

Using Background Knowledge for Ontology Evolution

Fouad Zablith, Marta Sabou, Mathieu d’Aquin, and Enrico Motta

Knowledge Media Institute (KMi), The Open University
Walton Hall, Milton Keynes, MK7 6AA, United Kingdom
{f.zablith, r.m.sabou, m.daquin, e.motta}@open.ac.uk

Abstract. One of the current bottlenecks for automating ontology evolution is resolving the right links between newly arising information and the existing knowledge in the ontology. Most of existing approaches mainly rely on the user when it comes to capturing and representing new knowledge. Our ontology evolution framework intends to reduce or even eliminate user input through the use of background knowledge. In this paper, we show how various sources of background knowledge could be exploited for relation discovery. We perform a relation discovery experiment focusing on the use of WordNet and Semantic Web ontologies as sources of background knowledge. We back our experiment with a thorough analysis that highlights various issues on how to improve and validate relation discovery in the future, which will directly improve the task of automatically performing ontology changes during evolution.

1 Introduction

Ontologies are fundamental building blocks of the Semantic Web and are often used as the knowledge backbones of advanced information systems. As such, they need to be kept up to date in order to reflect the changes that affect the life-cycle of such systems (e.g. changes in the underlying data sets, need for new functionalities, etc). This task, described as the “timely adaptation of an ontology to the arisen changes and the consistent management of these changes”, is called *ontology evolution* [11].

While it seems necessary to apply such a process consistently for most ontology-based systems, it is often a time-consuming and knowledge intensive task, as it requires a knowledge engineer to identify the need for change, perform appropriate changes on the base ontology and manage its various versions. Several research efforts have addressed various phases of this complex process. A first category of approaches [12, 16, 19, 20] are concerned with formalisms for representing changes and for facilitating the versioning process. Another category of work [2, 5, 14, 15, 17] aim to identify potential novel information that should be added to the ontology. They do this primarily by exploiting the changes occurring in the various data sources underlying an information system (e.g. databases, text corpora, etc), or by interpreting trends in the behavior of the users of the system [2, 5]. A few systems from this category also investigate methods that

propose appropriate changes to an ontology given a piece of novel information. These methods typically rely on agent negotiation/multi-agent systems [14, 15, 17] to propose concrete changes which then are verified by the ontology curator.

Our hypothesis is that this process of prescribing ways in which a piece of new information can be added to a base ontology could be made more cost effective by relying on various sources of background knowledge to reduce human intervention. These sources can have varying levels of formality, ranging from formal knowledge structures, such as ontologies available on the Semantic Web, to thesauri (e.g. WordNet), or even unstructured textual data from the Web. We believe that, by combining such complementary sources of knowledge, a large part of the ontology evolution process can be automated. Therefore, we propose an ontology evolution framework, called Evolva¹, which once a set of terms has been identified as potentially relevant concepts to add to the ontology, makes use of external sources of background knowledge to establish relations between these terms and the knowledge already present in the ontology.

In this paper, we present a first implementation of the *relation discovery for ontology changes* process of Evolva, exploiting WordNet and Semantic Web ontologies as sources of background knowledge. We describe initial experiments realized with this process, showing the feasibility of the approach and pointing out possible ways to better exploit external sources of knowledge, as well as the base ontology, in the evolution process.

We start by describing a motivating scenario in Section 2 and providing a brief overview of Evolva in Section 3. In Section 4 we detail our ideas about the use of background knowledge sources. We describe the implementation details of our prototype (Section 5) followed by a discussion of our experimental results obtained in the context of the example scenario (Section 6). We finalize the paper with some notes on related work and our conclusions (Sections 7 and 8).

2 Example Scenario: The KMi Semantic Web Portal

The Knowledge Media Institute’s (KMi) Semantic Web portal² is a typical example of an ontology based information system. The portal provides access to various data sources (e.g. staff and publication databases, news stories) by relying on an ontology that represents the academic domain, namely the AKT ontology³. The ontology has been originally built manually and is automatically populated by relying on a set of manually established mapping rules [13].

