

A large dataset for the evaluation of ontology matching

FAUSTO GIUNCHIGLIA¹, MIKALAI YATSKEVICH¹,
PAOLO AVESANI² and PAVEL SHIVAICO¹

¹*Department of Information Engineering and Computer Science (DISI), University of Trento, 38050 Povo, Trento, Italy;*

²*Fondazione Bruno Kessler, Via Sommarive 18, 38050 Povo, Trento, Italy; e-mail: fausto@disi.unitn.it, yatskevi@disi.unitn.it, avesani@fbk.eu, pavel@disi.unitn.it*

Abstract

Recently, the number of ontology matching techniques and systems has increased significantly. This makes the issue of their evaluation and comparison more severe. One of the challenges of the ontology matching evaluation is in building large-scale evaluation datasets. In fact, the number of possible correspondences between two ontologies grows quadratically with respect to the numbers of entities in these ontologies. This often makes the manual construction of the evaluation datasets demanding to the point of being infeasible for large-scale matching tasks. In this paper, we present an ontology matching evaluation dataset composed of thousands of matching tasks, called TaxME2. It was built semi-automatically out of the Google, Yahoo, and Looksmart web directories. We evaluated TaxME2 by exploiting the results of almost two-dozen of state-of-the-art ontology matching systems. The experiments indicate that the dataset possesses the desired key properties, namely it is *error-free*, *incremental*, *discriminative*, *monotonic*, and *hard* for the state-of-the-art ontology matching systems.

1 Introduction

Ontology matching is a critical operation in many applications, including ontology engineering, information integration, peer-to-peer information sharing, navigation, and query answering on the web (Euzenat & Shvaiko, 2007). It takes two graph-like structures as input, for instance, lightweight ontologies (Giunchiglia *et al.*, 2007) such as Google¹ and Yahoo² web directories, and produces, as output, an alignment that is a set of correspondences between the semantically related nodes of those graphs.

Many diverse solutions of matching have been proposed so far (see Rahm and Bernstein (2001), Kalfoglou and Schorlemmer (2003), Noy (2004), Doan and Halevy (2005), Shvaiko and Euzenat (2005) for recent surveys) while some examples of individual approaches addressing the matching problem can be found in Madhavan *et al.* (2001), Do and Rahm (2002), Bouquet *et al.* (2003), Doan *et al.* (2003), Noy and Musen (2003), Euzenat and Valtchev (2004), Ehrig *et al.* (2005), Gal *et al.* (2005), He and Chang (2006), Lambrix and Tan (2006), Qu *et al.* (2006), Tang *et al.* (2006), Giunchiglia *et al.* (2007), and Zhang and Bodenreider (2007)³. Finally, ontology matching has been given a book account in Euzenat and Shvaiko (2007). The rapid growth of various matching

¹ <http://www.google.com/Top/>

² <http://dir.yahoo.com/>

³ See <http://www.ontologymatching.org> for complete information on the topic.

approaches makes the issues of their evaluation and comparison more severe. In order to address these issues, the Ontology Alignment Evaluation Initiative (OAEI⁴), which is a coordinated international initiative that organizes the evaluation of the increasing number of ontology matching systems, was set up in 2005. The main goal of OAEI is to support the comparison of the systems and algorithms on the same basis, and to allow anyone to draw conclusions about the best matching strategies (Shvaiko *et al.*, 2007b).

One of the challenges of the ontology matching evaluation is how to build large-scale evaluation datasets; specifically, a large *set of reference correspondences* or *reference alignments* against which the results produced by ontology matching systems are to be compared. Notice that the number of possible correspondences grows quadratically with the number of entities to be compared. This often makes the manual construction of the reference correspondences demanding to the point of being infeasible for large-scale matching tasks.

The contributions of this paper are:

- a method for building semi-automatically large datasets enabling the assessment of quality results produced by ontology matching systems;
- the TaxME2 dataset composed of 4.639 matching tasks, which have been built out of the Google, Yahoo, and Looksmart web directories;
- an evaluation of the dataset within the OAEI campaigns of 2005, 2006 and 2007 with encouraging results, thus demonstrating, empirically, its strength.

