. Taneva, S. Thater, and G. Weikum. Robust Disambiguation of Named Entities in Text. In *Proceedings of the 2011 Conference on Empirical Methods in Natural Language Processing, EMNLP '11*, pages 782–792, Edinburgh, Scotland, UK., July 2011. Association for Computational Linguistics. URL <https://aclanthology.org/D11-1072>.
- S. Humeau, K. Shuster, M.-A. Lachaux, and J. Weston. Poly-encoders: Transformer Architectures and Pre-training Strategies for Fast and Accurate Multi-sentence Scoring. In *Proceedings of the 7th International Conference on Learning Representations, ICLR'19*, New Orleans, LA, USA, may 2019. URL <https://openreview.net/pdf?id=SkxgnnNFvH>.
- K. Järvelin and J. Kekäläinen. Cumulated gain-based evaluation of IR techniques. *ACM Trans. Inf. Syst.*, 20(4):422–446, Oct. 2002. ISSN 1046-8188. doi: 10.1145/582415.582418. URL <https://doi.org/10.1145/582415.582418>.
- F. Jelinek and R. L. Mercer. Interpolated estimation of Markov source parameters from sparse data. *Workshop on Pattern Recognition in Practice*, page 381–402, 1980.
- T. Joachims, L. Granka, B. Pan, H. Hembrooke, and G. Gay. Accurately interpreting clickthrough data as implicit feedback. In *Proceedings of the 28th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '05*, page 154–161, New York, NY, USA, 2005. Association for Computing Machinery. ISBN 1595930345. URL <https://doi.org/10.1145/1076034.1076063>.
- J. Johnson, M. Douze, and H. Jégou. Billion-scale similarity search with GPUs. *IEEE Transactions on Big Data*, 7(3):535–547, 2019.
- C. Kamphuis. Graph Databases for Information Retrieval. In *Advances in Information Retrieval*, pages 608–612, Cham, 2020. Springer International Publishing. ISBN 978-3-030-45442-5.
- C. Kamphuis and A. P. de Vries. Reproducible IR needs an (IR) (graph) query language. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 17–20, Aachen, 2019a. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2409/position03.pdf>.
- C. Kamphuis and A. P. de Vries. The OldDog Docker Image for OSIRRC at SIGIR 2019. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 47–49, Aachen, 2019b. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2409/docker07.pdf>.
- C. Kamphuis and A. P. de Vries. GeeseDB: A Python Graph Engine for Exploration and Search. In *Proceedings of the 2nd International Conference on Design of Experimental Search & Information REtrieval Systems, DESIRES '21*, pages 10–18, Aachen, 2021. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2950/paper-11.pdf>.
- C. Kamphuis, F. Hasibi, A. P. de Vries, and T. Crijns. Radboud University at TREC 2019. In *NIST Special Publication 1250: The Twenty-Eighth Text REtrieval Conference Proceedings (TREC 2019)*, TREC '19, Gaithersburg, Maryland, 2019. [S]: NIST. URL <https://trec.nist.gov/pubs/trec28/papers/RUIR.N.C.pdf>.
- C. Kamphuis, A. P. de Vries, L. Boytsov, and J. Lin. Which BM25 Do You Mean? A Large-Scale Reproducibility Study of Scoring Variants. In *Advances in Information Retrieval, ECIR '20*, pages 28–34, Cham, 2020. Springer International Publishing. ISBN 978-3-030-45442-5.

- C. Kamphuis, F. Hasibi, J. Lin, and A. P. de Vries. REBL: Entity Linking at Scale. In *Proceedings of the 3rd International Conference on Design of Experimental Search & Information REtrieval Systems*, DESIRES '22, Aachen, 2022. CEUR-WS.org. URL <https://desires.dei.unipd.it/2022/papers/paper-08.pdf>.
- C. Kamphuis, A. Lin, S. Yang, J. Lin, A. P. de Vries, and F. Hasibi. MMEAD: MS MARCO Entity Annotations and Disambiguations. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '23, New York, NY, USA, 2023. Association for Computing Machinery. ISBN 978-1-4503-9408-6. doi: 10.1145/3539618.3591887.
- F. Lancaster. *Information Retrieval Systems: Characteristics, Testing and Evaluation*. John Wiley and Sons, New York, 2 edition, 1979.
- F. Lancaster and E. Fayen. *Information Retrieval On-Line*. Melville Publishing Co., Los Angeles, California, 1973.
- P. Le and I. Titov. Improving Entity Linking by Modeling Latent Relations between Mentions. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1595–1604, Melbourne, Australia, July 2018. Association for Computational Linguistics. doi: 10.18653/v1/P18-1148. URL <https://aclanthology.org/P18-1148>.
- J. Lin. A Proposed Conceptual Framework for a Representational Approach to Information Retrieval. *SIGIR Forum*, 55(2), mar 2022. ISSN 0163-5840. URL <https://doi.org/10.1145/3527546.3527552>.
- J. Lin, J. Mackenzie, C. Kamphuis, C. Macdonald, A. Mallia, M. Siedlaczek, A. Trotman, and A. P. de Vries. Supporting Interoperability Between Open-Source Search Engines with the Common Index File Format. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '20, page 2149–2152, New York, NY, USA, 2020a. Association for Computing Machinery. ISBN 9781450380164. doi: 10.1145/3397271.3401404.
- J. Lin, X. Ma, S.-C. Lin, J.-H. Yang, R. Pradeep, and R. Nogueira. Pyserini: A Python Toolkit for Reproducible Information Retrieval Research with Sparse and Dense Representations. In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '21, page 2356–2362, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450380379. doi: 10.1145/3404835.3463238.
- J. Lin, R. Nogueira, and A. Yates. *Pretrained Transformers for Text Ranking*. Springer Cham, 1 edition, 2022. ISBN 978-3-031-02181-7. doi: <https://doi.org/10.1007/978-3-031-02181-7>.
- S. Lin, J. Yang, and J. Lin. Distilling Dense Representations for Ranking using Tightly-Coupled Teachers. *CoRR*, abs/2010.11386, 2020b. URL <https://arxiv.org/abs/2010.11386>.
- T. Lin, Mausam, and O. Etzioni. Entity Linking at Web Scale. In *Proceedings of the Joint Workshop on Automatic Knowledge Base Construction and Web-scale Knowledge Extraction (AKBC-WEKEX)*, pages 84–88, June 2012.
- T.-Y. Liu. Learning to Rank for Information Retrieval. In *Proceedings of the 33rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '10, page 904, New York, NY, USA, 2010. Association for Computing Machinery. ISBN 9781450301534. URL <https://doi.org/10.1145/1835449.1835676>.
- Y. Luan, J. Eisenstein, K. Toutanova, and M. Collins. Sparse, Dense, and Attentional Representations for Text Retrieval. *Transactions of the Association for Computational Linguistics*, 9: 329–345, 2021.

- Y. Lv and C. Zhai. Lower-Bounding Term Frequency Normalization. In *Proceedings of the 20th ACM International Conference on Information and Knowledge Management, CIKM '11*, page 7–16, New York, NY, USA, 2011a. Association for Computing Machinery. ISBN 9781450307178. doi: 10.1145/2063576.2063584.
- Y. Lv and C. Zhai. Adaptive Term Frequency Normalization for BM25. In *Proceedings of the 20th ACM International Conference on Information and Knowledge Management, CIKM '11*, page 1985–1988, New York, NY, USA, 2011b. Association for Computing Machinery. ISBN 9781450307178. doi: 10.1145/2063576.2063871.
- Y. Lv and C. Zhai. When Documents Are Very Long, BM25 Fails! In *Proceedings of the 34th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '11*, page 1103–1104, New York, NY, USA, 2011c. Association for Computing Machinery. ISBN 9781450307574. doi: 10.1145/2009916.2010070.
- C. Macdonald, R. L. Santos, and I. Ounis. On the Usefulness of Query Features for Learning to Rank. In *Proceedings of the 21st ACM International Conference on Information and Knowledge Management, CIKM '12*, page 2559–2562, New York, NY, USA, 2012. Association for Computing Machinery. ISBN 9781450311564. doi: 10.1145/2396761.2398691.
- C. Macdonald, N. Tonellotto, S. MacAvaney, and I. Ounis. PyTerrier: Declarative Experimentation in Python from BM25 to Dense Retrieval. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management, CIKM '21*, page 4526–4533, New York, NY, USA, 2021. Association for Computing Machinery. ISBN 9781450384469. doi: 10.1145/3459637.3482013.
- I. A. Macleod. Text Retrieval and the Relational Model. *Journal of the American Society for Information Science*, 42(3):155–165, 1991. doi: [https://doi.org/10.1002/\(SICI\)1097-4571\(199104\)42:3<155::AID-ASII>3.0.CO;2-H](https://doi.org/10.1002/(SICI)1097-4571(199104)42:3<155::AID-ASII>3.0.CO;2-H).
- G. Malewicz, M. H. Austern, A. J. Bik, J. C. Dehnert, I. Horn, N. Leiser, and G. Czajkowski. Pregel: A System for Large-Scale Graph Processing. In *Proceedings of the 2010 ACM SIGMOD International Conference on Management of Data, SIGMOD '10*, page 135–146, New York, NY, USA, 2010. Association for Computing Machinery. ISBN 9781450300322. URL <https://doi.org/10.1145/1807167.1807184>.
- A. Mallia, M. Siedlaczek, J. Mackenzie, and T. Suel. PISA: Performant Indexes and Search for Academia. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 50–56, Aachen, 2019. CEUR-WS.org. URL <https://ceur-ws.org/Vol-2409/docker08.pdf>.
- P. N. Mendes, M. Jakob, A. García-Silva, and C. Bizer. DBpedia Spotlight: Shedding Light on the Web of Documents. In *Proceedings of the 7th International Conference on Semantic Systems, I-Semantics '11*, page 1–8, 2011.
- A. Moffat and J. Zobel. Self-Indexing Inverted Files for Fast Text Retrieval. *ACM Trans. Inf. Syst.*, 14(4):349–379, Oct. 1996. ISSN 1046-8188. URL <https://doi.org/10.1145/237496.237497>.
- H. Mühleisen, T. Samar, J. Lin, and A. P. de Vries. Old Dogs Are Great at New Tricks: Column Stores for IR Prototyping. In *Proceedings of the 37th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '14*, page 863–866, New York, NY, USA, 2014. Association for Computing Machinery. ISBN 9781450322577. doi: 10.1145/2600428.2609460.
- Open Science Collaboration. Estimating the reproducibility of psychological science. *Science*, 349(6251):aac4716, 2015. doi: 10.1126/science.aac4716.

