



Approximate Definitional Constructs as Lightweight Evidence for Detecting Classes Among Wikipedia Articles

Marius Paşca

Google

Mountain View, California

mars@google.com

ABSTRACT

A lightweight method uses a few extraction patterns in the task of distinguishing Wikipedia articles that are classes (“*Walled garden*”, “*Garden*”) from other articles (“*High Hazels Park*”). The method acquires a set of classes, based on patterns targeting phrases that likely refer to either concepts being introduced or defined (“*a walled garden is a garden [..]*”); or to concepts used to introduce or define other concepts (“*a walled garden is a garden [..]*”). Experimental results over multiple evaluation sets are better, when relying on defined phrases alone vs. defining phrases alone; and further improved, when combining complementary evidence from both.

CCS CONCEPTS

- **Information systems** → Content analysis and feature selection;
- **Computing methodologies** → Information extraction.

KEYWORDS

Knowledge acquisition, open-domain information extraction, topic classification, classes, semantics

ACM Reference Format:

Marius Paşca. 2019. Approximate Definitional Constructs as Lightweight Evidence for Detecting Classes Among Wikipedia Articles. In *The 28th ACM International Conference on Information and Knowledge Management (CIKM '19), November 3–7, 2019, Beijing, China*. ACM, New York, NY, USA, 4 pages. <https://doi.org/10.1145/3357384.3358167>

1 INTRODUCTION

Motivation: Wikipedia and knowledge repositories derived from it are useful in a variety of tasks. They pertain to knowledge acquisition from text [12, 14, 18, 30–32], text analysis [22, 23] and information retrieval [3, 5, 6, 15, 16, 24, 28, 33] including commercial Web search, helping to potentially transform search results from sets of hyperlinks to relevant documents into sets of concepts directly relevant to users’ queries [26]. Most Wikipedia articles correspond to concepts that are instances (“*Wynnewood Valley Park Sensory Garden*”) as opposed to classes (“*Garden*”). This reflects in part the encyclopedic nature of Wikipedia and in part the expectation that instances naturally dominate classes in sheer count. But

multiple language editions of Wikipedia contain millions of articles each, making the subset of Wikipedia articles that are likely classes significant in size, both in absolute and relative even to resources whose focus is specifically not on instances but classes, such as WordNet [9]. Distinguishing articles that are classes would benefit Wikipedia and knowledge repositories derived from it.

Contributions: The method proposed in this paper relies on the text of Wikipedia articles, in order to automatically detect a subset of articles within Wikipedia that are classes. For this purpose, the method identifies text fragments likely containing noun phrases referring to countable concepts that are either being defined (“*a walled garden is a garden [..]*”), similarly to [20]; or used in defining other concepts (“*a walled garden is a garden [..]*”), unlike [20]. The phrases are disambiguated to their corresponding Wikipedia articles (titled “*Walled garden*”, “*Garden*”) based on link data readily available within Wikipedia. The method requires no linguistic preprocessing tools such as part of speech taggers, named entity recognizers, syntactic or semantic parsers or any other, making it simpler to port to other languages. Evaluation over multiple evaluation sets gives encouraging results, when exploiting both defined phrases as well as defining phrases.

2 DETECTION OF CLASSES

Task Definition: The task being addressed is the acquisition (i.e., selection) of a subset of Wikipedia articles that are classes. A class is a placeholder for a set of instances that share common properties. The acquisition is equivalent to attaching annotations to Wikipedia articles, whenever they are classes.