This paper is an expanded and updated version of an earlier conference paper (Avesani *et al.*, 2005), which originally provided the TaxME method, the corresponding dataset, and its preliminary evaluation. The key limitation of TaxME was that it allowed for assessing the Recall indicator only, which is a completeness measure. The most important extensions of this paper over the earlier work by Avesani *et al.* (2005) include: (i) TaxME2, that is a new method and the corresponding dataset, which allows for assessing not only Recall, but also Precision, which is a correctness measure; (ii) the new property of the dataset, that is, monotonicity; (iii) extensive evaluation of the TaxME2 dataset within the OAEI campaigns of 2005, 2006 and 2007.

The empirical evaluation highlighted that the TaxME2 dataset possesses five key properties: *complexity*, namely that it is hard for state-of-the-art matching systems; *incrementality*, namely that it is effective in revealing weaknesses of the state-of-the-art matching systems; *discrimination capability*, namely that it discriminates sufficiently among the various matching solutions; *monotonicity*, namely that the matching quality measures calculated on the subsets of the dataset do not differ substantially from the measures calculated on the whole dataset; and *correctness*, namely that it can be considered as a correct tool to support the improvement and research on the matching solutions.

The rest of the paper is organized as follows. Section 2 provides a brief introduction to the problems of ontology matching and ontology matching evaluation. Section 3 discusses our approach to building semi-automatically the dataset for assessing Recall, called TaxME. Section 4 presents an extension of TaxME to TaxME2, which allows for an assessment of Precision, beside Recall. Section 5 introduces the five key properties of the dataset. Section 6 presents the results of our experiments and shows that TaxME2 possesses the desired properties. Section 7 overviews the related work. Finally, Section 8 summarizes the major findings of the paper and outlines future work.

2 Basics

In this section, we briefly introduce the basic concepts at work by discussing first the ontology matching problem (Section 2.1) and then the ontology matching evaluation problem (Section 2.2).

⁴ <http://oei.ontologymatching.org/>

2.1 The ontology matching problem

An ontology typically provides a vocabulary that describes a domain of interest and a specification of the meaning of terms used in the vocabulary. Depending on the precision of this specification, the notion of ontology encompasses several data and conceptual models, including classifications or web directories, database schemas, and fully axiomatized theories.

Given two ontologies, a *correspondence* is a 5-tuple:

$$\langle id, e_1, e_2, n, R \rangle$$

such that:

- id is a unique identifier of the given correspondence;
- e_1 and e_2 are entities (e.g., classes, properties) of the first and the second ontology, respectively;
- n is a confidence measure (typically in $[0, 1]$) holding for the correspondence between e_1 and e_2 ;
- R is a relation (e.g., equivalence ($=$), more general ($\sqsupseteq$)) holding between e_1 and e_2 .

The correspondence $\langle id, e_1, e_2, n, R \rangle$ asserts that the relation R holds between the ontology entities e_1 and e_2 with confidence n . The higher the confidence, the higher is the likelihood of the relation holding.

Matching is the process of finding correspondences between entities of different ontologies. *Alignment* is a set of correspondences between two (or more) ontologies. The alignment is the output of the matching process (Euzenat & Shvaiko, 2007).

We view ontologies as graph-like structures (Giunchiglia & Shvaiko, 2003). Let us exemplify the matching problem with the help of fragments of two tree-like structures, such as Google and Looksmart (Figure 1). Notice that in the general case, the relation holding between the nodes of these directories is not the specialization relation, but the so-called classification or parent-child relation (Giunchiglia *et al.*, 2007). Let us suppose that the task is to merge these directories. Such situations occur, for example, when one e-commerce company is to acquire another one. A first step here is usually to identify the matching candidates or correspondences. For example, Basketball in O1 can be found equivalent to Basketball in O2, while Games in O2 is more general than Board_Games in O1. Then, based on the obtained correspondences articulation axioms can be generated in order to create a new directory covering the matched directories.

As discussed above, heterogeneity can be reduced in two steps: (i) match entities of different ontologies to determine alignments and (ii) process the alignments according to the application needs. In this paper, we focus on the evaluation of quality results in the first step.

