

Entity Extraction and Consolidation for Social Web Content Preservation

Stefan Dietze¹, Diana Maynard², Elena Demidova¹, Thomas Risse¹, Wim Peters²,
Katerina Doka³, Yannis Stavrakas³

¹ L3S Research Center, Leibniz University, Hannover, Germany
{dietze, nunes, demidova, risse}@l3s.de

² Department of Computer Science, University of Sheffield, Sheffield, UK
{diana, wim}@dcs.shef.ac.uk

³ IMIS, RC ATHENA, Artemidos 6, Athens 15125, Greece
katerina@cslab.ece.ntua.gr; yannis@imis.athenainnovation.gr

Abstract. With the rapidly increasing pace at which Web content is evolving, particularly social media, preserving the Web and its evolution over time becomes an important challenge. Meaningful analysis of Web content lends itself to an entity-centric view to organise Web resources according to the information objects related to them. Therefore, the crucial challenge is to extract, detect and correlate entities from a vast number of heterogeneous Web resources where the nature and quality of the content may vary heavily. While a wealth of *information extraction* tools aid this process, we believe that, the *consolidation* of automatically extracted data has to be treated as an equally important step in order to ensure high quality and non-ambiguity of generated data. In this paper we present an approach which is based on an iterative cycle exploiting Web data for (1) targeted archiving/crawling of Web objects, (2) entity extraction, and detection, and (3) entity correlation. The long-term goal is to preserve Web content over time and allow its navigation and analysis based on well-formed structured RDF data about entities.

Keywords. Knowledge Extraction, Linked Data, Data Consolidation, Data Enrichment, Web Archiving, Entity Recognition

1 Introduction

Given the ever increasing pace at which Web content is constantly evolving, adequate Web archiving and preservation have become a cultural necessity. Along with “common” challenges of digital preservation, such as media decay, technological obsolescence, authenticity and integrity issues, Web preservation has to deal with the sheer size and ever-increasing growth rate of Web content. This in particular applies to user-generated content and social media, which is characterized by a high degree of *diversity*, heavily *varying quality* and *heterogeneity*. Instead of following a *collect-all* strategy, archival organizations are striving to build *focused archives* that revolve around a particular *topic* and reflect the diversity of information people are interested in. Thus, focused archives largely revolve around the *entities* which define a topic or

area of interest, such as persons, organisations and locations. Hence, extraction of entities from archived Web content, in particular social media, is a crucial challenge in order to allow semantic search and navigation in Web archives and the relevance assessment of a given set of Web objects for a particular *focused crawl*.

However, while tools are available for information extraction from more formal text, social media affords particular challenges to knowledge acquisition. These are detailed more explicitly in Section 3. This calls for a range of specific strategies and techniques to *consolidate*, *enrich*, *disambiguate* and *interlink* extracted data. This in particular benefits from taking advantage of existing knowledge, such as Linked Open Data [1], to compensate for, disambiguate and remedy degraded information. While data consolidation techniques traditionally exist independent from named entity recognition (NER) technologies, their coherent integration into unified workflows is of crucial importance to improve the wealth of automatically extracted data on the Web. This becomes even more crucial with the emergence of an increasing variety of publicly available and end-user friendly knowledge extraction and NER tools such as DBpedia Spotlight¹, GATE², Open Calais³, Zemanta⁴.

In this paper, we introduce an integrated approach to extracting and consolidating structured knowledge about entities from archived Web content. This knowledge will in the future be used to facilitate semantic search of Web archives and to further guide the crawl. This work was developed in the EC-funded Integrating Project ARCOMEM⁵. Note, while temporal aspects related to term and knowledge evolution are substantial to Web preservation, these are currently under investigation [24] but out of scope for this paper.

2 Related Work

Entity recognition is one of the major tasks within information extraction and may encompass both NER and term extraction. Entity recognition may involve rule-based systems [13] or machine learning techniques [14]. Term extraction involves the identification and filtering of term candidates for the purpose of identifying domain-relevant terms or entities. The main aim in automatic term recognition is to determine whether a word or a sequence of words is a term that characterises the target domain. Most term extraction methods use a combination of linguistic filtering (e.g. possible sequences of part of speech tags) and statistical measures (e.g. tf.idf) [15] and [16], to determine the salience of each term candidate for each document in the corpus [23].