However, apart from the population process, the evolution of this ontology was performed entirely manually. Indeed, as in this scenario ontology population is bound by strict and limited mapping expressions, when a new type of term (i.e. not covered by the mapping rules) is extracted, the intervention of the ontology administrator is required to modify the mappings. Moreover, the ontology schema can only be updated by the administrator. Finally, with no mechanism

¹ An overview of Evolva can be found in [22].

² <http://semanticweb.kmi.open.ac.uk>

³ <http://kmi.open.ac.uk/projects/akt/ref-onto/index.html>

to support recording and managing changes, it is difficult to maintain a proper versioning of the ontology. Therefore, as this manual evolution of the ontology could not follow the changes in the underlying data (which happen on a daily basis), the ontology was finally left outdated.

3 Evolva: An Ontology Evolution Framework

Evolva is a complete ontology evolution framework that relies on various sources of background knowledge to support its process. We hereby provide a brief overview of its five components design (Figure 1), which is only partially implemented now. More details are available in [22].

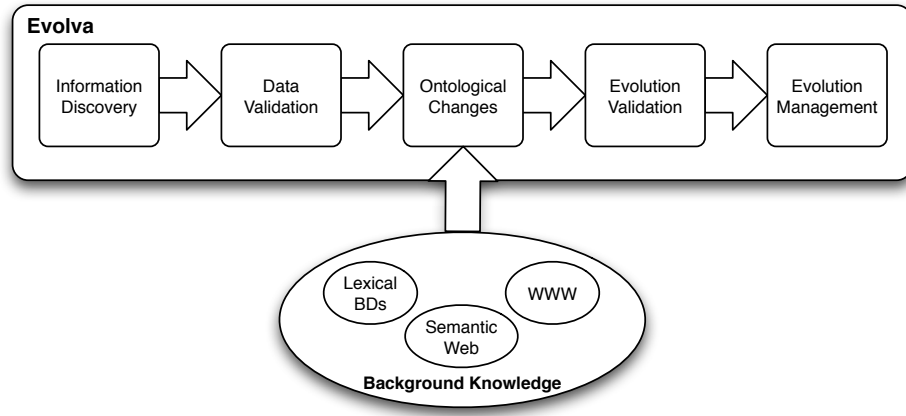


Fig. 1. The main components of Evolva.

Information Discovery. Our approach starts with discovering potentially new information from the data sources associated with the information system. Contrasting the ontology with the content of these sources is a way of detecting new knowledge that should be reflected by the base ontology. Data sources exist in various formats from unstructured data such as text documents or tags, to structured data such as databases and ontologies. This component handles each data source differently: (1) Text documents are processed using information extraction, ontology learning or entity recognition techniques. (2) Other external ontologies are subject to translation for language compatibility with the base ontology, and (3) database content is translated into ontology languages.

Data Validation. Discovered information are validated in this component. We rely on a set of heuristic rules such as the length of the extracted terms. This is especially needed for information discovered from text documents, as information extraction techniques are likely to introduce noise. For example, most of the two-letter terms extracted from KMi's news corpora are meaningless and should be

discarded. In the case of structured data, this validation is less needed as the type of information is explicitly defined.

Ontological Changes. This component is in charge of establishing relations between the extracted terms and the concepts in the base ontology. These relations are identified by exploring a variety of background knowledge sources, as we will describe in the next section. Appropriate changes will be performed directly to the base ontology and recorded by using the formal representations proposed by [12] and [19], which we will explore at a later stage of our research.

Evolution Validation. Performing ontology changes automatically may introduce inconsistencies and incoherences in the base ontology. Also, due to having multiple data sources, data duplication is likely to arise. Conflicting knowledge is highly possible to occur and should be handled by automated reasoning. As evolution is an ongoing process, many statements are time dependent and should be treated accordingly by applying temporal reasoning techniques.

Evolution Management. Managing the evolution will be about giving the ontology curator a degree of control over the evolution, as well as propagating changes to the dependent components of the ontology such as other ontologies or applications. User control will deal with tracking ontology changes, spotting and solving unresolved problems.

While this section presents an overview of our ontology evolution framework, our focus in this paper is on the relation discovery process that occurs in the *ontological changes* component.