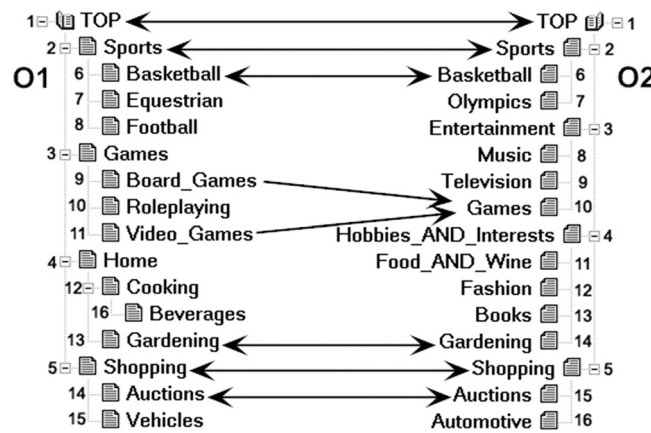


Figure 1 Fragments of Google and Looksmart web directories and some correspondences. The latter are expressed by arrows with single arrowheads standing for the more general relation ($\sqsupseteq$) and with double arrowheads standing for the equivalence relation ($=$)

2.2 The ontology matching evaluation problem

The ontology matching evaluation problem can be viewed as the problem of acquiring the reference correspondences that hold between entities of the ontologies. Given such reference correspondences, it would be straightforward to evaluate the quality of the results of a matching solution.

The commonly accepted measures for qualitative matching evaluation are based on the well-known information retrieval measures of relevance, such as Precision and Recall (van Rijsbergen, 1979). Let us consider Figure 2. The calculation of these measures is based on the comparison between the correspondences produced by a matching system (denoted S) and a complete set of reference correspondences (denoted H) considered to be correct. H is represented by the area inside the dotted circle. It is usually produced by humans. Finally, we denote the set of all possible correspondences, namely the cross product of the entities of two input ontologies, as M .

The correct correspondences found by a matching system are called the *true positives* (TP) and computed as follows:

$$TP = S \cap H \quad (1)$$

The incorrect correspondences found by a matching system are called the *false positives* (FP) and computed as follows:

$$FP = S - S \cap H \quad (2)$$

The correct correspondences missed by a matching system are called the *false negatives* (FN) and computed as follows:

$$FN = H - S \cap H \quad (3)$$

The incorrect correspondences not returned by a matching system are called the *true negatives* (TN) and computed as follows:

$$TN = M - S \cap H \quad (4)$$

We call the correspondences in H the *positive correspondences*, and the correspondences in N as defined in Equation (5), the *negative correspondences*.

$$N = M - H = TN + FP \quad (5)$$

Precision is a correctness measure. It varies in the $[0, 1]$ range; the higher the value, the smaller the set of wrong correspondences (false positives) which have been computed. It is calculated as follows:

$$Precision = \frac{|TP|}{|TP + FP|} = \frac{|H \cap S|}{|S|} \quad (6)$$

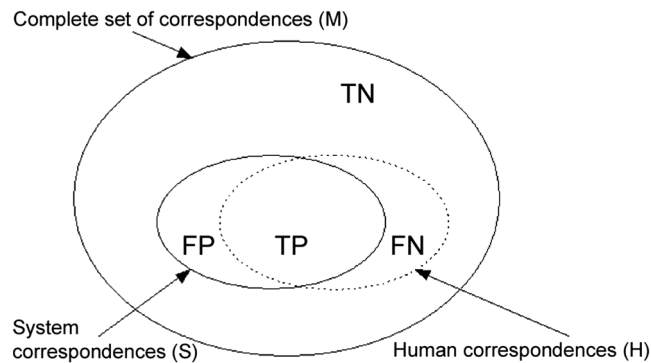


Figure 2 Basic sets of correspondences

Recall is a completeness measure. It varies in the $[0, 1]$ range; the higher the value, the smaller the set of correct correspondences (true positives) which have not been found. It is calculated as follows:

$$Recall = \frac{|TP|}{|TP + FN|} = \frac{|H \cap S|}{|H|} \quad (7)$$

Ontology matching systems are often not comparable based only on Precision or Recall. In fact, Recall can be maximized at the expense of Precision by returning all possible correspondences, that is, the cross product of the entities from two input ontologies. At the same time, higher Precision can be achieved at the expense of lower Recall by returning only few (correct) correspondences. Therefore, it is useful to consider both measures simultaneously or a combined measure, such as the F-measure.