Data consolidation has to cover a variety of areas such as enrichment, entity/identity resolution for disambiguation as well as clustering and correlation to consolidate disparate data. In addition, link prediction and discovery is of crucial importance to enable clustering and correlation of enriched data sources. A variety of

¹ <http://spotlight.dbpedia.org>

² <http://gate.ac.uk/>

³ <http://www.opencalais.com/>

⁴ <http://www.zemanta.com/>

⁵ <http://www.arcomem.eu>

methods for entity resolution have been proposed, using relationships among entities [7], string similarity metrics [6], as well as transformations [9]. An overview of the most important works in this area can be found in [8]. As opposed to entity correlation techniques exploited in this paper, text clustering of documents exploits feature vectors, to represent documents according to contained terms [10][11][12]. Clustering algorithms measure the similarity across the documents and assign the documents to the appropriate clusters based on this similarity. Similarly, vector-based approaches have been used to map distinct ontologies and datasets [2][3]. As opposed to text clustering, entity correlation and clustering takes advantage of background knowledge from related datasets to correlate previously extracted entities. Therefore, link discovery is another crucial area to be considered. Graph summarization predicts links in annotated RDF graphs. A detailed survey of link predictions techniques in complex networks and social network are presented by [4] and [5], respectively.

3 Challenges and overall approach

ARCOMEM follows a use case-driven approach based on scenarios aimed at creating focused Web archives. We deploy a document repository of crawled Web content and a structured RDF knowledge base containing metadata about entities detected in the archived content. Archivists can specify or modify crawl specifications (fundamentally consisting of selected sets of relevant entities and topics). The intelligent crawler will be able to learn about crawl intentions and to refine a crawling strategy on-the-fly. This is especially important for long running crawls with broader topics, such as the financial crisis or elections, where entities are changing more frequently and hence require regular adaptation of the crawl specification. End-user applications allow users to search and browse the archives by exploiting automatically extracted metadata about entities and topics.

Fundamental to both crawl strategy refinement and Web archive navigation is the efficient extraction of entities from archived Web content. In particular, social media poses a number of challenges for language analysis tools due to the degraded nature of the text, especially where tweets are concerned. In one study, the Stanford NER tagger dropped from 90.8% F1 to 45.88% when applied to a corpus of tweets [17]. [19] also demonstrate some of the difficulties in applying traditional POS tagging, chunking and NER techniques to tweets, while language identification tools typically also do not work well on short sentences. Problems are caused by incorrect spelling and grammar, made-up words, hashtags, @ signs and emoticons, unorthodox capitalisation, and spellings (e.g duplication of letters in words for emphasis, text speak). Since tokenisation, POS tagging and matching against pre-defined gazetteer lists are key to NER, it is important to resolve these problems: we adopt methods such as adapting tokenisers, using techniques from SMS normalisation, retraining language identifiers, use of case-insensitive matching in certain cases, using shallow techniques rather than full parsing, and using more flexible forms of matching.

Entity extraction and enrichment is covered by a set of dedicated components which have been incorporated into a dedicated processing chain (Figure 1) which handles

NER and consolidation (enrichment, clustering, disambiguation) as part of one coherent workflow.