- I. Ounis, G. Amati, V. Plachouras, B. He, C. Macdonald, and D. Johnson. Terrier Information Retrieval Platform. In *Advances in Information Retrieval*, pages 517–519, Berlin, Heidelberg, 2005. Springer Berlin Heidelberg. ISBN 978-3-540-31865-1.
- A. Overwijk, C. Xiong, and J. Callan. ClueWeb22: 10 Billion Web Documents with Rich Information. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '22, page 3360–3362, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450387323. URL <https://doi.org/10.1145/3477495.3536321>.
- L. Page, S. Brin, R. Motwani, and T. Winograd. The PageRank Citation Ranking: Bringing Order to the Web. Technical Report 1999-66, Stanford InfoLab, November 1999. URL <http://ilpubs.stanford.edu:8090/422/>. Previous number = SIDL-WP-1999-0120.
- M. F. Porter. An algorithm for suffix stripping. *Program: electronic library and information systems*, 14(3):130–137, 03 1980. ISSN 0033-0337. URL <https://doi.org/10.1108/eb046814>.
- M. Raasveldt and H. Mühleisen. DuckDB: An Embeddable Analytical Database. In *Proceedings of the 2019 International Conference on Management of Data*, SIGMOD '19, page 1981–1984, New York, NY, USA, 2019. Association for Computing Machinery. ISBN 9781450356435. doi: 10.1145/3299869.3320212.
- C. Raffel, N. Shazeer, A. Roberts, K. Lee, S. Narang, M. Matena, Y. Zhou, W. Li, and P. J. Liu. Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *Journal of Machine Learning Research*, 21(140):1–67, 2020. URL <http://jmlr.org/papers/v21/20-074.html>.
- R. Reinanda, E. Meij, and M. de Rijke. Mining, Ranking and Recommending Entity Aspects. In *Proceedings of the 38th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '15, page 263–272, New York, NY, USA, 2015. Association for Computing Machinery. ISBN 9781450336215. doi: 10.1145/2766462.2767724.
- S. E. Robertson and K. Spärck Jones. Relevance weighting of search terms. *Journal of the American Society for Information Science*, 27(3):129–146, 1976. doi: <https://doi.org/10.1002/asi.4630270302>.
- S. E. Robertson and H. Zaragoza. The Probabilistic Relevance Framework: BM25 and Beyond. *Found. Trends Inf. Retr.*, 3(4):333–389, apr 2009. ISSN 1554-0669. doi: 10.1561/15000000019.
- S. E. Robertson, S. Walker, S. Jones, M. Hancock-Beaulieu, and M. Gatford. Okapi at TREC-3. In *Overview of the third text Retrieval conference*, TREC-3, Gaithersburg, Maryland, USA, 1994. [S]: NIST. URL <https://trec.nist.gov/pubs/trec3/papers/city.ps.gz>.
- F. Rousseau and M. Vazirgiannis. Composition of TF Normalizations: New Insights on Scoring Functions for Ad Hoc IR. In *Proceedings of the 36th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '13, page 917–920, New York, NY, USA, 2013. Association for Computing Machinery. ISBN 9781450320344. doi: 10.1145/2484028.2484121.
- T. Sakai. On Fuhr's guideline for IR evaluation. *SIGIR Forum*, 54(1), Feb. 2021. ISSN 0163-5840. URL <https://doi.org/10.1145/3451964.3451976>.
- S. Sakr, A. Bonifati, H. Voigt, A. Iosup, K. Ammar, R. Angles, W. Aref, M. Arenas, M. Besta, P. A. Boncz, K. Daudjee, E. D. Valle, S. Dumbrava, O. Hartig, B. Haslhofer, T. Hegeman, J. Hidders, K. Hose, A. Iamnitchi, V. Kalavri, H. Kapp, W. Martens, M. T. Özsu, E. Peukert, S. Plantikow, M. Ragab, M. R. Ripeanu, S. Salihoglu, C. Schulz, P. Selmer, J. F. Sequeda, J. Shinavier, G. Szárnyas, R. Tommasini, A. Tumeo, A. Uta, A. L. Varbanescu, H.-Y. Wu,

- N. Yakovets, D. Yan, and E. Yoneki. The Future is Big Graphs: A Community View on Graph Processing Systems. *Commun. ACM*, 64(9):62–71, aug 2021. ISSN 0001-0782. URL <https://doi.org/10.1145/3434642>.
- G. Salton, A. Wong, and C. S. Yang. A Vector Space Model for Automatic Indexing. *Communications of the ACM*, 18(11):613–620, nov 1975. ISSN 0001-0782. URL <https://doi.org/10.1145/361219.361220>.
- H. Scells and G. Zuccon. ielab at the Open-Source IR Replicability Challenge 2019. In *Proceedings of the Open-Source IR Replicability Challenge co-located with 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval, OSIRRC@SIGIR 2019, Paris, France, July 25, 2019*, pages 57–61, Aachen, 2019. CEUR-WS.org. URL <http://ceur-ws.org/Vol-2409/docker09.pdf>.
- H.-J. Schek and P. Pistor. Data Structures for an Integrated Data Base Management and Information Retrieval System. In *Proceedings of the 8th International Conference on Very Large Data Bases, VLDB '82*, page 197–207, San Francisco, CA, USA, 1982. Morgan Kaufmann Publishers Inc. ISBN 0934613141.
- T. Schoegje, C. Kamphuis, K. Dercksen, D. Hiemstra, T. Pieters, and A. P. de Vries. Exploring task-based query expansion at the TREC-COVID track. *CoRR*, abs/2010.12674, 2020. URL <https://arxiv.org/abs/2010.12674>.
- C. Sciavolino, Z. Zhong, J. Lee, and D. Chen. Simple Entity-Centric Questions Challenge Dense Retrievers. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, EMNLP '21*, pages 6138–6148, Online and Punta Cana, Dominican Republic, Nov. 2021. Association for Computational Linguistics. URL <https://aclanthology.org/2021.emnlp-main.496>.
- D. Shehata. Information Retrieval with Entity Linking. Master's thesis, University of Waterloo, 2022. URL <http://hdl.handle.net/10012/18557>.
- A. Singhal, C. Buckley, and M. Mitra. Pivoted Document Length Normalization. In *Proceedings of the 19th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '96*, page 21–29, New York, NY, USA, 1996. Association for Computing Machinery. ISBN 0897917928. doi: 10.1145/243199.243206.
- I. Soboroff, S. Huang, and D. Harman. TREC 2018 News Track Overview. In *Proceedings of The Twenty-Seventh Text REtrieval Conference, TREC '18*, Gaithersburg, Maryland, USA, 2019. National Institute for Standards and Technology (NIST).
- F. Song and W. B. Croft. A General Language Model for Information Retrieval. In *Proceedings of the Eighth International Conference on Information and Knowledge Management, CIKM '99*, page 316–321, New York, NY, USA, 1999. Association for Computing Machinery. ISBN 1581131461. URL <https://doi.org/10.1145/319950.320022>.
- K. Spärck Jones. A Statistical Interpretation of Term Specificity and its Application in Retrieval. *Journal of Documentation*, 60:493–502, 1972.
- V. I. Spitzkovsky and A. X. Chang. A Cross-Lingual Dictionary for English Wikipedia Concepts. In *Proceedings of the Eighth International Conference on Language Resources and Evaluation, LREC '12*, pages 3168–3175, Istanbul, Turkey, May 2012. European Language Resources Association (ELRA). URL http://www.lrec-conf.org/proceedings/lrec2012/pdf/266_Paper.pdf.
- T. Strohman, D. Metzler, H. Turtle, and W. B. Croft. Indri: A language model-based search engine for complex queries. In *Proceedings of the International Conference on Intelligent Analysis*, number 6 in ICIA'05, pages 2–6. Washington, DC., 2005. URL <http://ciir.cs.umass.edu/pubfiles/ir-407.pdf>.