Detecting Definitional Constructs: The decision whether a Wikipedia article is a class relies on two types of (case insensitive) evidence available in the article, which apply to English and to many other languages, including Romance and Germanic languages:

Defined Phrases (*Dfd*): If a sentence from the article matches one of the following language-specific patterns, then the fragment *H* from the sentence may be a countable phrase being defined:

(*En*): [a|an] *H* [was|is] (*Fr*): [un|une] *H* [était|est]

(*Es*): [un|una] *H* [fue|es]

Defining Phrases (*Dng*): If a sentence from the article matches one of the following language-specific patterns, then the fragment *H* from the sentence may be a phrase being used to define another:

(*En*): [was|is] [a|an] *H* (*Fr*): [était|est] [un|une] *H*

(*Es*): [fue|es] [un|una] *H*

Locating Phrase Boundaries: The patterns that target defined phrases (*Dfd*) encode sufficient information to determine the left and right boundaries of the fragments *H*. That is not the case, however, for patterns that extract defining phrases (*Dng*): the right-side boundary of the fragments *H* remains to be determined.

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

CIKM '19, November 3–7, 2019, Beijing, China

© 2019 Copyright held by the owner/author(s).

ACM ISBN 978-1-4503-6976-3/19/11.

<https://doi.org/10.1145/3357384.3358167>

For simplicity, as an alternative to using other text analysis tools, the right-side boundaries of fragments H in defining phrase patterns are determined in two steps. First, H is further required to be identical to the anchor text of any of the outgoing internal links from the article. An outgoing internal link is a hyperlink from the article to another Wikipedia article. If multiple such anchor texts exist, the longest is retained. Second, H is further required to be such that its last token is immediately followed by a separating token. A separating token is approximated to be either one of a set of known prepositions in the respective language (e.g., “about”, “from”, “in” etc., for English); or one of a small set of relative pronouns (“whose”, “which”, “that”, for English); or a potential sentence terminator (i.e., a dot). If the requirements in both steps are satisfied, the right-side boundary of the defining phrase H has been identified, as being the end of the anchor text. Otherwise, the pattern match is ignored and the fragment H discarded.

If a sentence matches a pattern and a fragment H is successfully extracted, the fragment H constitutes a defined phrase (if extracted with Dfd) or defining phrase (if extracted with Dng).

To illustrate, consider the sentence “A Philosophical Garden is a garden whose design reflects [...]” from the Wikipedia article “Philosophical garden”. The sentence matches a Dfd pattern and, more interestingly, a Dng pattern. For the latter, the fragment H that starts immediately after “is a” in the sentence a) is identical to the anchor text “garden” of an outgoing internal link from the article “Philosophical garden” to the article “Garden”; and b) is immediately followed by a separating token, namely the relative pronoun “whose”. Therefore, the fragment “garden” is extracted from the sentence as a defining phrase. In another example, the sentence “We Will Follow: A Tribute to U2 is a U2 tribute album recorded by [...]” from the Wikipedia article “We Will Follow: A Tribute to U2” also initially matches a Dng pattern. The longest fragment H immediately after “is a” in the sentence that is identical to the anchor text of any outgoing internal link is “U2”. The token immediately following it in the sentence is “tribute”. Since the token is not a separating token, the initial pattern match and the fragment “U2” are discarded.

The proposed patterns are inspired by the popular lexical patterns introduced in [2, 13] and still in widespread use in extracting IsA relations from text [10, 25, 27, 31]. In comparison to [13], where patterns target pairs of a hyponym and hypernym as done in [13], the patterns here target only one side of such possible pairs, regardless of whether the other side could or could not be extracted from the text surrounding the patterns in the articles. The set of patterns is not meant to be exhaustive in extracting classes.

Mapping Phrases to Articles: Defined phrases as well as defining phrases are mapped to their corresponding Wikipedia articles, effectively disambiguating them. The procedure is different, depending on whether the phrases are defined or defining phrases.

Defined phrases are mapped to the article from which they are extracted, if the phrases are identical to the article title after normalization to lowercase. The defined phrase “Philosophical Garden”, extracted from the sentence “A Philosophical Garden is a garden whose design reflects [...]”, is mapped to the Wikipedia article containing the sentence, namely “Philosophical garden”.

Defining phrases are mapped to the article to which the outgoing internal link, whose anchor text gave the longest match against

the sentence, points. In the same example sentence “A Philosophical Garden is a garden whose design reflects [...]”, the defining phrase “garden” is matched against the hyperlink whose anchor text was matched against the phrase, namely the article “Garden”.