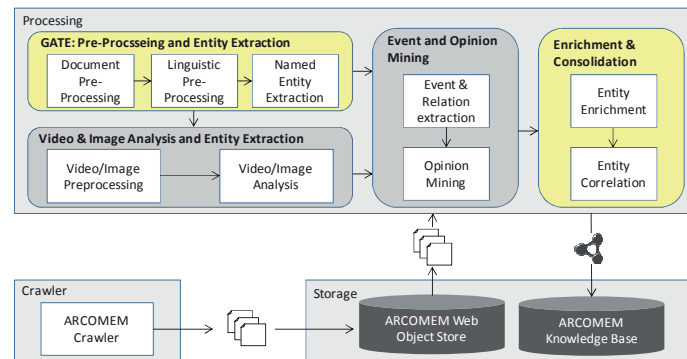


Fig. 1. Entity extraction and consolidation processing chain

The ARCOMEM storage composed of the *object store* and the *knowledge base* handles (a) binary data, in the form of Web objects, which represent the original content collected by the crawler; and (b) semi-structured data, in the form of RDF⁶ triples (Web object annotations). Storage is based on a distributed solution that combines the MapReduce [9] paradigm and NoSQL databases and is realised based on HBase⁷ (see also [25]). The *ARCOMEM data model*⁸ provides an RDF schema to reflect the informational needs for knowledge capturing, crawling, and preservation (see [20] for details).

Within the ARCOMEM model, "entity" encompasses both traditional Named Entities and also single and multi-word terms: the recognition of both is done using GATE tools. GATE has been chosen over other NLP tools primarily for its coverage, extensibility and flexibility: it has a wide range of NLP components, which are easily modifiable for the demands of the project, unlike tools such as OpenCalais and DBpedia Spotlight which are more limited in scope. While extracted data is already classified and labelled as a result of the extraction process, it is nevertheless (i) heterogeneous, i.e. not well interlinked, (ii) ambiguous and (iii) provides only very limited information. This is due to data being extracted by different components and during independent processing cycles, since the tools in GATE have no possibility to perform co-reference on entities generated asynchronously across multiple documents. For instance, during one particular cycle, the text analysis component might detect an entity from the term "Ireland", while during later cycles, entities based on the term "Republic of Ireland" or the German term "Irland" might be extracted, together with, the entity "Dublin". These would all be classified as entities of type *Location* and correctly stored in the data store as disparate entities described according to the data

⁶ <http://www.w3.org/RDF/>

⁷ Apache Foundation; The Apache HBase Project: <http://hbase.apache.org/>

⁸ <http://www.gate.ac.uk/ns/ontologies/arcomem-datamodel.rdf>

model. Thus, *Enrichment and Consolidation* (Fig. 1) follows three aims: (a) *enrich existing entities* with related publicly available knowledge; (b) *disambiguation*, and (c) identify *data correlations* such as the ones illustrated above. This is achieved by mapping isolated entities to concepts (nodes) within reference datasets (*enrichment*) and exploiting the corresponding graphs to discover correlations. Therefore, we exploit publicly available data from the Linked Open Data cloud which offers a vast amount of data of both domain-specific and domain-independent nature (the current release consists of 31 billion distinct triples, i.e. RDF statements⁹).

4 Implementation

For entity recognition, we use a modified version of ANNIE [18] to find mentions of *Person*, *Location*, *Organization*, *Date*, *Time*, *Money* and *Percent*. We included extra subtypes of *Organization* such as *Band* and *Political Party*, and have made various modifications to deal with the problems specific to social media such as incorrect English (see [21] for more details). The entity extraction framework can be divided into the following components (GATE component in Fig. 1) which are executed sequentially over a corpus of documents:

- Document Pre-processing (document format analysis, content detection)
- Linguistic Pre-processing (language detection, tokenisation, POS tagging etc)
- Named Entity Extraction: Term Extraction (generation of ranked list of terms and thresholding) & NER (gazetteers, rule-based grammars and co-reference)

For term extraction, we use an adapted version of TermRaider¹⁰. This considers noun phrases (NPs) as candidate terms (as determined by linguistic pre-processing), and ranks them in order of termhood according to 3 different scoring functions: (1) basic tf.idf (2) an augmented tf.idf which also takes into account the tf.idf score of any hyponyms of a candidate term, and (3) the Kyoto score based on [22] which takes into account the number of hyponyms of a candidate term occurring in the document. All are normalised to represent a value between 0 and 100. A candidate term is not considered an entity if it matches or is contained within an existing Named Entity, to avoid duplication. Also, we have set a threshold score above which we consider a candidate term to be valid. This threshold is a parameter which can be manually changed at any time – currently it is set to an augmented score of 45, i.e. only terms with a score of 45 or greater will be used by later processes.