In particular, F-measure is a global measure of the matching quality. It varies in the $[0, 1]$ range. It allows for comparison of the systems by their Precision and Recall at the point where their F-measure is maximal. Here, we use F-measure, which is a harmonic mean of Precision and Recall; that is, each of these measures is given equal importance. It is calculated as follows:

$$F\text{-measure} = \frac{2 \times Recall \times Precision}{Recall + Precision} \quad (8)$$

In order to calculate Precision, Recall, and F-measure, the complete reference alignment H must be known in advance. This opens up a problem of the reference alignments acquisition. The problem is that the construction of H is usually a manual process, which, in the case of matching, is quadratic with respect to the size of the ontologies to be matched. This manual process often turns out to be unfeasible for large datasets. For instance, each of the web directories, such as Google, Yahoo, and Looksmart, has entities in the order of 10^5 . This means that construction of H would require the manual evaluation of 10^{10} correspondences.

3 A dataset for evaluating Recall

In this section, we first outline the key idea behind the TaxME approach (Section 3.1). Then, we present the details of how the TaxME dataset has been built (Section 3.2).

3.1 The TaxME approach

We propose a semi-automatic method, called TaxME⁵, for an approximation of a reference alignment for tree-like structures, such as web directories. The key idea is to rely on a reference interpretation for entities (nodes), constructed by analyzing which documents have been classified in which nodes. The assumption is that the semantics of nodes can be derived from their pragmatics, namely from analyzing the documents that are classified under the given nodes. The working hypothesis is that the meaning of two nodes is equivalent if the sets of documents classified under those nodes have a meaningful overlap. The basic idea is, therefore, to compute the relationship measures based on the co-occurrence of documents.

Let us consider the example presented in Figure 3. Let N_1 be a node in the first ontology and N_2 be a node in the second ontology. D_1 and D_2 stand for the sets of documents classified under the nodes N_1 and N_2 , respectively. A_2 denotes the documents classified in the ancestor node of N_2 , and C_1 denotes the documents classified in the children nodes of N_1 .

⁵ The abbreviation TaxME stands for ‘TAXonomy Mapping Evaluation’, though in this paper, we use the terminology of Euzenat and Shvaiko (2007) and keep this abbreviation as originally introduced in Avesani *et al.* (2005) for historical reasons.

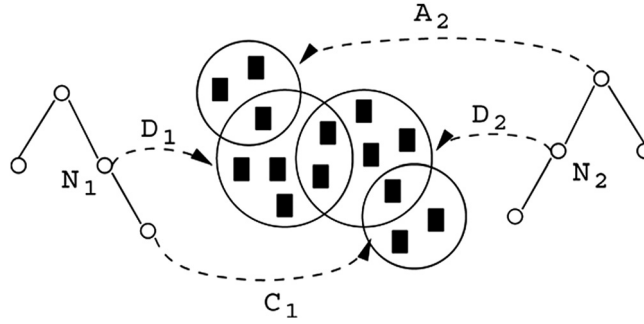


Figure 3 TaxME. Illustration of a document-driven similarity assessment

The equivalence measure, Eq , is defined as follows:

$$Eq(N_1, N_2) = \frac{|D_1 \cap D_2|}{|D_1 \cup D_2| - |D_1 \cap D_2|} \quad (9)$$

Notice that the range of $Eq(N_1, N_2)$ is $[0 \infty]$. The intuition is that the more D_1 and D_2 overlap, the bigger is the $Eq(N_1, N_2)$; with $Eq(N_1, N_2)$ becoming infinite when $D_1 = D_2$. $Eq(N_1, N_2)$ is normalized within $[0 1]$. The special case of $D_1 = D_2$ is approximated to 1.