Acquisition from Wikipedia: Through mapping to their corresponding articles, defined phrases and defining phrases are effectively converted (disambiguated) into defined and defining articles. Those articles are extracted as classes, based on defined (Dfd) or defining (Dng) constructs as evidence. The set of articles extracted as classes is the union of the sets of classes extracted by any of the individual types of evidence, namely defined vs. defining phrases.

3 EXPERIMENTAL SETTING

Data Source: The experiments rely on a snapshot of Wikipedia from which disambiguation or redirect pages are discarded.

Evaluation Sets: Three evaluation sets introduced in [20] serve as the source data for testing the proposed method. Each evaluation set consists in pairs of a Wikipedia article in English and a gold label indicating whether the article is a class or not. The evaluation sets are derived automatically from WordNet [9] (S_W); or through manual annotation of samples of Wikipedia articles (S_D , S_Q). They contain 5,735 (S_W), 2,000 (S_D) and 1,000 (S_Q) entries, divided into 547 and 5,188 (S_W), 73 and 1,927 (S_D) and 362 and 628 (S_Q) gold classes and gold non-classes respectively.

Experimental Runs: The evaluation relies on a variety of experimental runs. They extract classes among articles in a combination of one or more of { En , Fr , Es }, for English vs. French vs. Spanish articles respectively; and using a combination of one or more of { Dfd , Dng }, for extraction based on occurrences of defined vs. defining phrases. For example, whereas run $E_{[En,Dfd]}$ extracts classes from English articles based on defined phrases, $E_{[Fr,Dfd \cup Dng]}$ extracts from French articles based on defined and defining phrases. When enabling multiple languages or multiple types of evidence, they are required to be satisfied disjunctively (any) rather than conjunctively (all), before extracting an article as a class. Only for the purpose of evaluation, unless stated differently, when an article, in a language edition other than English (e.g., “Glande endocrine” in French), is extracted as a class by any experimental run, it is automatically either mapped to its equivalent English article, if any (e.g., “Endocrine gland”); or discarded.

Baseline Runs: In the first baseline, B_{Acl} , all Wikipedia articles in English are uniformly extracted as classes. A separate baseline that uniformly does not extract any Wikipedia articles as classes is not considered or evaluated, since its precision would be undefined and its recall zero. The second baseline, B_{Wrb} , is a rule-based method [20] that uses defined phrases as well as two other types of evidence based on occurrences of the title of a Wikipedia article in the article text, to decide whether the article is a class or not. By comparing against the second baseline, the proposed method is also transitively compared to other baselines (e.g., [19, 34]) against which the baseline itself was compared in [20].

4 EVALUATION RESULTS

Evaluation Metrics: Given the output set of classes extracted by a particular run, where the classes are Wikipedia articles, its quality and coverage are computed automatically relative to one of

Examples of Extracted Classes
Language (X?): <i>En</i> :
A-type main-sequence star, Arena, Consort of instruments, Contact sport, Incentive, Inductor, Limit ordinal, Paddle steamer, Passion (music), Path graph, Round barn, Satyr, Square knot (mathematics), Station (Australian agriculture), Synonym (taxonomy), Taskbar, Town square, Volcanic plug, Wide receiver
Language (X?): <i>Fr</i> :
Avion militaire, Bidon (récipient), Blouson, Bouteille, Colorant, Corregimiento, Cratère volcanique, Église-halle, Facteur de risque, Fonction gaussienne, Hutte, Matelas, Nuage de points (statistique), Organisme de normalisation, Pli (géologie), Poste de traite, Réseau de Bravais, Salle d'exposition, Scintillateur, Service web, Visiocasque
Language (X?): <i>Es</i> :
Algoritmo cuántico, Aminoácido, Arseniato, Banco, Cuórum, Diccionario, Ecomuseo, Epítome, Estación de trabajo de audio digital, Función analítica, Grupo abeliano, Hernia, Mancha (suciedad), Numerónimo, Nunatak, Proceso estocástico, Representación de grupo, Salón del automóvil, Sistema electoral, Sustancia simple, Tormenta, Tubérculo

Table 1: Examples of Wikipedia articles extracted as classes simultaneously by *Dfd* and *Dng* from some language (X?)