The entity extraction generates RDF data describing NEs and terms according to the ARCOMEM data model which is pushed to our knowledge base and directly digested by our *Enrichment & Consolidation* component (Fig. 1). The latter exploits (a) the *entity label* and (b) the *entity type* to expand, disambiguate and correlate extracted data. Note that an entity/event label might correspond directly to a label of

⁹ <http://lod-cloud.net/state>

¹⁰ <http://gate.ac.uk/projects/arcomem/TermRaider.html>

one unique node in a structured dataset (as is likely for an entity of type person labelled “Angela Merkel”), but might also correspond to more than one node/concept, as is the case for most of the events in our dataset. For instance, the event labeled “Jean Claude Trichet gives keynote at ECB summit” will most likely be enriched with links to concepts representing the ECB as well as Jean Claude Trichet. Our approach is based on the following steps (reflected in Fig. 1):

S1. *Entity enrichment*

S1.a. Translation: we determine the language of the entity label, and, if necessary, translate it into English using an online translation service.

S1.b. Enrichment: co-referencing with related entities in reference datasets.

S2. *Entity correlation and clustering*

In order to obtain enrichments for these entities we perform queries on external knowledge bases. Our current enrichment approach uses DBpedia¹¹ and Freebase¹² as reference datasets, though it is envisaged to expand this approach with additional and more domain-specific datasets, e.g., event-specific ones. DBpedia and Freebase are particularly well-suited due to their vast size, the availability of disambiguation techniques which can utilise the variety of multilingual labels available in both datasets for individual data items and the level of inter-connectedness of both datasets, allowing the retrieval of a wealth of related information for particular items. In the case of DBpedia, we make use of the DBpedia Spotlight service which enables an approximate string matching with adjustable confidence level in the interval [0,1]. As part of our evaluation (Section 6), we experimentally selected a confidence level of 0.6 which provided the best balance of precision and recall. Note that Spotlight offers NER capabilities complementary to GATE. However, these were only utilised in cases where entities/events were not in a rather atomic form, as is often the case for events which mostly consists of free text descriptions such the one mentioned above.

Freebase contains about 22 million entities and more than 350 millions facts in about 100 domains. Keyword queries over Freebase are particularly ambiguous due to the size and the structure of the dataset. In order to reduce query ambiguity, we used the Freebase API and restricted the types of the entities to be matched using a manually defined type mapping from ARCOMEM to Freebase entity types. For example, we mapped the ARCOMEM type “person” to the “people/person” type of Freebase, and the ARCOMEM type “location” to the Freebase types “location/continent”, “location/location” and “location/country”. For instance, an ARCOMEM entity of type “Person” with the label “Angela Merkel” is mapped to the Freebase MQL query that retrieves one unique Freebase entity with the mid= “/m/0jl0g”. With respect to data correlation, we distinct *direct* as well as *indirect* correlations. Please note, that a *correlation* does not describe any notion of equivalence (e.g. similar to *owl:sameAs*) but merely a meaningful level of relatedness.

Fig. 2 depicts both cases, direct as well as indirect correlations. Direct correlations are identified by means of equivalent and shared enrichments, i.e., any entities/events

¹¹ <http://dbpedia.org/>

¹² <http://www.freebase.com/>

sharing the same enrichments are supposedly correlated and hence clustered. A direct correlation is visible between the entity of type Person labeled “Jean Claude Trichet” and the event “Trichet warns of systemic debt crisis”. In addition, the retrieved enrichments associate the ARCOMEM entities and associated Web objects with the knowledge, i.e., data graph, available in associated reference datasets.