- N. Thakur, N. Reimers, A. Rüklé, A. Srivastava, and I. Gurevych. BEIR: A Heterogeneous Benchmark for Zero-shot Evaluation of Information Retrieval Models. In *Proceedings of the Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*, NeurIPS '21, 2021. URL <https://openreview.net/forum?id=wCu6T5xFJeJ>.
- N. Thakur, N. Reimers, and J. Lin. Injecting Domain Adaptation with Learning-to-hash for Effective and Efficient Zero-shot Dense Retrieval. In *Workshop on Reaching Efficiency in Neural Information Retrieval*, ReNeuIR'23, July 2023.
- H. D. Tran and A. Yates. Dense Retrieval with Entity Views. In *Proceedings of the 31st ACM International Conference on Information & Knowledge Management, CIKM '22*, page 1955–1964, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450392365. URL <https://doi.org/10.1145/3511808.3557285>.
- A. Trotman, X. Jia, and M. Crane. Towards an Efficient and Effective Search Engine. In *Proceedings of the SIGIR 2012 Workshop on Open Source Information Retrieval*, OSIR@SIGIR'12, pages 40–47, 2012.
- A. Trotman, A. Puurula, and B. Burgess. Improvements to BM25 and Language Models Examined. In *Proceedings of the 2014 Australasian Document Computing Symposium, ADCS '14*, page 58–65, New York, NY, USA, 2014. Association for Computing Machinery. ISBN 9781450330008. doi: 10.1145/2682862.2682863.
- J. M. van Hulst, F. Hasibi, K. Dercksen, K. Balog, and A. P. de Vries. REL: An Entity Linker Standing on the Shoulders of Giants. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '20, page 2197–2200, New York, NY, USA, 2020. Association for Computing Machinery. ISBN 9781450380164. URL <https://doi.org/10.1145/3397271.3401416>.
- C. J. van Rijsbergen. *Information Retrieval*. Butterworths, London, 2 edition, 1979. URL <http://www.dcs.gla.ac.uk/Keith/Preface.html>.
- E. M. Voorhees and D. M. Tice. The TREC-8 Question Answering Track Evaluation. In *Text Retrieval Conference TREC-8*. NIST, 2000.
- J. Webber. A Programmatic Introduction to Neo4j. In *Proceedings of the 3rd Annual Conference on Systems, Programming, and Applications: Software for Humanity, SPLASH '12*, page 217–218, New York, NY, USA, 2012. Association for Computing Machinery. ISBN 9781450315630. URL <https://doi.org/10.1145/2384716.2384777>.
- L. Wu, F. Petroni, M. Josifoski, S. Riedel, and L. Zettlemoyer. Scalable Zero-shot Entity Linking with Dense Entity Retrieval. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing, EMNLP '20*, pages 6397–6407, Online, Nov. 2020. Association for Computational Linguistics. URL <https://aclanthology.org/2020.emnlp-main.519>.
- C. Xiong, J. Callan, and T.-Y. Liu. Word-Entity Duet Representations for Document Ranking. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '17, page 763–772, New York, NY, USA, 2017. Association for Computing Machinery. ISBN 9781450350228. doi: 10.1145/3077136.3080768.
- C. Xu, D. Guo, N. Duan, and J. McAuley. LaPraDoR: Unsupervised Pretrained Dense Retriever for Zero-Shot Text Retrieval. In *Findings of the Association for Computational Linguistics: ACL 2022*, pages 3557–3569, Dublin, Ireland, May 2022. Association for Computational Linguistics. URL <https://aclanthology.org/2022.findings-acl.281>.
- I. Yamada, H. Shindo, H. Takeda, and Y. Takefuji. Joint Learning of the Embedding of Words and Entities for Named Entity Disambiguation. In *Proceedings of the 20th SIGNLL Conference on Computational Natural Language Learning*, pages 250–259, Berlin, Germany, Aug. 2016. Association for Computational Linguistics. URL <https://aclanthology.org/K16-1025>.

- P. Yang and J. Lin. Anserini at TREC 2018: Centre, common core, and news tracks. In *Proceedings of the Twenty-Seventh Text REtrieval Conference (TREC 2018)*, Gaithersburg, MD, TREC '18, Gaithersburg, Maryland, USA, 2019. National Institute for Standards and Technology (NIST).
- P. Yang, H. Fang, and J. Lin. Anserini: Reproducible Ranking Baselines Using Lucene. *Journal of Data and Information Quality*, 10(4), 2018a. ISSN 1936-1955. URL <https://doi.org/10.1145/3239571>.
- W. Yang, K. Lu, P. Yang, and J. Lin. Critically Examining the "Neural Hype": Weak Baselines and the Additivity of Effectiveness Gains from Neural Ranking Models. In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR'19, page 1129–1132, New York, NY, USA, 2019. Association for Computing Machinery. ISBN 9781450361729. doi: 10.1145/3331184.3331340.
- Y. Yang, O. Irsoy, and K. S. Rahman. Collective Entity Disambiguation with Structured Gradient Tree Boosting. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 777–786, New Orleans, Louisiana, June 2018b. Association for Computational Linguistics. doi: 10.18653/v1/N18-1071. URL <https://aclanthology.org/N18-1071>.
- C. Zhai and J. Lafferty. Two-Stage Language Models for Information Retrieval. In *Proceedings of the 25th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '02, page 49–56, New York, NY, USA, 2002. Association for Computing Machinery. ISBN 1581135610. URL <https://doi.org/10.1145/564376.564387>.
- C. Zhai and J. Lafferty. A Study of Smoothing Methods for Language Models Applied to Information Retrieval. *ACM Trans. Inf. Syst.*, 22(2):179–214, Apr. 2004. ISSN 1046-8188. doi: 10.1145/984321.984322. URL <https://doi.org/10.1145/984321.984322>.
- M. Zukowski, M. van de Wiel, and P. Boncz. Vectorwise: A Vectorized Analytical DBMS. In *Proceedings of the 2012 IEEE 28th International Conference on Data Engineering*, ICDE '12, page 1349–1350, USA, 2012. IEEE Computer Society. ISBN 9780769547473. doi: 10.1109/ICDE.2012.148.

Summary

Finding relevant information in a large collection of documents can be challenging, especially when only text is considered when determining relevancy. This research leverages graph data to express information needs that consider more information than just text data. In some cases, instead of using inverted indexes for the data representation in our work, we use database management systems to store data.

First, we show that relational database systems are suited for retrieval experiments. A prototype system we built implements many proposed improvements to the BM25 ranking algorithm. In a large-scale reproduction study, we compare these improvements and find that the differences in effectiveness are smaller than we would expect, given the literature. We can easily change between versions of BM25 by rewriting the SQL query slightly, validating the usefulness of relational databases for reproducible IR research.

Then, we extend the data model to a graph data model. Using a graph data model; we can include more diverse data than just text. We show that we can more easily express complex information needs with a corresponding graph query language than when a relational language is used. This model is built on top of an embedded database system, allowing fast materialization of output data and using it for further steps. One of the aspects we capture in the graph is information about entities. We use the Radboud Entity Linking (REL) system to connect entity information with documents. In order to efficiently annotate a large document collection with REL, we improved its efficiency. After these improvements, we used REL to create annotations for the MS MARCO document and passage collections. We can significantly improve recall for harder MS MARCO queries using these annotations. These entities are also used for an interactive demonstration where the geographical data of entities is used.

Samenvatting

Het vinden van relevante informatie in een grote verzameling van documenten is een zeer uitdagende taak, zeker wanneer alleen tekst in overweging wordt genomen bij het bepalen of een document relevant is. Dit onderzoek maakt gebruik van grafen om informatiebehoeften uit te drukken waar meer in overweging genomen moet worden dan alleen tekst. In sommige gevallen gebruiken we, voor de data representatie, in plaats van *inverted indexes*, data management systemen om data op te slaan.

Allereerst, laten we zien dat relationele database systemen geschikt zijn voor *information retrieval* experimenten. Een door ons gebouwd prototype systeem implementeert een hele reeks verbeteringen aan het BM25-rangschikking algoritme die zijn voorgesteld in de literatuur. In een grootschalig reproductie onderzoek vergelijken we deze verbeteringen en vinden we dat de verschillen in effectiviteit kleiner zijn dan we op grond van de literatuur zouden verwachten. We kunnen gemakkelijk wisselen tussen versies van BM25 door de SQL-query slechts weinig te herschrijven met simpele variaties op een onderliggende SQL query.

Hiermee valideren we het nut van relationele databases voor reproduceerbaar IR-onderzoek. Vervolgens breiden we het datamodel uit naar een graaf datamodel. Met dit graaf datamodel kunnen we meer diverse gegevens uitdrukken dan alleen tekst. We kunnen complexe informatiebehoeften makkelijker uitdrukken met een bijbehorende graaf querytaal, dan wanneer een relationele taal wordt gebruikt. Dit model is gebouwd bovenop een *embedded* database systeem, hierdoor kunnen we data dat geproduceerd wordt door dit systeem snel voor een andere applicatie gebruiken.