4 The Role of Background Knowledge in Evolva

A core task in most ontology evolution scenarios is the integration of new knowledge into the base ontology. We focus on those scenarios in which such new knowledge is extracted as a set of emerging terms from textual corpora, databases, or domain ontologies. Traditionally this process of integrating a new set of emerging terms is performed by the ontology curator. For a given term, he/she would rely on his/her own knowledge of the domain to identify, in the base ontology, elements related to the term, as well as the actual relations they share. As such, it is a time consuming process, which requires the ontology curator to know well the ontology, as well as being an expert in the domain it covers.

Evolva makes use of various background knowledge sources to identify relations between new terms and ontology elements. The hypothesis is that a large part of the process of updating an ontology with new terms can be automated by using these sources as an alternative to the curator's domain knowledge. We have identified several potential sources of background knowledge. For example, thesauri such as WordNet have been long used as a reference resource for establishing relations between two given concepts, based on the relation that exists between their synsets. Because WordNet's dictionary can be downloaded and accessed locally by the system and because a variety of relation discovery techniques have been proposed and optimized, exploring this resource is quite fast. Online ontologies constitute another source of background knowledge which has

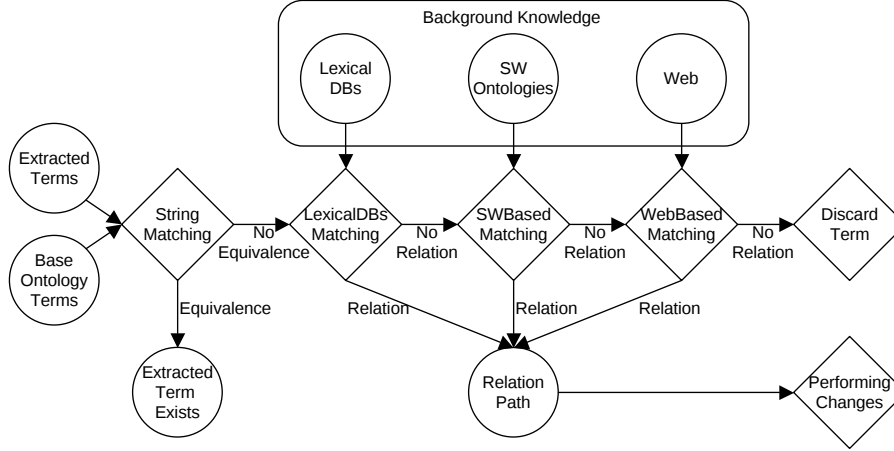


Fig. 2. Finding relations between new terms and the base ontology in Evolva.

been recently explored to support various tasks such as ontology matching [18] or enrichment [4]. While the initial results in employing these ontologies are encouraging, these techniques are still novel and in need of further optimizations (in particular regarding time-performance). Finally, the Web itself has been recognized as a vast source of information that can be exploited for relation discovery through the use of so-called lexico-syntactic patterns [6]. Because they rely on unstructured, textual sources, these techniques are more likely to introduce noise than the previously mentioned techniques which rely on already formalized knowledge. Additionally, these techniques are time consuming given that they operate at Web scale.

Taking into account these considerations, we devised a relation discovery process that combines various background knowledge sources with the goal of optimizing time-performance and precision. As shown in Figure 2, the relation discovery starts from quick methods that are likely to return good results, and continues with slower methods which are likely to introduce a higher percentage of noise: (1) The process begins with string matching for detecting already existing terms in the ontology. This will identify equivalence relations between the new terms and the ontology elements. (2) Extracted elements that do not exist in the base ontology are passed to a module that performs relation discovery by exploring WordNet’s synset hierarchy. (3) Terms that could not be incorporated by using WordNet are passed to the next module which explores Semantic Web ontologies. (4) If no relation is found, we resort to the slower and more noisy methods which explore the Web itself through search engines’ APIs and lexico-syntactic patterns [6]. In case no relation is found at the final level, the extracted term is discarded or, optionally, forwarded for manual check.

5 Implementation of Evolva’s Relation Discovery

We have partially implemented the algorithm presented in Figure 2 by making use of methods for exploring two main background knowledge sources: WordNet and online ontologies. We have not yet implemented methods for exploiting the Web as a source of knowledge. The first part of the implementation performs the string matching between the extracted terms and the ontology elements. We rely on the Jaro distance metric similarity [8] which takes into account the number and positions of the common characters between a term and an ontology concept label. This string similarity technique performs well on short strings, and offers a way to find a match between strings that are slightly different only because of typos or the use of different naming conventions.