Evidence (Y?): Examples of Extracted Classes
Language (X?): <i>En</i> :
<i>Dfd</i> : Analyser, CAT(k) group, Cable-backed bow, Caret, Core concern, Covered warrant, Debenture (sport), Electrophysiology study, Eyelash curler, Immediately-invoked function expression, Koniscope, Message consumer, Motivational poster, NIH shift, Pole figure, Superphone
<i>Dng</i> : 3D computer graphics software, Bighorn sheep, Headworks, Issa (clan), Mainair Blade, Maple leaf, Metropolitan regions in Germany, Nucleoside analogue, Peptide, Prefecture-level city, Revocation, Separable state, Suicide attempt, Table setting
Language (X?): <i>Fr</i> :
<i>Dfd</i> : Ailier (handball), Anomère, Bague ecclésiastique, Civil township, Douelle (Lot), Excentrique (visserie), Intégrale paramétrique, Isotopologue, Logarithme d'une matrice, Organe de réserve, Organisme de bassin versant, Plancher stalagmitique, Roulement de tambour
<i>Dng</i> : Études supérieures, Algèbre sur un corps, Aspirant, Coefficient de transfert thermique, Croix de la Passion, Encorbellement, Fugue, Grades de la Marine nationale (France), Immeuble de grande hauteur, Jeu d'aventure graphique, Langage, Maître (arts martiaux), Matelot
Language (X?): <i>Es</i> :
<i>Dfd</i> : Agujero de tono, Bolsa Gamow, Brazaletes (tela), Célula Hfr, Cursor (base de datos), Diagrama de casos de uso, Ficha de datos de seguridad, Girómetro, Infusor de té, Láser Nd-YAG, Movimiento armónico complejo, Operación ternaria, Peine de frecuencias ópticas, Unión civil
<i>Dng</i> : Brote de rayos gamma, Cartel (organización ilícita), Crisis diplomática, Desastre natural, Embalse de usos múltiples, Forma musical, Juego en la nube, Lipoproteína, Minoría étnica, Periférico (informática), Phlebovirus, Sémola, Sal ácida, Título académico

Table 2: Examples of Wikipedia articles extracted as classes only by one of *Dfd* or *Dng* from some language (X?)

the evaluation sets. Precision is the fraction of extracted classes that also appear in the evaluation set (as gold classes or gold non-classes), which are gold classes. Recall is the fraction of gold classes from the evaluation set, which are extracted.

Run	Eval Set	Scores		
		P	R	F
Baseline runs:				
B_{Acl}	S_W	0.093	1.000	0.171
	S_D	0.036	1.000	0.070
	S_Q	0.362	1.000	0.532
B_{Wrb}	S_W	0.938	0.824	0.877
	S_D	0.914	0.438	0.593
	S_Q	0.973	0.702	0.816
Experimental runs:				
$E_{[En,Dfd]}$	S_W	0.966	0.638	0.768
	S_D	0.933	0.384	0.544
	S_Q	0.989	0.517	0.679
$E_{[En,Dng]}$	S_W	0.837	0.804	0.820
	S_D	1.000	0.096	0.175
	S_Q	0.979	0.387	0.554
$E_{[En,Dfd \cup Dng]}$	S_W	0.830	0.855	0.842
	S_D	0.938	0.411	0.571
	S_Q	0.979	0.652	0.783
$E_{[LL,Dfd]}$	S_W	0.963	0.736	0.834
	S_D	0.909	0.411	0.566
	S_Q	0.986	0.605	0.750
$E_{[LL,Dng]}$	S_W	0.806	0.828	0.817
	S_D	1.000	0.123	0.220
	S_Q	0.980	0.412	0.580
$E_{[LL,Dfd \cup Dng]}$	S_W	0.799	0.875	0.835
	S_D	0.919	0.466	0.618
	S_Q	0.977	0.704	0.819