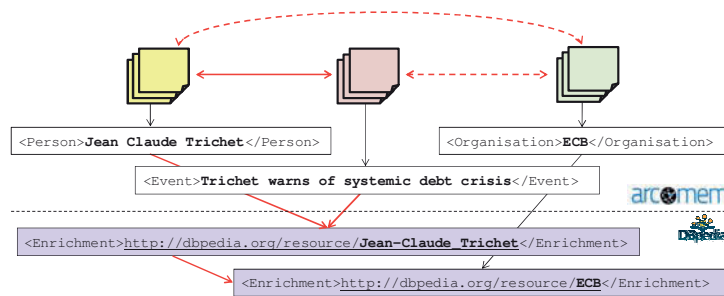


Fig. 2.Enrichment and correlation example: ARCOMEM Web objects, entities/events, associated DBpedia enrichments and identified correlations

For instance, the DBpedia resource of the European Central Bank (<http://DBpedia.org/resource/ECB>) provides additional facts (e.g., a classification as organisation, its members, or previous presidents) in a structured, and therefore, machine-processable form. Exploiting the graphs of underlying reference datasets allows us to identify additional, *indirect correlations*. While linguistic/syntactic approaches would fail to detect a relationship between the two enrichments above (Trichet, ECB) and hence their corresponding entities and Web objects, by analysing the DBpedia graph we are able to uncover a close relationship between the two (Trichet being the former ECB president). Hence, computing the *relatedness* of enrichments would allow us to detect indirect correlations to create a relationship (dashed line) between highly related entities/events, beyond mere equivalence.

Our current implementation is limited to detect direct correlations, while ongoing experiments based on graph analysis mechanisms aim to automatically measure *semantic relatedness* of entities in reference datasets to detect indirect relations. While in a large graph, all nodes are connected with each other in some way, a key research challenge is the investigation of appropriate graph navigation and analysis techniques to uncover indirect but semantically meaningful relationships between resources within reference datasets, and hence ARCOMEM entities and Web objects.

5 Results & evaluation

For our experiments, we used a dataset composed of English and German archived Web objects constituting a sample of crawls relating to the financial crisis¹³. The English content covered 32 Facebook posts, 41,000 tweets and 800 user comments from

¹³ Parts of the archived crawls are available at <http://collections.europarchive.org/arcomem/>.

greekcrisis.net. The German content consisted of archived data from the Austrian Parliament¹⁴ consisting of 326 documents (mostly PDF, some HTML).

Our extraction and enrichment experiments resulted in an evaluation dataset¹⁵ of 99,569 unique entities involving the types *Event*, *Location*, *Money*, *Organization*, *Person*, *Time*. Using the procedure described above, we obtained enrichments for 1,358 of the entities in our dataset using DBpedia (484 entities) and Freebase (975 entities). In total, we obtained 5,291 Freebase enrichments and 491 DBpedia enrichments. These enrichments built 5,801 entity-enrichment pairs, 5,039 with Freebase and 492 with DBpedia.

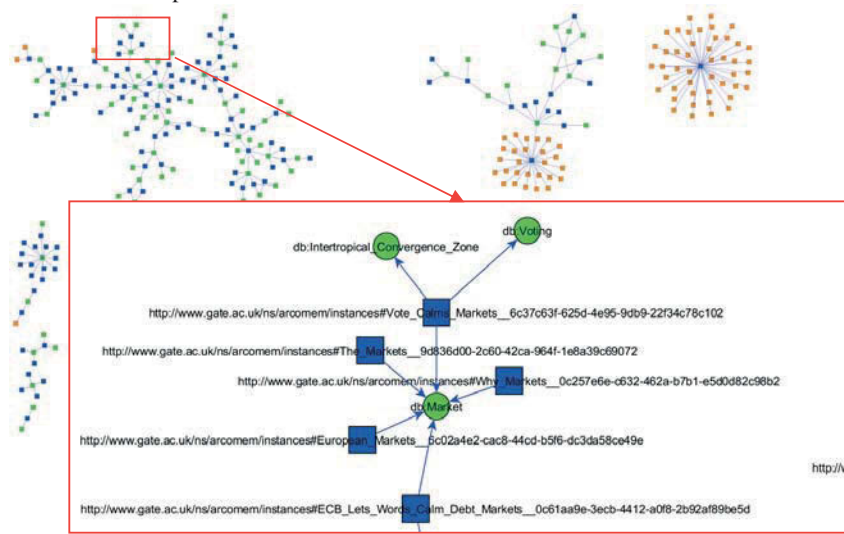


Fig. 3. Generated ARCOMEM graph and clusters

5.1 Entity extraction evaluation

We have performed initial evaluations on the various text analysis components. We manually annotated a small corpus of 20 Facebook posts (in English) from the dataset described above with named entities to form a gold standard corpus. This contained 93 instances of Named Entities. For evaluating TermRaider, we took a larger set of 80 documents from the financial crisis dataset, from which, TermRaider produced 1003 term candidates (merged from the results of the three different scoring systems). Three human annotators selected valid terms from that list, and we produced a gold standard of 315, comprising each term candidate selected by at least two annotators (221 terms selected by exactly two annotators and 94 selected by all three). While inter-annotator agreement was thus quite low, this is normal for a term extraction task