The WordNet based relation discovery uses the Wu and Palmer similarity [21] for identifying the best similarity measure between the two terms. This measure is computed according to the following formula:

$$Sim(C1, C2) = \frac{2 * N3}{N1 + N2 + 2 * N3}$$

where C1 and C2 are the concepts to check for similarity, N1 is the number of nodes on the path from C1 to the least common superconcept (C3) of C1 and C2, N2 is the number of nodes between C2 and C3, and N3 is the number of nodes on the path from C3 to the root [21]. For those terms that are most closely related to each other, we derive a relation by exploring WordNet’s hierarchy using a functionality built into its Java library⁴. This will result in a relation between a term, as well as an inference path which lead to its discovery.

The terms that could not be related to the base ontology are forwarded to the next module which makes use of online ontologies. For this component, we rely on the Scarlet relation discovery engine⁵. It is worth to note that we handle ontologies at the level of statements, i.e. ontologies are not processed as one block of statements. Thus we focus on knowledge reuse without taking care of the validation of the sources as a whole with respect to the base ontology. Scarlet [18] automatically selects and explores online ontologies *to discover relations between two given concepts*. For example, when relating two concepts labeled *Researcher* and *AcademicStaff*, Scarlet 1) identifies (at run-time) online ontologies that can provide information about how these two concepts inter-relate and then 2) combines this information to infer their relation. [18] describes two increasingly sophisticated strategies to identify and to exploit online ontologies for relation discovery. Hereby, we rely on the first strategy that derives a relation between two concepts if this relation is defined within a single online ontology, e.g., stating that *Researcher* $\sqsubseteq$ *AcademicStaff*. Besides subsumption relations, Scarlet is also able to identify disjoint and named relations. All relations are obtained

⁴ <http://jwordnet.sourceforge.net/>

⁵ <http://scarlet.open.ac.uk/>

by using derivation rules which explore not only direct relations but also relations deduced by applying subsumption reasoning within a given ontology. For example, when matching two concepts labeled *Drinking Water* and *tap_water*, appropriate anchor terms are discovered in the TAP ontology and the following subsumption chain in the external ontology is used to deduce a subsumption relation: *DrinkingWater* $\sqsubseteq$ *FlatDrinkingWater* $\sqsubseteq$ *TapWater*. Note, that as in the case of WordNet, the derived relations are accompanied by a path of inferences that lead to them.

6 Relation Discovery Experiment

We performed an experimental evaluation of the current implementation of the relation discovery module on the data sets provided by the KMi scenario. Our goal was to answer three main questions. First, we wanted to get an insight into the efficiency, in particular in terms of precision, of the relation discovery relying on our two main background knowledge sources: WordNet and online ontologies. Second, we wished to understand the main reasons behind the incorrect relations, leading to ways of identifying these automatically. Tackling these issues would further increase the precision of the identified relations and bring us closer to a full automation of this task. Finally, as a preparation for implementing Evolva’s algorithm for performing ontology changes, we also wanted to identify a few typical cases of relations to integrate into the base ontology.

6.1 Experimental Data

We relied on 20 documents from KMi’s news repository as a source of potentially new information. We used Text2Onto’s extraction algorithm [7] and discovered 520 terms from these text documents.

The base ontology which we wish to evolve (i.e. KMi’s ontology) currently contains 256 concepts, although it has not been updated for well over one year, since April 2007. By using the Jaro matcher we identified that 21 of the extracted terms have exact correspondences within the base ontology and that 7 are closely related to some concepts (i.e. their similarity coefficient is above the threshold of 0.92).

6.2 Evaluation of the WordNet Based Relation Discovery

Out of the 492 remaining new terms, 162 have been related to concepts of the ontology thanks to the WordNet based relation discovery module. Some of these relations were duplicates as they related the same pair of term and concept through the relation of different synsets. For evaluation purposes, we eliminated duplicate relations and obtained 413 distinct relations (see examples in Table 1).