Table 3: Precision and recall over the evaluation sets. Computed for baseline runs, as well as for experimental runs with extraction from English (*En*) or from all languages (*LL*, i.e., $En \cup Fr \cup Es$); and using various combinations of defined (*Dfd*) or defining (*Dng*) phrases or both of them ($Dfd \cup Dng$) as evidence (P=precision; R=recall; F=balanced F-score)

Precision and Recall: Table 1 shows articles extracted as classes simultaneously by both types of evidence *Dfd* and *Dng*. Conversely, Table 2 considers *Dfd* and *Dng* individually rather than in combination, showing random samples of articles extracted as classes by one of the two types of evidence but not the other.

Table 3 measures the performance of the experimental runs from the method proposed here, against baseline runs from alternative methods. Among experimental runs, combining evidence from all languages gives better performance, as measured by F-scores, than when extracting from each language individually. The results suggest that *Dfd* and *Dng* offer complementary clues which, when combined, extract a larger number of correct classes. The baseline run B_{Acl} achieves high recall at the expense of low precision. Comparatively, the experimental runs extract better classes, as indicated by F-scores that are higher across the evaluation sets. Furthermore, the proposed method gives competitive results relative to B_{Wrb} . This is encouraging, since the latter relies on more types of evidence in more languages than the proposed method (cf. [20]).

5 RELATED WORK

Previous work in open-domain information extraction [7, 8, 17] often uses Wikipedia data [29, 30].

Articles in Wikipedia are organized into fine-grained categories, which in turn are organized into iteratively coarser-grained categories. Collecting Wikipedia categories as classes would not be a sufficient solution to the problem of identifying classes in Wikipedia. Wikipedia categories often do not correspond to classes. In addition, Wikipedia categories without a corresponding article, of which there are many, have little utility, as suggested by existing work where Wikipedia serves as the reference resource in some task virtually always relying on articles rather than categories [1, 4, 11, 21, 23], with few exceptions [16].

Most methods that distinguish classes in Wikipedia [19, 34] require access to a part of speech tagger, a syntactic parser or a named entity recognizer and apply to English data only. In contrast, our method does not need access to any linguistic processing tools and is applicable to multiple languages. It proposes and investigates the approximation of occurrences of both defined and defining phrases, as evidence towards identifying classes; and achieves competitive results when combining them, relative to using defined phrases but not defining phrases as reported in [20]. The latter also distinguishes classes without any linguistic processing tools, based on various types of evidence within Wikipedia including occurrences of defined phrases.

6 CONCLUSION

The lightweight detection of a small number of definitional constructs in Wikipedia articles, coupled with their disambiguation based on internal hyperlinks in the articles, extracts classes at encouraging precision and recall. The constructs (patterns) being used are simple and by no means exhaustive. Nevertheless, they identify two types of evidence that, together, act as complementary signals. Current work investigates other types of relevant constructs, besides defined and defining phrases, that occur in Wikipedia articles and are evidence that the articles are classes.