Een van de aspecten die we in de graaf vastleggen, is informatie over entiteiten. We gebruiken het Radboud Entity Linking (REL) systeem om entiteit informatie te koppelen aan documenten. Om een grote documentverzameling efficiënt te annoteren met REL hebben we de efficiëntie van REL verbeterd. Na deze verbeteringen hebben we REL gebruikt om annotaties te maken voor de MS MARCO-document- en passage-collecties. Met behulp van deze annotaties kunnen we de *recall* voor moeilijkere MS

MARCO-query's aanzienlijk verbeteren. Deze entiteiten worden ook gebruikt voor een interactieve demonstratie waarbij de geografische gegevens van entiteiten worden gebruikt.

Acknowledgements

On October 1st, 2018, I started my PhD research. The work done in the following years eventually led to this thesis. I attribute this to a trifecta of excellent supervision, collaboration, and support. For this, I would like to thank some people:

First and foremost, thank you, Arjen. Thank you for the opportunity to be your student. Throughout the years, I have learned a lot from you. I have always enjoyed working together with you, and I still do. I will always be appreciative of what you have offered me.

I want to thank the manuscript committee for reading my thesis carefully and providing me with great feedback. Your input really helped me.

I am grateful for the collaborations that helped my research. Jimmy, I really enjoyed working together on a couple of occasions. I liked seeing how you approach research projects. Working with you showed me how I can approach work more structured, helping me become more efficient in my work. Faegheh, our collaborations were a pleasure as well. I have learned a lot from you working together. You have a great attention to detail, without losing sight of the bigger picture.

When I started my research, Arjen and I formed a two-person Information Retrieval research group. Over the years, this changed to a large group of Information Retrieval researchers. Emma, you were the first person to join the group, and we have been office mates throughout our appointments as PhD students. We were in the same situation, making it possible to always share our work experiences with someone going through the same. The Information Retrieval research group grew rapidly, making the job even more delightful. So thank you, Djoerd, Martha, Faegheh, Harrie, Ivan, Koen, Neghin, Hideaki, Norman, Tim, Gijs, Daria, Mohanna, Heydar, Thomas, and Mick.

Our research group was, of course, part of a larger group of researchers in the department of data science at the institute of computing science. All colleagues at DAS were very kind and knowledgeable. I really enjoyed all our lunch sessions together, as well as all board game nights and other

informal events. I am very lucky that I was able to work in such a great environment.

Then, I would say thank you to my current colleagues at Spinque. I really enjoy working at a company with so many great people. At Spinque, I can work on projects where I can use the things I learned during my PhD. Thank you also for your support when finishing my thesis.

Mart, Koen, and Pieter, thank you for being great friends not only during my PhD research, but already throughout the bachelor's and master's programs. Mart, we have known each other even longer. Thank you for being such a good friend for such a long time. To all my other friends, thank you for being my friend. I am very lucky to have you all in my life.

Mijn ouders wil ik bedanken voor een warm en fijn thuis. Op een goed fundament kun je bouwen, en die hebben jullie zeker gelegd. Stijn en Anne, ik ben heel blij met jullie als broer en zus. Jullie inspireren mij dagelijks. Ook mijn schoonfamilie wil ik bedanken, jullie staan altijd voor mij klaar.

Ten slotte, Maudy, bedankt voor jouw onuitputtelijke steun en liefde. Joep, de afgelopen maanden met jou waren geweldig. Ik kijk er naar uit om de rest van het leven samen met jullie door te brengen.

Research Data Management

This thesis research has been carried out under the research data management policy of the Institute for Computing and Information Science of the Radboud University, the Netherlands.¹

The following research datasets have been produced during this PhD research:

- Resources for Chapter 3
 - Code for OldDog [Kamphuis and de Vries, 2019b]:
chriskamphuis/olddog: (v1.0.0). Zenodo. 10.5281/zenodo.3255060
 - Code for the OldDog docker [Kamphuis and de Vries, 2019b]:
osirrc/olddog-docker: (v1.0.0). Zenodo. 10.5281/zenodo.3255060
- Resources for Chapter 4
 - Code for GeeseDB [Kamphuis and de Vries, 2021]:
informagi/GeeseDB: (v0.0.2). Zenodo. 10.5281/zenodo.7892326
- Resources for Chapter 5
 - Code for REBL [Kamphuis et al., 2022]:
informagi/REBL: (v0.0.1). Zenodo. 10.5281/zenodo.7892359
- Resources for Chapter 6
 - Code for MMEAD [Kamphuis et al., 2023]:
informagi/mmead: (v0.1.0). Zenodo. 10.5281/zenodo.7897027
 - Data for MMEAD [Kamphuis et al., 2023]:
informagi/mmead: (v0.1.0). Zenodo. 10.5281/zenodo.7896782

¹<https://www.ru.nl/rdm/>, last accessed September 2025

Curriculum Vitæ

Chris Kamphuis

- 1993 Born March 22th, Oldenzaal
- 2005–2012 VWO at the Twents Carmelcollege De Thij, Oldenzaal
- 2012–2016 Bachelor's degree Artificial Intelligence at the Radboud University, Nijmegen
- 2016–2018 Master's degree (cum laude) Computing Science (Data Science track) at the Radboud University, Nijmegen
- 2018–2023 PhD candidate at the Institute for Computing and Information Sciences (iCIS) of Radboud University, Nijmegen
- 2024– Data Scientist at Spinque, Utrecht

SIKS Dissertations

- 2016 01 Syed Saiden Abbas (RUN), Recognition of Shapes by Humans and Machines
- 02 Michiel Christiaan Meulendijk (UU), Optimizing medication reviews through decision support: prescribing a better pill to swallow
- 03 Maya Sappelli (RUN), Knowledge Work in Context: User Centered Knowledge Worker Support
- 04 Laurens Rietveld (VUA), Publishing and Consuming Linked Data
- 05 Evgeny Sherkhonov (UvA), Expanded Acyclic Queries: Containment and an Application in Explaining Missing Answers
- 06 Michel Wilson (TUD), Robust scheduling in an uncertain environment
- 07 Jeroen de Man (VUA), Measuring and modeling negative emotions for virtual training
- 08 Matje van de Camp (TiU), A Link to the Past: Constructing Historical Social Networks from Unstructured Data
- 09 Archana Nottamkandath (VUA), Trusting Crowdsourced Information on Cultural Artefacts
- 10 George Karafotias (VUA), Parameter Control for Evolutionary Algorithms
- 11 Anne Schuth (UvA), Search Engines that Learn from Their Users
- 12 Max Knobbout (UU), Logics for Modelling and Verifying Normative Multi-Agent Systems
- 13 Nana Baah Gyan (VUA), The Web, Speech Technologies and Rural Development in West Africa - An ICT4D Approach
- 14 Ravi Khadka (UU), Revisiting Legacy Software System Modernization
- 15 Steffen Michels (RUN), Hybrid Probabilistic Logics - Theoretical Aspects, Algorithms and Experiments

- 16 Guangliang Li (UvA), Socially Intelligent Autonomous Agents that Learn from Human Reward
- 17 Berend Weel (VUA), Towards Embodied Evolution of Robot Organisms
- 18 Albert Meroño Peñuela (VUA), Refining Statistical Data on the Web
- 19 Julia Efremova (TU/e), Mining Social Structures from Genealogical Data
- 20 Daan Odijk (UvA), Context & Semantics in News & Web Search
- 21 Alejandro Moreno Céleri (UT), From Traditional to Interactive Playspaces: Automatic Analysis of Player Behavior in the Interactive Tag Playground
- 22 Grace Lewis (VUA), Software Architecture Strategies for Cyber-Foraging Systems
- 23 Fei Cai (UvA), Query Auto Completion in Information Retrieval
- 24 Brend Wanders (UT), Repurposing and Probabilistic Integration of Data; An Iterative and data model independent approach
- 25 Julia Kiseleva (TU/e), Using Contextual Information to Understand Searching and Browsing Behavior
- 26 Dilhan Thilakarathne (VUA), In or Out of Control: Exploring Computational Models to Study the Role of Human Awareness and Control in Behavioural Choices, with Applications in Aviation and Energy Management Domains
- 27 Wen Li (TUD), Understanding Geo-spatial Information on Social Media
- 28 Mingxin Zhang (TUD), Large-scale Agent-based Social Simulation - A study on epidemic prediction and control
- 29 Nicolas Höning (TUD), Peak reduction in decentralised electricity systems - Markets and prices for flexible planning
- 30 Ruud Mattheij (TiU), The Eyes Have It
- 31 Mohammad Khelghati (UT), Deep web content monitoring
- 32 Eelco Vriezekolk (UT), Assessing Telecommunication Service Availability Risks for Crisis Organisations
- 33 Peter Bloem (UvA), Single Sample Statistics, exercises in learning from just one example
- 34 Dennis Schunselaar (TU/e), Configurable Process Trees: Elicitation, Analysis, and Enactment
- 35 Zhaochun Ren (UvA), Monitoring Social Media: Summarization, Classification and Recommendation
- 36 Daphne Karreman (UT), Beyond R2D2: The design of nonverbal interaction behavior optimized for robot-specific morphologies