¹⁴ <http://www.parliament.gv.at/>

¹⁵ The SPARQL endpoint of our dataset (extracted entities and enrichments) is available at <http://arcomem.l3s.uni-hannover.de:9988/openrdf-sesame/repositories/arcomem-rdf?query>.

as it is extremely subjective; however, in future we will tighten the annotation guidelines and provide further training to the annotators with the aim of reaching a better consensus.

For the NE recognition evaluation, we compared the system annotations with the gold standard. The system achieved a Precision of 80% and a Recall of 68% on the task of NE detection (i.e. detecting whether an entity was present or not, regardless of its type). On the task of type determination (getting the correct type of the entity (Person, Organization, Location etc.)), the system performed with 98.8% Precision and 98.5% Recall. Overall (for the two tasks combined), this gives NE recognition scores of 79% Precision and 67% Recall. However, the results are slightly low because this actually includes Sentence detection also. Normally, Sentence detection is 100% accurate (or near enough), but in this case, it is subject to the language detection issue, because we only perform the entity detection on sentences deemed to be relevant (in the language of the task and which corresponds to the relevant part of the document - in this case, the actual text of the postings by the users). 26 of the missing system annotations in the document were outside the span of the sentences annotated, so could not have been annotated. Excluding these increased Recall from 68% to 83.9% for NE detection (shown in the table as "NE detection (adjusted)"), and from 67% to 73.5% for the complete NE recognition task (shown in the table as "Full NE recognition (adjusted)").

Table 1. NER evaluation results

Task	Precision	Recall	F1
NE detection	80%	68%	74%
NE detection (adjusted)	80%	83.9%	81.9%
Type determination	98.8%	98.5%	98.6%
Full NE recognition	79%	67%	72.5%
Full NE recognition (adjusted)	79%	82.1%	80.5%

For term recognitions, we compared the TermRaider output for each scoring system with the gold standard set of terms, at different levels of the ranked list, as shown in Figure 4. For the terms above the threshold, we achieved Precision scores of 31% and Recall of 90% for tf.idf, 73% Precision and 50% Recall for augmented tf.idf and 63% Precision and 17% Recall for the Kyoto score. For any further processing, we only use the terms scored by the augmented tf.idf above the threshold.

5.2 Enrichment and correlation evaluation

For this evaluation we randomly selected a set of entity-enrichment pairs. Our evaluation was performed manually by 6 judges including graduate computer science students and researchers. The judges were asked to assign scores to each entity-enrichment pair, with "0" for *incorrect*, and "1" for *correct*. We judge an enrichment as correct if it partially defines a specific dimension of the entity/event, that is, an enrichment does not need to completely match an entity. For instance, enrichments

referring to [http://dbpedia.org/resource/Doctor_\(title\)](http://dbpedia.org/resource/Doctor_(title)) and http://dbpedia.org/page/Angela_Merkel and enriching an entity of type Person labelled “Dr Angela Merkel” were both equally ranked as correct. This is due to entities and events being potentially related to multiple enrichments, each enriching a particular facet of the source entity/event. Each entity/enrichment pair was shown to at least 3 judges and an average of their scores was built to alleviate bias. In case an entity label did not make sense to a judge, we assumed that there has been an error in the extraction phase. In this case we asked the judges to mark the corresponding entity as invalid and excluded it from the evaluation.

We computed the average scores of entity-enrichment pairs across judges and averaged the scores obtained for each entity type. Table 4 presents the average scores of the enrichment-entity pairs obtained using DBpedia and Freebase for different ARCOMEM entity types.