We evaluated a sample of randomly selected 205 relations (i.e. half of the total) in three parallel evaluations performed by three of the authors of this

Extracted Term	Ontology Concept	Relation	Relation Path
Contact	Person	$\sqsubseteq$	contact $\sqsubseteq$ representative $\sqsubseteq$ negotiator $\sqsubseteq$ communicator $\sqsubseteq$ person
Business	Partnership	$\sqsubseteq$	business $\sqsubseteq$ partnership
Child	Person	$\sqsubseteq$	child $\sqsubseteq$ person

Table 1. Examples of relations derived by using WordNet.

paper. This manual evaluation⁶, which is not part of our evolution framework, helped to identify those relations which we considered correct or false, as well as those for which we could not decide on a correctness value (“Don’t know”). Our results are shown in Table 2. We computed a precision value for each evaluator, however, because there was a considerable variation between these, we decided to also compute a precision value on the sample on which they all agreed. Even though, because of the rather high disagreement level between evaluators (more than 50%), we cannot draw a generally valid conclusion from these values. Nevertheless, they already give us an indication that, even in the worst case scenario, more than half of the obtained relations would be correct. Moreover, this experiment helped us to identify typical incorrect relations that could be filtered out automatically. These will be discussed in Section 6.4.

	Evaluator 1	Evaluator 2	Evaluator 3	Agreed by all
Correct	106	137	132	76
False	96	53	73	26
Don’t know	2	15	0	0
Precision	53 %	73 %	65 %	75 %

Table 2. Evaluation results for the relations derived from WordNet.

6.3 Evaluation Results for Scarlet

The Scarlet based relation discovery processed the 327 terms for which no relation has been found in WordNet. It identified 786 relations of different types (subsumption, disjointness, named relations) for 68 of these terms (see some examples in Table 3). Some of these relations were duplicates, as the same relation can often be derived from several online ontologies. Duplicate elimination lead to 478 distinct relations.

For the evaluation, we randomly selected 240 of the distinct relations (i.e. 50% of them). They were then evaluated in the same setting as the WordNet-based

⁶ To our knowledge, there are no benchmarks of similar experimental data against which our results could be compared.

relations. Our results are shown in Table 4, where, as in the case of the WordNet-based relations, precision values were computed both individually and for the jointly agreed relations. These values were in the same ranges as for WordNet. One particular issue we faced here was the evaluation of the named relations. These proved difficult because the names of the relations did not always make their meanings clear. Different evaluators provided different interpretations for these and thus increased the disagreement levels. Therefore, again, we cannot provide a definitive conclusion of the performance of this particular algorithm. Nevertheless, each evaluator identified more correct than incorrect relations.

6.4 Error Analysis

One of the main goals of our experiment was to identify typical errors and to envisage ways to avoid them. We hereby describe some of our observations.

As already mentioned, in addition to the actual relation discovered between a new term and an ontology concept, our method also provides the path that lead to this relation, either in the WordNet synsets hierarchy or in the external ontology in the case of Scarlet. Related to that, a straightforward observation was that there seem to be a correlation between the length of this path and the correctness of the relation, i.e. relations derived from longer paths are more likely to be incorrect. To verify this intuition, we inspected groups of relations with different path lengths and for each computed the percentage of correct, false, un-ranked relations, as well as the relations on which an agreement was not reached. These results are shown in Table 5. As expected, we observe that the percentage of correct relations decreases for relations with longer paths (although, a similar observation cannot be derived for the incorrect relations). We also note that the percentages of relations which were not ranked and of those on which no agreement was reached are higher for relations established through a longer path. This indicates that relations generated from longer paths are more difficult to interpret, and so, may be less suitable for automatic integration.