- 37 Giovanni Sileno (UvA), Aligning Law and Action - a conceptual and computational inquiry
 - 38 Andrea Minuto (UT), Materials that Matter - Smart Materials meet Art & Interaction Design
 - 39 Merijn Bruijnes (UT), Believable Suspect Agents; Response and Interpersonal Style Selection for an Artificial Suspect
 - 40 Christian Detweiler (TUD), Accounting for Values in Design
 - 41 Thomas King (TUD), Governing Governance: A Formal Framework for Analysing Institutional Design and Enactment Governance
 - 42 Spyros Martzoukos (UvA), Combinatorial and Compositional Aspects of Bilingual Aligned Corpora
 - 43 Saskia Koldijk (RUN), Context-Aware Support for Stress Self-Management: From Theory to Practice
 - 44 Thibault Sellam (UvA), Automatic Assistants for Database Exploration
 - 45 Bram van de Laar (UT), Experiencing Brain-Computer Interface Control
 - 46 Jorge Gallego Perez (UT), Robots to Make you Happy
 - 47 Christina Weber (UL), Real-time foresight - Preparedness for dynamic innovation networks
 - 48 Tanja Buttler (TUD), Collecting Lessons Learned
 - 49 Gleb Polevoy (TUD), Participation and Interaction in Projects. A Game-Theoretic Analysis
 - 50 Yan Wang (TiU), The Bridge of Dreams: Towards a Method for Operational Performance Alignment in IT-enabled Service Supply Chains
-
- 2017 01 Jan-Jaap Oerlemans (UL), Investigating Cybercrime
 - 02 Sjoerd Timmer (UU), Designing and Understanding Forensic Bayesian Networks using Argumentation
 - 03 Daniël Harold Telgen (UU), Grid Manufacturing; A Cyber-Physical Approach with Autonomous Products and Reconfigurable Manufacturing Machines
 - 04 Mrunal Gawade (CWI), Multi-core Parallelism in a Column-store
 - 05 Mahdiah Shadi (UvA), Collaboration Behavior
 - 06 Damir Vandic (EUR), Intelligent Information Systems for Web Product Search
 - 07 Roel Bertens (UU), Insight in Information: from Abstract to Anomaly

- 08 Rob Konijn (VUA), Detecting Interesting Differences: Data Mining in Health Insurance Data using Outlier Detection and Subgroup Discovery
- 09 Dong Nguyen (UT), Text as Social and Cultural Data: A Computational Perspective on Variation in Text
- 10 Robby van Delden (UT), (Steering) Interactive Play Behavior
- 11 Florian Kunneman (RUN), Modelling patterns of time and emotion in Twitter #anticipointment
- 12 Sander Leemans (TU/e), Robust Process Mining with Guarantees
- 13 Gijs Huisman (UT), Social Touch Technology - Extending the reach of social touch through haptic technology
- 14 Shoshannah Tekofsky (TiU), You Are Who You Play You Are: Modelling Player Traits from Video Game Behavior
- 15 Peter Berck (RUN), Memory-Based Text Correction
- 16 Aleksandr Chuklin (UvA), Understanding and Modeling Users of Modern Search Engines
- 17 Daniel Dimov (UL), Crowdsourced Online Dispute Resolution
- 18 Ridho Reinanda (UvA), Entity Associations for Search
- 19 Jeroen Vuurens (UT), Proximity of Terms, Texts and Semantic Vectors in Information Retrieval
- 20 Mohammadbashir Sedighi (TUD), Fostering Engagement in Knowledge Sharing: The Role of Perceived Benefits, Costs and Visibility
- 21 Jeroen Linssen (UT), Meta Matters in Interactive Storytelling and Serious Gaming (A Play on Worlds)
- 22 Sara Magliacane (VUA), Logics for causal inference under uncertainty
- 23 David Graus (UvA), Entities of Interest — Discovery in Digital Traces
- 24 Chang Wang (TUD), Use of Affordances for Efficient Robot Learning
- 25 Veruska Zamborlini (VUA), Knowledge Representation for Clinical Guidelines, with applications to Multimorbidity Analysis and Literature Search
- 26 Merel Jung (UT), Socially intelligent robots that understand and respond to human touch
- 27 Michiel Joosse (UT), Investigating Positioning and Gaze Behaviors of Social Robots: People's Preferences, Perceptions and Behaviors
- 28 John Klein (VUA), Architecture Practices for Complex Contexts

- 29 Adel Alhuraibi (TiU), From IT-BusinessStrategic Alignment to Performance: A Moderated Mediation Model of Social Innovation, and Enterprise Governance of IT"
 - 30 Wilma Latuny (TiU), The Power of Facial Expressions
 - 31 Ben Ruijl (UL), Advances in computational methods for QFT calculations
 - 32 Thaer Samar (RUN), Access to and Retrievability of Content in Web Archives
 - 33 Brigit van Loggem (OU), Towards a Design Rationale for Software Documentation: A Model of Computer-Mediated Activity
 - 34 Maren Scheffel (OU), The Evaluation Framework for Learning Analytics
 - 35 Martine de Vos (VUA), Interpreting natural science spreadsheets
 - 36 Yuanhao Guo (UL), Shape Analysis for Phenotype Characterisation from High-throughput Imaging
 - 37 Alejandro Montes Garcia (TU/e), WiBAF: A Within Browser Adaptation Framework that Enables Control over Privacy
 - 38 Alex Kayal (TUD), Normative Social Applications
 - 39 Sara Ahmadi (RUN), Exploiting properties of the human auditory system and compressive sensing methods to increase noise robustness in ASR
 - 40 Altaf Hussain Abro (VUA), Steer your Mind: Computational Exploration of Human Control in Relation to Emotions, Desires and Social Support For applications in human-aware support systems
 - 41 Adnan Manzoor (VUA), Minding a Healthy Lifestyle: An Exploration of Mental Processes and a Smart Environment to Provide Support for a Healthy Lifestyle
 - 42 Elena Sokolova (RUN), Causal discovery from mixed and missing data with applications on ADHD datasets
 - 43 Maaïke de Boer (RUN), Semantic Mapping in Video Retrieval
 - 44 Garm Lucassen (UU), Understanding User Stories - Computational Linguistics in Agile Requirements Engineering
 - 45 Bas Testerink (UU), Decentralized Runtime Norm Enforcement
 - 46 Jan Schneider (OU), Sensor-based Learning Support
 - 47 Jie Yang (TUD), Crowd Knowledge Creation Acceleration
 - 48 Angel Suarez (OU), Collaborative inquiry-based learning
-
- 2018 01 Han van der Aa (VUA), Comparing and Aligning Process Representations
 - 02 Felix Mannhardt (TU/e), Multi-perspective Process Mining

- 03 Steven Bosems (UT), Causal Models For Well-Being: Knowledge Modeling, Model-Driven Development of Context-Aware Applications, and Behavior Prediction
- 04 Jordan Janeiro (TUD), Flexible Coordination Support for Diagnosis Teams in Data-Centric Engineering Tasks
- 05 Hugo Huurdeman (UvA), Supporting the Complex Dynamics of the Information Seeking Process
- 06 Dan Ionita (UT), Model-Driven Information Security Risk Assessment of Socio-Technical Systems
- 07 Jieting Luo (UU), A formal account of opportunism in multi-agent systems
- 08 Rick Smetsers (RUN), Advances in Model Learning for Software Systems
- 09 Xu Xie (TUD), Data Assimilation in Discrete Event Simulations
- 10 Julienka Mollee (VUA), Moving forward: supporting physical activity behavior change through intelligent technology
- 11 Mahdi Sargolzaei (UvA), Enabling Framework for Service-oriented Collaborative Networks
- 12 Xixi Lu (TU/e), Using behavioral context in process mining
- 13 Seyed Amin Tabatabaei (VUA), Computing a Sustainable Future
- 14 Bart Joosten (TiU), Detecting Social Signals with Spatiotemporal Gabor Filters
- 15 Naser Davarzani (UM), Biomarker discovery in heart failure
- 16 Jaebok Kim (UT), Automatic recognition of engagement and emotion in a group of children
- 17 Jianpeng Zhang (TU/e), On Graph Sample Clustering
- 18 Henriette Nakad (UL), De Notaris en Private Rechtspraak
- 19 Minh Duc Pham (VUA), Emergent relational schemas for RDF
- 20 Manxia Liu (RUN), Time and Bayesian Networks
- 21 Aad Slootmaker (OU), EMERGO: a generic platform for authoring and playing scenario-based serious games
- 22 Eric Fernandes de Mello Araújo (VUA), Contagious: Modeling the Spread of Behaviours, Perceptions and Emotions in Social Networks
- 23 Kim Schouten (EUR), Semantics-driven Aspect-Based Sentiment Analysis
- 24 Jered Vroon (UT), Responsive Social Positioning Behaviour for Semi-Autonomous Telepresence Robots
- 25 Riste Gligorov (VUA), Serious Games in Audio-Visual Collections
- 26 Roelof Anne Jelle de Vries (UT), Theory-Based and Tailor-Made: Motivational Messages for Behavior Change Technology