Table 2. Enrichment evaluation results

Entity Type	Avg. Score DBpedia	Avg. Score Freebase	Avg. Score Total
Location	0.94	0.94	0.94
Money	0.63	-	0.63
Organization	0.93	1	0.97
Person	0.72	0.89	0.8
Time	1	-	1
Total	0.84	0.94	0.89

Our initial clustering approach simply correlated entities/events which share equivalent enrichments. In total we generated 1013 clusters with 2.85 entities on average, with a minimum of 2 and a maximum of 112 entities. Ambiguous enrichments led to redundant clusters and require additional disambiguation. For instance, a location entity labelled “Berlin” might be (correctly) enriched with <http://rdf.freebase.com/ns/m/0xfhc> and <http://rdf.freebase.com/ns/m/047ckrl> (each referring to a different location “Berlin”) requiring additional disambiguation to clean up the clusters. To this end, we exploit graph analysis methods to detect closeness of enrichments originating from the same object. For instance, measuring the relatedness of two location entities “Berlin” and “Angela Merkel” used to annotate the same Web object will allow us to disambiguate enrichments.

6 Discussion and future works

In this paper we have presented our current strategy for entity extraction and enrichment as realised within the ARCOMEM project, aimed at creating a large knowledge base of structured knowledge about archived heterogeneous Web content. Based on an integrated processing chain, we tackle entity consolidation and enrichment as implicit activity in the information extraction workflow.

The results of the entity extraction show respectable scores for this kind of social media data on which NLP techniques typically struggle. However, current work is

focusing on better handling of degraded English (tokenisation, language recognition etc) and especially of tweets, which should improve the entity extraction further. The enrichment results indicate a comparably good quality of generated enrichments. The results obtained from DBpedia Spotlight provided a lower recall, but introduced less ambiguous enrichments due to Spotlight's inherent disambiguation feature. On the other hand, partially matched keywords reduce the precision results. As future work, we foresee different directions to improve quality of the enrichment results. For example, one possibility is to use structured DBpedia queries to restrict entity types, similar to the approach used for Freebase. We also consider the introduction of subtypes of entities to further increase granularity of the types to be matched.

In addition, while preservation of Web content over time has to consider temporal aspects, evolution of entities and terms as well as time-dependent disambiguation are important research areas currently under investigation [24]. While our current data consolidation approach only detects direct relationships between entities sharing the same enrichments, our main efforts are dedicated to investigate graph analysis mechanisms. Thus, we aim to further take advantage of knowledge encoded in large reference graphs to automatically identify semantically meaningful relationships between disparate entities extracted during different processing cycles. Given the increasing use of both automated NER tools and reference datasets such as DBpedia, WordNet or Freebase, there is an increasing need for consolidating automatically extracted information on the Web which we aim to facilitate with our work.

Acknowledgments

This work is partly funded by the European Union under FP7 grant agreement n° 270239 (ARCOMEM).

References

- [1] Bizer, C., T. Heath, Berners-Lee, T. (2009). Linked data - The Story So Far. Special Issue on Linked data, International Journal on Semantic Web and Information Systems.
- [2] Dietze, S., and Domingue, J. (2008) Exploiting Conceptual Spaces for Ontology Integration, Workshop: Data Integration through Semantic Technology (DIST2008) Workshop at 3rd Asian Semantic Web Conference (ASWC) 2008, Bangkok, Thailand.
- [3] Dietze, S., Gugliotta, A., and Domingue, J. (2009) Exploiting Metrics for Similarity-based Semantic Web Service Discovery, IEEE 7th International Conference on Web Services (ICWS 2009), Los Angeles, CA, USA.
- [4] Lü, L., Zhou, T.: Link prediction in complex networks: a survey, *Physica A* 390 (2011), 1150–1170.
- [5] Hasan, M. A., Zaki, M. J.: A survey of link prediction in social networks. In C. Aggarwal, editor, *Social Network Data Analytics*, pages 243–276. Springer,
- [6] Cohen, W. W., Ravikumar, P. D., Fienberg, S. E.. A comparison of string distance metrics for name-matching tasks. In *IWeb*, 2003.
- [7] Dong, X., Halevy, A., Madhavan, J., Reference reconciliation in complex information spaces. In *SIGMOD*, 2005.