Another observation was that several relations were derived for the *Thing* concept (e.g. *Lecturer* $\sqsubseteq$ *Thing*). While these relations cannot be considered incorrect, they are of little relevance for the domain ontology, as they would not

No.	Extracted Term	Ontology Concept	Relation	Relation Path
1	Funding	Grant	$\sqsubseteq$	funding $\sqsubseteq$ grant
2	Region	Event	occurredIn	region $\sqsubseteq$ place $\leftarrow$ occurredIn- event
3	Hour	Duration	$\sqsubseteq$	hour $\sqsubseteq$ duration
4	Broker	Person	isOccupationOf	broker -isOccupationOf $\rightarrow$ person
5	Lecturer	Book	editor	lecturer $\sqsubseteq$ academicStaff $\sqsubseteq$ employee $\sqsubseteq$ person $\leftarrow$ editor-book
6	Innovation	Event	$\sqsubseteq$	innovation $\sqsubseteq$ activity $\sqsubseteq$ event

Table 3. Examples of relations discovered using Scarlet.

	Evaluator 1	Evaluator 2	Evaluator 3	Agreed by all
Correct	118	126	81	62
False	96	56	57	17
Don't know	11	47	102	8
Precision	56 %	70 %	59 %	79 %

Table 4. Evaluation results for the relations derived with Scarlet.

contribute in making it evolve in a useful way. Therefore, they should simply be discarded. Similarly, relations that contained in their path abstract concepts such as *Event*, *Individual* or *Resource* tended to be incorrect.

Relation Path Length	True	False	Don't Know	No agreement
1	33 %	10 %	0 %	58 %
2	26 %	8 %	4 %	64 %
3	30 %	5 %	5 %	63 %
4	23 %	10 %	2 %	67 %
5	20 %	3 %	9 %	69 %

Table 5. Correlation between the length of the path and the correctness of a relation.

Finally, during the evaluation we also identified a set of relations to concepts that are not relevant for the domain (e.g. death, doubt). While they sometimes lead to correct relations (e.g. $Death \sqsubseteq Event$), these were rather irrelevant for the domain and thus should be avoided. We concluded that it would be beneficial to include a filtering step that eliminates, prior to the relation discovery step, those terms which are less relevant for the base ontology.

6.5 Observations on Integrating Relations into the Base Ontology

A particularity of the use of Scarlet is that different relations are derived from different online ontologies, reflecting various perspectives and subscribing to different design decisions.

One side effect of exploring multiple knowledge sources is that the derived knowledge is sometimes redundant. Duplicates often appear when two or more ontologies state the same relation between two concepts. These are easy to eliminate for subsumption and disjoint relations, but become non-trivial for named relations.

Another side effect is that we can derive contradictory relations between the same pair of concepts originating from different ontologies. For example, between *Process* and *Event* we found three different relations: “disjoint”, “subClassOf”

and “superClassOf”. Such a case is a clear indication that at least one of the relations should be discarded, as they cannot be all integrated into the ontology.

As we mentioned previously, both our methods provide a relation as well as an inference path that lead to its derivation. This makes the integration with the base ontology easier as more information is available.

An interesting situation arises when a part of the path supporting the relation contradicts the base ontology. For example, the second relation in the path relating *Innovation* to *Event*, Row 6 of Table 3, contradicts the base ontology where *Event* and *Activity* are siblings. This is a nice illustration of how the base ontology can be used as a context for checking the validity of a relation. Indeed, we could envision a mechanism that increases the confidence value for those paths which have a high correlation with the ontology (i.e. when they “agree” at least on some parts).

In the process of matching a path to an ontology, we can encounter situations where some elements of the path only have a partial syntactic match with the labels of some ontology concepts. Referring to Row 5 of Table 3, some of the terms in the relation path connecting *Lecturer* to *Book* partially map to labels in the subsumption hierarchy of the base ontology:

$$\begin{aligned} \textit{LecturerInAcademia} &\sqsubseteq \textit{AcademicStaffMember} \sqsubseteq \\ &\textit{HigherEducationalOrganizationEmployee} \sqsubseteq \textit{EducationalEmployee} \sqsubseteq \\ &\textit{Employee} \sqsubseteq \textit{AffiliatedPerson} \sqsubseteq \textit{Person} \end{aligned}$$

While our Jaro based matcher could not identify a match between *Lecturer* and *LecturerInAcademia*, this association can be done by taking into account the discovered path and the base ontology, therefore avoiding the addition of already existing concepts, and giving further indications on the way to integrate the discovered relations.