- 27 Maikel Leemans (TU/e), Hierarchical Process Mining for Scalable Software Analysis
 - 28 Christian Willemse (UT), Social Touch Technologies: How they feel and how they make you feel
 - 29 Yu Gu (TiU), Emotion Recognition from Mandarin Speech
 - 30 Wouter Beek (VUA), The "K" in "semantic web" stands for "knowledge": scaling semantics to the web
-
- 2019 01 Rob van Eijk (UL), Web privacy measurement in real-time bidding systems. A graph-based approach to RTB system classification
 - 02 Emmanuelle Beauxis Aussalet (CWI, UU), Statistics and Visualizations for Assessing Class Size Uncertainty
 - 03 Eduardo Gonzalez Lopez de Murillas (TU/e), Process Mining on Databases: Extracting Event Data from Real Life Data Sources
 - 04 Ridho Rahmadi (RUN), Finding stable causal structures from clinical data
 - 05 Sebastiaan van Zelst (TU/e), Process Mining with Streaming Data
 - 06 Chris Dijkshoorn (VUA), Nichesourcing for Improving Access to Linked Cultural Heritage Datasets
 - 07 Soude Fazeli (TUD), Recommender Systems in Social Learning Platforms
 - 08 Frits de Nijs (TUD), Resource-constrained Multi-agent Markov Decision Processes
 - 09 Fahimeh Alizadeh Moghaddam (UvA), Self-adaptation for energy efficiency in software systems
 - 10 Qing Chuan Ye (EUR), Multi-objective Optimization Methods for Allocation and Prediction
 - 11 Yue Zhao (TUD), Learning Analytics Technology to Understand Learner Behavioral Engagement in MOOCs
 - 12 Jacqueline Heinerman (VUA), Better Together
 - 13 Guanliang Chen (TUD), MOOC Analytics: Learner Modeling and Content Generation
 - 14 Daniel Davis (TUD), Large-Scale Learning Analytics: Modeling Learner Behavior & Improving Learning Outcomes in Massive Open Online Courses
 - 15 Erwin Walraven (TUD), Planning under Uncertainty in Constrained and Partially Observable Environments
 - 16 Guangming Li (TU/e), Process Mining based on Object-Centric Behavioral Constraint (OCBC) Models

- 17 Ali Hurriyetoglu (RUN), Extracting actionable information from microtexts
- 18 Gerard Wagenaar (UU), Artefacts in Agile Team Communication
- 19 Vincent Koeman (TUD), Tools for Developing Cognitive Agents
- 20 Chide Groenouwe (UU), Fostering technically augmented human collective intelligence
- 21 Cong Liu (TU/e), Software Data Analytics: Architectural Model Discovery and Design Pattern Detection
- 22 Martin van den Berg (VUA), Improving IT Decisions with Enterprise Architecture
- 23 Qin Liu (TUD), Intelligent Control Systems: Learning, Interpreting, Verification
- 24 Anca Dumitrache (VUA), Truth in Disagreement - Crowdsourcing Labeled Data for Natural Language Processing
- 25 Emiel van Miltenburg (VUA), Pragmatic factors in (automatic) image description
- 26 Prince Singh (UT), An Integration Platform for Synchromodal Transport
- 27 Alessandra Antonaci (OU), The Gamification Design Process applied to (Massive) Open Online Courses
- 28 Esther Kuindersma (UL), Cleared for take-off: Game-based learning to prepare airline pilots for critical situations
- 29 Daniel Formolo (VUA), Using virtual agents for simulation and training of social skills in safety-critical circumstances
- 30 Vahid Yazdanpanah (UT), Multiagent Industrial Symbiosis Systems
- 31 Milan Jelisavcic (VUA), Alive and Kicking: Baby Steps in Robotics
- 32 Chiara Sironi (UM), Monte-Carlo Tree Search for Artificial General Intelligence in Games
- 33 Anil Yaman (TU/e), Evolution of Biologically Inspired Learning in Artificial Neural Networks
- 34 Negar Ahmadi (TU/e), EEG Microstate and Functional Brain Network Features for Classification of Epilepsy and PNES
- 35 Lisa Facey-Shaw (OU), Gamification with digital badges in learning programming
- 36 Kevin Ackermans (OU), Designing Video-Enhanced Rubrics to Master Complex Skills
- 37 Jian Fang (TUD), Database Acceleration on FPGAs

- 38 Akos Kadar (OU), Learning visually grounded and multilingual representations
-
- 2020 01 Armon Toubman (UL), Calculated Moves: Generating Air Combat Behaviour
 - 02 Marcos de Paula Bueno (UL), Unraveling Temporal Processes using Probabilistic Graphical Models
 - 03 Mostafa Deghani (UvA), Learning with Imperfect Supervision for Language Understanding
 - 04 Maarten van Gompel (RUN), Context as Linguistic Bridges
 - 05 Yulong Pei (TU/e), On local and global structure mining
 - 06 Preethu Rose Anish (UT), Stimulation Architectural Thinking during Requirements Elicitation - An Approach and Tool Support
 - 07 Wim van der Vegt (OU), Towards a software architecture for reusable game components
 - 08 Ali Mirsoleimani (UL), Structured Parallel Programming for Monte Carlo Tree Search
 - 09 Myriam Traub (UU), Measuring Tool Bias and Improving Data Quality for Digital Humanities Research
 - 10 Alifah Syamsiyah (TU/e), In-database Preprocessing for Process Mining
 - 11 Sepideh Mesbah (TUD), Semantic-Enhanced Training Data Augmentation Methods for Long-Tail Entity Recognition Models
 - 12 Ward van Breda (VUA), Predictive Modeling in E-Mental Health: Exploring Applicability in Personalised Depression Treatment
 - 13 Marco Virgolin (CWI), Design and Application of Gene-pool Optimal Mixing Evolutionary Algorithms for Genetic Programming
 - 14 Mark Raasveldt (CWI/UL), Integrating Analytics with Relational Databases
 - 15 Konstantinos Georgiadis (OU), Smart CAT: Machine Learning for Configurable Assessments in Serious Games
 - 16 Ilona Wilmont (RUN), Cognitive Aspects of Conceptual Modelling
 - 17 Daniele Di Mitri (OU), The Multimodal Tutor: Adaptive Feedback from Multimodal Experiences
 - 18 Georgios Methenitis (TUD), Agent Interactions & Mechanisms in Markets with Uncertainties: Electricity Markets in Renewable Energy Systems
 - 19 Guido van Capelleveen (UT), Industrial Symbiosis Recommender Systems
 - 20 Albert Hankel (VUA), Embedding Green ICT Maturity in Organisations

- 21 Karine da Silva Miras de Araujo (VUA), Where is the robot?: Life as it could be
 - 22 Maryam Masoud Khamis (RUN), Understanding complex systems implementation through a modeling approach: the case of e-government in Zanzibar
 - 23 Rianne Conijn (UT), The Keys to Writing: A writing analytics approach to studying writing processes using keystroke logging
 - 24 Lenin da Nóbrega Medeiros (VUA/RUN), How are you feeling, human? Towards emotionally supportive chatbots
 - 25 Xin Du (TU/e), The Uncertainty in Exceptional Model Mining
 - 26 Krzysztof Leszek Sadowski (UU), GAMBIT: Genetic Algorithm for Model-Based mixed-Integer opTimization
 - 27 Ekaterina Muravyeva (TUD), Personal data and informed consent in an educational context
 - 28 Bibeg Limbu (TUD), Multimodal interaction for deliberate practice: Training complex skills with augmented reality
 - 29 Ioan Gabriel Bucur (RUN), Being Bayesian about Causal Inference
 - 30 Bob Zadok Blok (UL), Creatief, Creatiever, Creatiefst
 - 31 Gongjin Lan (VUA), Learning better – From Baby to Better
 - 32 Jason Rhuggenaath (TU/e), Revenue management in online markets: pricing and online advertising
 - 33 Rick Gilsing (TU/e), Supporting service-dominant business model evaluation in the context of business model innovation
 - 34 Anna Bon (UM), Intervention or Collaboration? Redesigning Information and Communication Technologies for Development
 - 35 Siamak Farshidi (UU), Multi-Criteria Decision-Making in Software Production
-
- 2021 01 Francisco Xavier Dos Santos Fonseca (TUD), Location-based Games for Social Interaction in Public Space
 - 02 Rijk Mercur (TUD), Simulating Human Routines: Integrating Social Practice Theory in Agent-Based Models
 - 03 Seyyed Hadi Hashemi (UvA), Modeling Users Interacting with Smart Devices
 - 04 Ioana Jivet (OU), The Dashboard That Loved Me: Designing adaptive learning analytics for self-regulated learning
 - 05 Davide Dell'Anna (UU), Data-Driven Supervision of Autonomous Systems
 - 06 Daniel Davison (UT), "Hey robot, what do you think?" How children learn with a social robot