- [8] Elmagarmid, A. K., Ipeirotis, P. G., Verykios, V. S., Duplicate record detection: A survey. *TKDE*, 19(1), 2007.
- [9] Tejada, S., Knoblock, C. A., Minton, S., Learning domain-independent string transformation weights for high accuracy object identification. In *KDD*, 2002.
- [10] Boley, D., Principal Direction Divisive Partitioning. *Data Mining and Knowledge Discovery*, 2(4), 1998.
- [11] Broder, A., Glassman, S., Manasse, M., Zweig, G., Syntactic Clustering of the Web. In *Proceedings of the 6th International World Wide Web Conference*, pages 1997.
- [12] Hotho, A., Maedche, A., Staab, S., Ontology-based Text Clustering. In *Proceedings of the IJCAI Workshop on Text Learning: Beyond Supervision*, 2001.
- [13] Maynard, D., Tablan, V., Ursu, C., Cunningham, H., Wilks, Y., Named Entity Recognition from Diverse Text Types. *Recent Advances in Natural Language Processing 2001 Conference*, Tzigrav Chark, Bulgaria, 2001
- [14] Li, Y., Bontcheva, K., Cunningham, H., Adapting SVM for Data Sparseness and Imbalance: A Case Study on Information Extraction. *Natural Language Engineering*, 15(02), 241-271, 2009.
- [15] Buckley, C., G. Salton, G., Term-weighting approaches in automatic text retrieval. *Information Processing and Management*, 24(5):513–523, 1988.
- [16] Maynard, D., Li, Y., Peters, W., NLP techniques for term extraction and ontology population. In: Buitelaar, P. and Cimiano, P. (eds.), *Ontology Learning and Population: Bridging the Gap between Text and Knowledge*, pp. 171-199, IOS Press, Amsterdam (2008)
- [17] Lui, M., Baldwin, T., 2011. Cross-domain feature selection for language identification. In *Proceedings of 5th International Joint Conference on Natural Language Processing*, pages 553–561, November.
- [18] Cunningham, H., Maynard, D., Bontcheva, K., Tablan, V., 2002. GATE: A Framework and Graphical Development Environment for Robust NLP Tools and Applications. *Proceedings of the 40th Anniversary Meeting of the Association for Computational Linguistics (ACL'02)*.
- [19] Ritter, A., Clark, S., Mausam, Etzioni, O., 2011. Named entity recognition in tweets: An experimental study. In *Proc. of Empirical Methods for Natural Language Processing (EMNLP)*, Edinburgh, UK
- [20] Risse, T., Dietze, S., Peters, W., Doka, K., Stavarakas, Y., Senellart, P., Exploiting the Social and Semantic Web for guided Web Archiving, *The International Conference on Theory and Practice of Digital Libraries 2012 (TPDL2012)*, Cyprus, September 2012.
- [21] Maynard, D., Bontcheva, K., Rout, D., Challenges in developing opinion mining tools for social media. In *Proceedings of @NLP can u tag #usergeneratedcontent?! Workshop at LREC 2012*, May 2012, Istanbul, Turkey.
- [22] Bosma, W., Vossen, P., 2010. Bootstrapping languageneutral term extraction. In *7th Language Resources and Evaluation Conference (LREC)*, Valletta, Malta.
- [23] Deane, P. A nonparametric method for extraction of candidate phrasal terms, In *Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics*, 2005.
- [24] Tahmasebi, N., Risse, T., Dietze, S. (2011) Towards Automatic Language Evolution Tracking: A Study on Word Sense Tracking, *Joint Workshop on Knowledge Evolution and Ontology Dynamics 2011 (EvoDyn2011)*, at the 10th International Semantic Web Conference (ISWC2011), Bonn, Germany.
- [25] Weiss, C., Karras, P., Bernstein, A., Hexastore: sextuple indexing for semantic web data management. *Proceedings of the VLDB Endowment*, 1(1):1008–1019, 2008.