REFERENCES

- [1] R. Blanco, G. Ottaviano, and E. Meij. 2015. Fast and Space-Efficient Entity Linking in Queries. In *Proceedings of the 8th ACM Conference on Web Search and Data Mining (WSDM-15)*. Shanghai, China, 179–188.
- [2] N. Calzolari and E. Picchi. 1988. Acquisition of Semantic Information from an On-line Dictionary. In *Proceedings of the 12th International Conference on Computational Linguistics (COLING-88)*. Budapest, Hungary, 87–92.
- [3] D. Chen, A. Fisch, J. Weston, and A. Bordes. 2017. Reading Wikipedia to Answer Open-Domain Questions. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (ACL-17)*. Vancouver, Canada, 1870–1879.
- [4] A. Chisholm and B. Hachey. 2015. Entity disambiguation with Web links. *Transactions of the Association for Computational Linguistics* 3 (2015), 145–156.
- [5] X. Du and C. Cardie. 2018. Harvesting Paragraph-level Question-Answer Pairs from Wikipedia. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (ACL-18)*. Melbourne, Australia, 1907–1917.
- [6] F. Ensan and E. Bagheri. 2017. Document Retrieval Model Through Semantic Linking. In *Proceedings of the 10th ACM Conference on Web Search and Data Mining (WSDM-17)*. Cambridge, United Kingdom, 181–190.
- [7] P. Ernst, A. Siu, and G. Weikum. 2018. HighLife: Higher-Arity Fact Harvesting. In *Proceedings of the 2018 Web Conference (WWW-18)*. Lyon, France, 1013–1022.
- [8] O. Etzioni, A. Fader, J. Christensen, S. Soderland, and Mausam. 2011. Open Information Extraction: The Second Generation. In *Proceedings of the 22nd International Joint Conference on Artificial Intelligence (IJCAI-11)*. Barcelona, Spain, 3–10.
- [9] C. Fellbaum (Ed.). 1998. *WordNet: An Electronic Lexical Database and Some of its Applications*. MIT Press.
- [10] T. Flati, D. Vannella, T. Pasini, and R. Navigli. 2014. Two Is Bigger (and Better) Than One: the Wikipedia Bitaxonomy Project. In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics (ACL-14)*. Baltimore, Maryland, 945–955.
- [11] O. Ganea, M. Ganea, A. Lucchi, C. Eickhoff, and T. Hofmann. 2016. Probabilistic Bag-Of-Hyperlinks Model for Entity Linking. In *Proceedings of the 25th World Wide Web Conference (WWW-16)*. Montreal, Canada, 927–938.
- [12] P. Gupta, S. Rajaram, H. Schütze, and T. Runkler. 2019. Neural Relation Extraction Within and Across Sentence Boundaries. In *Proceedings of the 33rd National Conference on Artificial Intelligence (AAAI-19)*. Honolulu, Hawaii, 6513–6520.
- [13] M. Hearst. 1992. Automatic acquisition of hyponyms from large text corpora. In *Proceedings of the 14th International Conference on Computational Linguistics (COLING-92)*. Nantes, France, 539–545.
- [14] J. Hoffart, F. Suchanek, K. Berberich, and G. Weikum. 2013. YAGO2: a Spatially and Temporally Enhanced Knowledge Base from Wikipedia. *Artificial Intelligence Journal. Special Issue on Artificial Intelligence, Wikipedia and Semi-Structured Resources* 194 (2013), 28–61.
- [15] J. Hu, G. Wang, F. Lochovsky, J. Sun, and Z. Chen. 2009. Understanding User's Query Intent with Wikipedia. In *Proceedings of the 18th World Wide Web Conference (WWW-09)*. Madrid, Spain, 471–480.
- [16] D. Ma, Y. Chen, K. Chang, and X. Du. 2018. Leveraging Fine-Grained Wikipedia Categories for Entity Search. In *Proceedings of the 2018 Web Conference (WWW-18)*. Lyon, France, 1623–1632.
- [17] Mausam, M. Schmitz, S. Soderland, R. Bart, and O. Etzioni. 2012. Open Language Learning for Information Extraction. In *Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning (EMNLP-CoNLL-12)*. Jeju Island, Korea, 523–534.