- 07 Armel Lefebvre (UU), Research data management for open science
- 08 Nardie Fanchamps (OU), The Influence of Sense-Reason-Act Programming on Computational Thinking
- 09 Cristina Zaga (UT), The Design of Robothings. Non-Anthropomorphic and Non-Verbal Robots to Promote Children's Collaboration Through Play
- 10 Quinten Meertens (UvA), Misclassification Bias in Statistical Learning
- 11 Anne van Rossum (UL), Nonparametric Bayesian Methods in Robotic Vision
- 12 Lei Pi (UL), External Knowledge Absorption in Chinese SMEs
- 13 Bob R. Schadenberg (UT), Robots for Autistic Children: Understanding and Facilitating Predictability for Engagement in Learning
- 14 Negin Samaeemofrad (UL), Business Incubators: The Impact of Their Support
- 15 Onat Ege Adali (TU/e), Transformation of Value Propositions into Resource Re-Configurations through the Business Services Paradigm
- 16 Esam A. H. Ghaleb (UM), Bimodal emotion recognition from audio-visual cues
- 17 Dario Dotti (UM), Human Behavior Understanding from motion and bodily cues using deep neural networks
- 18 Remi Wieten (UU), Bridging the Gap Between Informal Sense-Making Tools and Formal Systems - Facilitating the Construction of Bayesian Networks and Argumentation Frameworks
- 19 Roberto Verdecchia (VUA), Architectural Technical Debt: Identification and Management
- 20 Masoud Mansoury (TU/e), Understanding and Mitigating Multi-Sided Exposure Bias in Recommender Systems
- 21 Pedro Thiago Timbó Holanda (CWI), Progressive Indexes
- 22 Sihang Qiu (TUD), Conversational Crowdsourcing
- 23 Hugo Manuel Proença (UL), Robust rules for prediction and description
- 24 Kaijie Zhu (TU/e), On Efficient Temporal Subgraph Query Processing
- 25 Eoin Martino Grua (VUA), The Future of E-Health is Mobile: Combining AI and Self-Adaptation to Create Adaptive E-Health Mobile Applications
- 26 Benno Kruit (CWI/VUA), Reading the Grid: Extending Knowledge Bases from Human-readable Tables

- 27 Jelte van Waterschoot (UT), Personalized and Personal Conversations: Designing Agents Who Want to Connect With You
 - 28 Christoph Selig (UL), Understanding the Heterogeneity of Corporate Entrepreneurship Programs
-
- 2022 01 Judith van Stegeren (UT), Flavor text generation for role-playing video games
 - 02 Paulo da Costa (TU/e), Data-driven Prognostics and Logistics Optimisation: A Deep Learning Journey
 - 03 Ali el Hassouni (VUA), A Model A Day Keeps The Doctor Away: Reinforcement Learning For Personalized Healthcare
 - 04 Ünal Aksu (UU), A Cross-Organizational Process Mining Framework
 - 05 Shiwei Liu (TU/e), Sparse Neural Network Training with In-Time Over-Parameterization
 - 06 Reza Refaei Afshar (TU/e), Machine Learning for Ad Publishers in Real Time Bidding
 - 07 Sambit Praharaj (OU), Measuring the Unmeasurable? Towards Automatic Co-located Collaboration Analytics
 - 08 Maikel L. van Eck (TU/e), Process Mining for Smart Product Design
 - 09 Oana Andreea Inel (VUA), Understanding Events: A Diversity-driven Human-Machine Approach
 - 10 Felipe Moraes Gomes (TUD), Examining the Effectiveness of Collaborative Search Engines
 - 11 Mirjam de Haas (UT), Staying engaged in child-robot interaction, a quantitative approach to studying preschoolers' engagement with robots and tasks during second-language tutoring
 - 12 Guanyi Chen (UU), Computational Generation of Chinese Noun Phrases
 - 13 Xander Wilcke (VUA), Machine Learning on Multimodal Knowledge Graphs: Opportunities, Challenges, and Methods for Learning on Real-World Heterogeneous and Spatially-Oriented Knowledge
 - 14 Michiel Overeem (UU), Evolution of Low-Code Platforms
 - 15 Jelmer Jan Koorn (UU), Work in Process: Unearthing Meaning using Process Mining
 - 16 Pieter Gijbbers (TU/e), Systems for AutoML Research
 - 17 Laura van der Lubbe (VUA), Empowering vulnerable people with serious games and gamification

- 18 Paris Mavromoustakos Blom (TiU), Player Affect Modelling and Video Game Personalisation
 - 19 Bilge Yigit Ozkan (UU), Cybersecurity Maturity Assessment and Standardisation
 - 20 Fakhra Jabeen (VUA), Dark Side of the Digital Media - Computational Analysis of Negative Human Behaviors on Social Media
 - 21 Seethu Mariyam Christopher (UM), Intelligent Toys for Physical and Cognitive Assessments
 - 22 Alexandra Sierra Rativa (TiU), Virtual Character Design and its potential to foster Empathy, Immersion, and Collaboration Skills in Video Games and Virtual Reality Simulations
 - 23 Ilir Kola (TUD), Enabling Social Situation Awareness in Support Agents
 - 24 Samaneh Heidari (UU), Agents with Social Norms and Values - A framework for agent based social simulations with social norms and personal values
 - 25 Anna L.D. Latour (UL), Optimal decision-making under constraints and uncertainty
 - 26 Anne Dirkson (UL), Knowledge Discovery from Patient Forums: Gaining novel medical insights from patient experiences
 - 27 Christos Athanasiadis (UM), Emotion-aware cross-modal domain adaptation in video sequences
 - 28 Onuralp Ulusoy (UU), Privacy in Collaborative Systems
 - 29 Jan Kolkmeier (UT), From Head Transform to Mind Transplant: Social Interactions in Mixed Reality
 - 30 Dean De Leo (CWI), Analysis of Dynamic Graphs on Sparse Arrays
 - 31 Konstantinos Traganos (TU/e), Tackling Complexity in Smart Manufacturing with Advanced Manufacturing Process Management
 - 32 Cezara Pastrav (UU), Social simulation for socio-ecological systems
 - 33 Brinn Hekkelman (CWI/TUD), Fair Mechanisms for Smart Grid Congestion Management
 - 34 Nimat Ullah (VUA), Mind Your Behaviour: Computational Modelling of Emotion & Desire Regulation for Behaviour Change
 - 35 Mike E.U. Ligthart (VUA), Shaping the Child-Robot Relationship: Interaction Design Patterns for a Sustainable Interaction
-
- 2023 01 Bojan Simoski (VUA), Untangling the Puzzle of Digital Health Interventions

- 02 Mariana Rachel Dias da Silva (TiU), Grounded or in flight? What our bodies can tell us about the whereabouts of our thoughts
- 03 Shabnam Najafian (TUD), User Modeling for Privacy-preserving Explanations in Group Recommendations
- 04 Gineke Wiggers (UL), The Relevance of Impact: bibliometric-enhanced legal information retrieval
- 05 Anton Bouter (CWI), Optimal Mixing Evolutionary Algorithms for Large-Scale Real-Valued Optimization, Including Real-World Medical Applications
- 06 António Pereira Barata (UL), Reliable and Fair Machine Learning for Risk Assessment
- 07 Tianjin Huang (TU/e), The Roles of Adversarial Examples on Trustworthiness of Deep Learning
- 08 Lu Yin (TU/e), Knowledge Elicitation using Psychometric Learning
- 09 Xu Wang (VUA), Scientific Dataset Recommendation with Semantic Techniques
- 10 Dennis J.N.J. Soemers (UM), Learning State-Action Features for General Game Playing
- 11 Fawad Taj (VUA), Towards Motivating Machines: Computational Modeling of the Mechanism of Actions for Effective Digital Health Behavior Change Applications
- 12 Tessel Bogaard (VUA), Using Metadata to Understand Search Behavior in Digital Libraries
- 13 Injy Sarhan (UU), Open Information Extraction for Knowledge Representation
- 14 Selma Čaušević (TUD), Energy resilience through self-organization
- 15 Alvaro Henrique Chaim Correia (TU/e), Insights on Learning Tractable Probabilistic Graphical Models
- 16 Peter Blomsma (TiU), Building Embodied Conversational Agents: Observations on human nonverbal behaviour as a resource for the development of artificial characters
- 17 Meike Nauta (UT), Explainable AI and Interpretable Computer Vision – From Oversight to Insight
- 18 Gustavo Penha (TUD), Designing and Diagnosing Models for Conversational Search and Recommendation
- 19 George Aalbers (TiU), Digital Traces of the Mind: Using Smartphones to Capture Signals of Well-Being in Individuals

- 20 Arkadiy Dushatskiy (TUD), Expensive Optimization with Model-Based Evolutionary Algorithms applied to Medical Image Segmentation using Deep Learning
 - 21 Gerrit Jan de Bruin (UL), Network Analysis Methods for Smart Inspection in the Transport Domain
 - 22 Alireza Shojaifar (UU), Volitional Cybersecurity
 - 23 Theo Theunissen (UU), Documentation in Continuous Software Development
 - 24 Agathe Balayn (TUD), Practices Towards Hazardous Failure Diagnosis in Machine Learning
 - 25 Jurian Baas (UU), Entity Resolution on Historical Knowledge Graphs
 - 26 Loek Tonnaer (TU/e), Linearly Symmetry-Based Disentangled Representations and their Out-of-Distribution Behaviour
 - 27 Ghada Sokar (TU/e), Learning Continually Under Changing Data Distributions
 - 28 Floris den Hengst (VUA), Learning to Behave: Reinforcement Learning in Human Contexts
 - 29 Tim Draws (TUD), Understanding Viewpoint Biases in Web Search Results
-
- 2024 01 Daphne Miedema (TU/e), On Learning SQL: Disentangling concepts in data systems education
 - 02 Emile van Krieken (VUA), Optimisation in Neurosymbolic Learning Systems
 - 03 Feri Wijayanto (RUN), Automated Model Selection for Rasch and Mediation Analysis
 - 04 Mike Huisman (UL), Understanding Deep Meta-Learning
 - 05 Yiyong Gou (UM), Aerial Robotic Operations: Multi-environment Cooperative Inspection & Construction Crack Autonomous Repair
 - 06 Azqa Nadeem (TUD), Understanding Adversary Behavior via XAI: Leveraging Sequence Clustering to Extract Threat Intelligence
 - 07 Parisa Shayan (TiU), Modeling User Behavior in Learning Management Systems
 - 08 Xin Zhou (UvA), From Empowering to Motivating: Enhancing Policy Enforcement through Process Design and Incentive Implementation
 - 09 Giso Dal (UT), Probabilistic Inference Using Partitioned Bayesian Networks