A final interesting observation relates to the appropriate abstraction level where a named relation should be added. We listed in Row 5 of Table 3 a relation path where *Lecturer* inherits a named relation to *Book* from its superclass, *Person*. Because *Person* also exists in the base ontology, we think that it is more appropriate to add the relation to this concept rather than to the more specific concept.

7 Related Work

As mentioned in the introduction, the work presented in this paper is mostly related to those ontology evolution approaches which aim to automate the process of incorporating potentially novel information into a base ontology. In particular DINO [14, 15] makes use of ontology alignment and agent-negotiation techniques to achieve such integration. These are then validated by a human user. Similarly, Dynamo [17] uses an adaptive multi-agent system architecture to evolve ontologies. The system considers the extracted entities as individual agents related to other entities (agents) through a certain relationship. Unfortunately, Dynamo

only generates concept hierarchies and its output depends on the order of the input data fed in the system. Unlike these systems, Evolva explores various background knowledge sources to reduce the amount of human intervention required for ontology evolution.

Similarly to Evolva, background knowledge has been used to successfully support various other tasks. Ontology matching techniques have been built to exploit such sources: WordNet [9], medical domain ontologies [3] or online ontologies [18]. Online ontologies have also been used for supporting the enrichment of folksonomy tagspaces [4] and ontology learning [1].

8 Conclusion and Future Work

Ontology evolution is a tedious and time consuming task, especially at the level of introducing new knowledge to the ontology. Most of current ontology evolution approaches rely on the ontology curator's expertise to come up with the right integration decisions. We discussed in this paper how background knowledge can support Evolva, our ontology evolution framework, for automating the process of relation discovery. Our technique is based on gradually harvesting background knowledge sources in order to exploit the most specific and easily accessible sources first and later investigate more generic but noisier sources. In our experiments, we explored WordNet and Semantic Web ontologies (through the Scarlet relation discovery engine).

Our relation discovery experiments using WordNet and Scarlet helped identifying various possible ways in which the overall quality of the process can be improved. First, it became evident that relations established with abstract concepts or concepts that are poorly related to the base ontology have a low relevance. We could avoid deriving such relations in the first place by simply maintaining a list of common abstract concepts (e.g. *Thing*, *Resource*) that should be avoided. Another approach would be to check the relevance of the terms with respect to the ontology, by measuring their co-occurrence in Web documents with concepts from the base ontology.

In addition, the observation of the result of the relation discovery process can lead to the design of a number of heuristic-based methods to improve the general quality of the output. For example, we observed a correlation between the length of the path from which a relation was derived and its quality. Note, however, that more in-depth analysis is needed to verify such hypothesis.

Finally, and most interestingly, the base ontology itself can be used for validating the correctness of a relation. Indeed, an overlap between the statements in the path and the base ontology is an indication that the relation is likely to be correct, and, inversely, if contradictions exist between the path and the ontology, the relation should be discarded. We have also shown cases where during the integration of a path in the ontology, some concepts can be considered equivalent even if their labels only partially match at a syntactic level.

While the work presented in this paper has shown the feasibility of exploiting external background knowledge sources to automate, at least partially, the ontology evolution process, there are still a number of aspects that need to be considered to make this approach fully operational. First, the way different sources are combined (here, WordNet and Scarlet) should be better studied, so that the method could better benefit from the complementarity of local, well established knowledge bases and of dynamic, distributed and heterogeneous Semantic Web ontologies. While we use a linear combination of background knowledge sources in our experiment, it would be worth to test a parallel combination to assess its effect on precision. Moreover, we plan to enhance the matching and term anchoring process by introducing word sense disambiguation techniques [10] that should increase precision. Also, one element not considered in this paper concerns the computational performance of our approach to ontology evolution. Additional work is currently ongoing to optimize particularly complex components like Scarlet.

Acknowledgements

This work was funded by the NeOn project sponsored under EC grant number IST-FF6-027595.