- [18] V. Nastase and M. Strube. 2008. Decoding Wikipedia Categories for Knowledge Acquisition. In *Proceedings of the 23rd National Conference on Artificial Intelligence (AAAI-08)*. Chicago, Illinois, 1219–1224.
- [19] V. Nastase and M. Strube. 2013. Transforming Wikipedia into a Large Scale Multilingual Concept Network. *Artificial Intelligence* 194 (2013), 62–85.
- [20] M. Paşca. 2018. Finding Needles in an Encyclopedic Haystack: Detecting Classes Among Wikipedia Articles. In *Proceedings of the 2018 Web Conference (WWW-18)*. Lyon, France, 1267–1276.
- [21] X. Pan, T. Cassidy, U. Hermjakob, H. Ji, and K. Knight. 2015. Unsupervised Entity Linking with Abstract Meaning Representation. In *Proceedings of the 2015 Conference of the North American Association for Computational Linguistics (NAACL-HLT-15)*. Denver, Colorado, 1130–1139.
- [22] L. Ratnov and D. Roth. 2012. Learning-Based Multi-Sieve Co-Reference Resolution with Knowledge. In *Proceedings of the Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning (EMNLP-CoNLL-12)*. Jeju Island, Korea, 1234–1244.
- [23] L. Ratnov, D. Roth, D. Downey, and M. Anderson. 2011. Local and Global Algorithms for Disambiguation to Wikipedia. In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics (ACL-11)*. Portland, Oregon, 1375–1384.
- [24] U. Scaella, P. Ferragina, A. Marino, and M. Ciarmita. 2012. Topical Clustering of Search Results. In *Proceedings of the 5th ACM Conference on Web Search and Data Mining (WSDM-12)*. Seattle, Washington, 223–232.
- [25] J. Seitner, C. Bizer, K. Eckert, S. Faralli, R. Meusel, H. Paulheim, and S. Ponzetto. 2016. A Large Database of Hyponymy Relations Extracted from the Web. In *Proceedings of the 10th Conference on Language Resources and Evaluation (LREC-16)*. Portoroz, Slovenia, 360–367.
- [26] A. Singhal. 2012. Introducing the Knowledge Graph: Things, not Strings. Corporate blog.
- [27] Y. Sun, A. Singla, D. Fox, and A. Krause. 2015. Building Hierarchies of Concepts via Crowdsourcing. In *Proceedings of the 24th International Joint Conference on Artificial Intelligence (IJCAI-15)*. Buenos Aires, Argentina, 844–851.
- [28] C. Tan, F. Wei, P. Ren, W. Lv, and M. Zhou. 2017. Entity Linking for Queries by Searching Wikipedia Sentences. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing (EMNLP-17)*. Copenhagen, Denmark, 68–77.
- [29] D. Tsurel, D. Pelleg, I. Guy, and D. Shahaf. 2017. Fun Facts: Automatic Trivia Fact Extraction from Wikipedia. In *Proceedings of the 10th ACM Conference on Web Search and Data Mining (WSDM-17)*. Cambridge, United Kingdom, 345–354.
- [30] F. Wu and D. Weld. 2010. Open Information Extraction using Wikipedia. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics (ACL-10)*. Uppsala, Sweden, 118–127.
- [31] W. Wu, H. Li, H. Wang, and K. Zhu. 2012. Probase: a Probabilistic Taxonomy for Text Understanding. In *Proceedings of the 2012 International Conference on Management of Data (SIGMOD-12)*. Scottsdale, Arizona, 481–492.
- [32] Y. Yan, N. Okazaki, Y. Matsuo, Z. Yang, and M. Ishizuka. 2009. Unsupervised Relation Extraction by Mining Wikipedia Texts Using Information from the Web. In *Proceedings of the 47th Annual Meeting of the Association for Computational Linguistics (ACL-IJCNLP-09)*. Singapore, 1021–1029.
- [33] S. Zhang and K. Balog. 2018. Ad Hoc Table Retrieval Using Semantic Similarity. In *Proceedings of the 2018 Web Conference (WWW-18)*. Lyon, France, 1553–1562.
- [34] C. Zirn, V. Nastase, and M. Strube. 2008. Distinguishing Between Instances and Classes in the Wikipedia Taxonomy. In *Proceedings of the 5th European Semantic Web Conference (ESWC-08)*. Tenerife, Spain, 376–387.