- 10 Cristina-Iulia Bucur (VUA), Linkflows: Towards Genuine Semantic Publishing in Science
- 11 withdrawn
- 12 Peide Zhu (TUD), Towards Robust Automatic Question Generation For Learning
- 13 Enrico Liscio (TUD), Context-Specific Value Inference via Hybrid Intelligence
- 14 Larissa Capobianco Shimomura (TU/e), On Graph Generating Dependencies and their Applications in Data Profiling
- 15 Ting Liu (VUA), A Gut Feeling: Biomedical Knowledge Graphs for Interrelating the Gut Microbiome and Mental Health
- 16 Arthur Barbosa Câmara (TUD), Designing Search-as-Learning Systems
- 17 Razieh Alidoosti (VUA), Ethics-aware Software Architecture Design
- 18 Laurens Stoop (UU), Data Driven Understanding of Energy-Meteorological Variability and its Impact on Energy System Operations
- 19 Azadeh Mozafari Mehr (TU/e), Multi-perspective Conformance Checking: Identifying and Understanding Patterns of Anomalous Behavior
- 20 Ritsart Anne Plantenga (UL), Omgang met Regels
- 21 Federica Vinella (UU), Crowdsourcing User-Centered Teams
- 22 Zeynep Ozturk Yurt (TU/e), Beyond Routine: Extending BPM for Knowledge-Intensive Processes with Controllable Dynamic Contexts
- 23 Jie Luo (VUA), Lamarck's Revenge: Inheritance of Learned Traits Improves Robot Evolution
- 24 Nirmal Roy (TUD), Exploring the effects of interactive interfaces on user search behaviour
- 25 Alisa Rieger (TUD), Striving for Responsible Opinion Formation in Web Search on Debated Topics
- 26 Tim Gubner (CWI), Adaptively Generating Heterogeneous Execution Strategies using the VOILA Framework
- 27 Lincen Yang (UL), Information-theoretic Partition-based Models for Interpretable Machine Learning
- 28 Leon Helwerda (UL), Grip on Software: Understanding development progress of Scrum sprints and backlogs
- 29 David Wilson Romero Guzman (VUA), The Good, the Efficient and the Inductive Biases: Exploring Efficiency in Deep Learning Through the Use of Inductive Biases

- 30 Vijanti Ramautar (UU), Model-Driven Sustainability Accounting
 - 31 Ziyu Li (TUD), On the Utility of Metadata to Optimize Machine Learning Workflows
 - 32 Vinicius Stein Dani (UU), The Alpha and Omega of Process Mining
 - 33 Siddharth Mehrotra (TUD), Designing for Appropriate Trust in Human-AI interaction
 - 34 Robert Deckers (VUA), From Smallest Software Particle to System Specification - MuDForM: Multi-Domain Formalization Method
 - 35 Sicui Zhang (TU/e), Methods of Detecting Clinical Deviations with Process Mining: a fuzzy set approach
 - 36 Thomas Mulder (TU/e), Optimization of Recursive Queries on Graphs
 - 37 James Graham Nevin (UvA), The Ramifications of Data Handling for Computational Models
 - 38 Christos Koutras (TUD), Tabular Schema Matching for Modern Settings
 - 39 Paola Lara Machado (TU/e), The Nexus between Business Models and Operating Models: From Conceptual Understanding to Actionable Guidance
 - 40 Montijn van de Ven (TU/e), Guiding the Definition of Key Performance Indicators for Business Models
 - 41 Georgios Siachamis (TUD), Adaptivity for Streaming Dataflow Engines
 - 42 Emmeke Veltmeijer (VUA), Small Groups, Big Insights: Understanding the Crowd through Expressive Subgroup Analysis
 - 43 Cedric Waterschoot (KNAW Meertens Instituut), The Constructive Conundrum: Computational Approaches to Facilitate Constructive Commenting on Online News Platforms
 - 44 Marcel Schmitz (OU), Towards learning analytics-supported learning design
 - 45 Sara Salimzadeh (TUD), Living in the Age of AI: Understanding Contextual Factors that Shape Human-AI Decision-Making
 - 46 Georgios Stathis (Leiden University), Preventing Disputes: Preventive Logic, Law & Technology
 - 47 Daniel Daza (VUA), Exploiting Subgraphs and Attributes for Representation Learning on Knowledge Graphs
 - 48 Ioannis Petros Samiotis (TUD), Crowd-Assisted Annotation of Classical Music Compositions
-

- 2025 01 Max van Haastrecht (UL), Transdisciplinary Perspectives on Validity: Bridging the Gap Between Design and Implementation for Technology-Enhanced Learning Systems
- 02 Jurgen van den Hoogen (JADS), Time Series Analysis Using Convolutional Neural Networks
- 03 Andra-Denis Ionescu (TUD), Feature Discovery for Data-Centric AI
- 04 Rianne Schouten (TU/e), Exceptional Model Mining for Hierarchical Data
- 05 Nele Albers (TUD), Psychology-Informed Reinforcement Learning for Situated Virtual Coaching in Smoking Cessation
- 06 Daniël Vos (TUD), Decision Tree Learning: Algorithms for Robust Prediction and Policy Optimization
- 07 Ricky Maulana Fajri (TU/e), Towards Safer Active Learning: Dealing with Unwanted Biases, Graph-Structured Data, Adversary, and Data Imbalance
- 08 Stefan Bloemheuvel (TiU), Spatio-Temporal Analysis Through Graphs: Predictive Modeling and Graph Construction
- 09 Fadime Kaya (VUA), Decentralized Governance Design - A Model-Based Approach
- 10 Zhao Yang (UL), Enhancing Autonomy and Efficiency in Goal-Conditioned Reinforcement Learning
- 11 Shahin Sharifi Noorian (TUD), From Recognition to Understanding: Enriching Visual Models Through Multi-Modal Semantic Integration
- 12 Lijun Lyu (TUD), Interpretability in Neural Information Retrieval
- 13 Fuda van Diggelen (VUA), Robots Need Some Education: on the complexity of learning in evolutionary robotics
- 14 Gennaro Gala (TU/e), Probabilistic Generative Modeling with Latent Variable Hierarchies
- 15 Michiel van der Meer (UL), Opinion Diversity through Hybrid Intelligence
- 16 Monika Grewal (TU Delft), Deep Learning for Landmark Detection, Segmentation, and Multi-Objective Deformable Registration in Medical Imaging
- 17 Matteo De Carlo (VUA), Real Robot Reproduction: Towards Evolving Robotic Ecosystems
- 18 Anouk Neerincx (UU), Robots That Care: How Social Robots Can Boost Children's Mental Wellbeing
- 19 Fang Hou (UU), Trust in Software Ecosystems

- 20 Alexander Melchior (UU), Modelling for Policy is More Than Policy Modelling (The Useful Application of Agent-Based Modelling in Complex Policy Processes)
- 21 Mandani Ntekeuli (UM), Bridging Individual and Group Perspectives in Psychopathology: Computational Modeling Approaches using Ecological Momentary Assessment Data
- 22 Hilde Weerts (TU/e), Decoding Algorithmic Fairness: Towards Interdisciplinary Understanding of Fairness and Discrimination in Algorithmic Decision-Making
- 23 Roderick van der Weerd (VUA), IoT Measurement Knowledge Graphs: Constructing, Working and Learning with IoT Measurement Data as a Knowledge Graph
- 24 Zhong Li (UL), Trustworthy Anomaly Detection for Smart Manufacturing
- 25 Kyana van Eijndhoven (TiU), A Breakdown of Breakdowns: Multi-Level Team Coordination Dynamics under Stressful Conditions
- 26 Tom Pepels (UM), Monte-Carlo Tree Search is Work in Progress
- 27 Danil Provodin (JADS, TU/e), Sequential Decision Making Under Complex Feedback
- 28 Jinke He (TU Delft), Exploring Learned Abstract Models for Efficient Planning and Learning
- 29 Erik van Haeringen (VUA), Mixed Feelings: Simulating Emotion Contagion in Groups
- 30 Myrthe Reuver (VUA), A Puzzle of Perspectives: Interdisciplinary Language Technology for Responsible News Recommendation
- 31 Gebrekirstos Gebreselassie Gebremeskel (RUN), Spotlight on Recommender Systems: Contributions to Selected Components in the Recommendation Pipeline
- 32 Ryan Brate (UU), Words Matter: A Computational Toolkit for Charged Terms
- 33 Merle Reimann (VUA), Speaking the Same Language: Spoken Capability Communication in Human-Agent and Human-Robot Interaction
- 34 Eduard C. Groen (UU), Crowd-Based Requirements Engineering
- 35 Urja Khurana (VUA), From Concept To Impact: Toward More Robust Language Model Deployment
- 36 Anna Maria Wegmann (UU), Say the Same but Differently: Computational Approaches to Stylistic Variation and Paraphrasing