References

1. H. Alani. Position Paper: Ontology Construction from Online Ontologies. In *Proc. of WWW*, 2006.
2. H. Alani, S. Harris, and B. O’Neil. Winnowing ontologies based on application use. *Proceedings of 3rd European Semantic Web Conference (ESWC-06)*, Montenegro, 2006.
3. Z. Aleksovski, M. Klein, W. ten Katen, and F. van Harmelen. Matching Unstructured Vocabularies using a Background Ontology. In S. Staab and V. Svatek, editors, *Proc. of EKAW*, LNAI. Springer-Verlag, 2006.
4. S. Angeletou, M. Sabou, L. Specia, and E. Motta. Bridging the Gap Between Folksonomies and the Semantic Web: An Experience Report. In *Proc. of the ESWC Workshop on Bridging the Gap between Semantic Web and Web 2.0*, 2007.
5. S. Bloehdorn, P. Haase, Y. Sure, and J. Voelker. *Ontology Evolution*, pages 51–70. John Wiley & Sons, June 2006.
6. P. Cimiano, S. Handschuh, and S. Staab. Towards the self-annotating web. *Proceedings of the 13th international conference on World Wide Web*, pages 462–471, 2004.
7. P. Cimiano and J. Völker. Text2Onto - a framework for ontology learning and data-driven change discovery. *Proceedings of the 10th International Conference on Applications of Natural Language to Information Systems (NLDB-05)*, Alicante, pages 15–17, 2005.
8. W. W. Cohen, P. Ravikumar, and S. E. Fienberg. A comparison of string distance metrics for name-matching tasks. *Proceedings of the IJCAI-2003 Workshop on Information Integration on the Web (IIWeb-03)*, 2003.

9. F. Giunchiglia, P. Shvaiko, and M. Yatskevich. Discovering Missing Background Knowledge in Ontology Matching. In *Proc. of ECAI*, 2006.
10. J. Gracia, V. Lopez, M. d'Aquin, M. Sabou, E. Motta, and E. Mena. Solving semantic ambiguity to improve semantic web based ontology matching. *Workshop on Ontology Matching, The 6th International Semantic Web Conference and the 2nd Asian Semantic Web Conference*, 2007.
11. P. Haase and L. Stojanovic. Consistent evolution of owl ontologies. *Proceedings of the Second European Semantic Web Conference, Heraklion, Greece, 2005*, pages 182–197, 2005.
12. M. Klein. *Change Management for Distributed Ontologies*. PhD thesis, Vrije Universiteit in Amsterdam, 2004.
13. Y. Lei, M. Sabou, V. Lopez, J. Zhu, V. Uren, and E. Motta. An infrastructure for acquiring high quality semantic metadata. *Proceedings of the 3rd European Semantic Web Conference*, 2006.
14. V. Novacek, L. Laera, and S. Handschuh. Semi-automatic integration of learned ontologies into a collaborative framework. *International Workshop on Ontology Dynamics (IWOD-07)*, 2007.
15. V. Novacek, L. Laera, S. Handschuh, J. Zemanek, M. Volkel, R. Bendaoud, M. R. Hacene, Y. Toussaint, B. Delecroix, and A. Napoli. D2. 3.8 v2 report and prototype of dynamics in the ontology lifecycle. 2008.
16. N. F. Noy, A. Chugh, W. Liu, and M. A. Musen. A framework for ontology evolution in collaborative environments. *Proceedings of the 5th Int. Semantic Web Conference (ISWC-06). Athens, Georgia, USA*, pages 544–558, 2006.
17. K. Ottens and P. Glize. A multi-agent system for building dynamic ontologies. *Proceedings of the 6th international joint conference on Autonomous agents and multiagent systems*, 2007.
18. M. Sabou, M. d'Aquin, and E. Motta. Exploring the Semantic Web as Background Knowledge for Ontology Matching. *Journal on Data Semantics*, XI, 2008.
19. L. Stojanovic. *Methods and Tools for Ontology Evolution*. PhD thesis, FZI - Research Center for Information Technologies at the University of Karlsruhe, 2004.
20. D. Vrandecic, H. S. Pinto, Y. Sure, and C. Tempich. The diligent knowledge processes. *Journal of Knowledge Management*, 9:85–96, 2005.
21. Z. Wu and M. Palmer. Verb semantics and lexical selection. *Proceedings of the 32nd Annual Meeting of the Association for Computational Linguistics*, pages 133–138, 1994.
22. F. Zablit. Dynamic ontology evolution. *To appear in: Proceedings of the ISWC Doctoral Consortium*, 2008.