Let us now discuss the generalization relation. Given two nodes N_1 and N_2 , and the related document sets D_1 and D_2 , we use two additional sets: (i) the set of documents classified in the ancestor node of N_2 , denoted as A_2 , and (ii) the set of documents classified in the children nodes of N_1 , denoted as C_1 . The generalization relationship holds when the first node has to be considered more general than the second node. Intuitively, this happens when the documents classified under the first node occur in the ancestor of the second node, or the documents classified under the second node occur in the subtree of the first node. Following this intuition, we can formalize the generalization measure, Mg (where Mg stands for more general), as follows:

$$Mg(N_1, N_2) = \frac{|(A_2 \cap D_1) \cup (C_1 \cap D_2)|}{|D_1 \cup D_2|} \quad (10)$$

The specialization relationship measure $Lg(N_1, N_2)$, where Lg stands for less general, can be easily formulated by exploiting the ‘symmetry’ of the problem. In particular, the first node is more specific than the second node when the meaning associated to the first node can be subsumed by the meaning of the second node. Intuitively, this happens when the documents classified under the first nodes occur in the subtree of the second node, or the documents classified under the second node occur in the ancestor of the first node.

The definitions above allow us to compute a relationship measure between two nodes of two different ontologies. Such a measure relies on the assumption that if two nodes classify the same set of documents, the meaning associated to the nodes is reasonably the same. Of course, this assumption is true for a virtually infinite set of documents. In a real world, we have to deal with a finite set of documents, and therefore, this way of proceeding is error-prone. Nevertheless, our claim is that the approximation introduced by our assumption is balanced by the benefit of scaling with the annotation of large directories.

3.2 The TaxME dataset

The reference alignment for the TaxME dataset is computed based on Google, Yahoo, and Looksmart. These web directories possess many useful properties: (i) they are widely known, (ii) they cover overlapping topics, (iii) they are heterogeneous, (iv) they incorporate typical real-world modeling and terminological errors, (v) they are large, and (vi) they address the same space of contents. Therefore, the working hypothesis of document co-occurrence is sustainable. Naturally, different web directories do not cover the same portion of the web but the overlap is meaningful.

Table 1 Number of nodes and documents processed during the TaxME construction process

	Google	Looksmart	Yahoo
Number of nodes	335.902	884.406	321.585
Number of URLs	2.425.215	8.498.157	872.410

The nodes are considered as categories denoted by lexical labels. The tree structures are considered as hierarchical relations. The URLs classified under a given node are taken to denote documents. Table 1 summarizes the total amount of the processed data.

Let us discuss the process used to build the TaxME reference alignment. In particular, it is organized in five steps.

- Step 1.** We crawled all three web directories, both the hierarchical structure, and the web contents. Then we computed the subset of URLs classified by all of them.
- Step 2.** We pruned the downloaded web directories by removing all the URLs that were not referred by all the three web directories.
- Step 3.** We performed an additional pruning by removing all the nodes with a number of URLs under a given threshold for obvious reasons. We used the threshold of 10.
- Step 4.** We manually recognized potential overlaps between two branches of two different web directories. Examples here include:

```

Google: Top > Science > Biology
Looksmart: Top > Science-and-Health > Biology

Yahoo: Top > Computers-and-Internet > Internet
Looksmart: Top > Computing > Internet

Google: Top > Reference > Education
Yahoo: Top > Education

```

We recognized 50 potential overlaps as exemplified above, and for each of them we ran an exhaustive assessment on all the possible pairs between the two related subtrees. Such a heuristic allowed us to reduce the search space and not consider all the possible correspondences coming from the cross product of the entities of two input directories. Specifically, we focused the analysis only on smaller subtrees where the overlaps were the most evident.

- Step 5.** For each of the subtree pairs selected, an exhaustive assessment of the correspondences holding between nodes was performed. This is done by exploiting the equivalence measure and the corresponding generalization and specialization measures. These are normalized within [0 1]. The final TaxME similarity measure is computed as the maximum out of these three measures, namely as follows:

$$Sim_{TaxME} = \max(Eq(N_1, N_2), Lg(N_1, N_2), Mg(N_1, N_2)) \quad (11)$$

We discarded all the pairs where none of the three relationships were detected. The distribution of the correspondences constructed using Sim_{TaxME} is shown in Figure 4.

Figure 4 indicates that the number of correspondences is stable and grows substantially, of two orders of magnitude, only with a value of the measure less than 0.1. As a pragmatic decision, the correspondences with Sim_{TaxME} above 0.5 were taken to constitute the reference alignment of TaxME. This process allowed us to obtain 2.265 pairwise relationships defined using the document-driven interpretation. Half are equivalence relationships and half are generalization relationships (notice that, by definition, the generalization and specialization measures are symmetric).

The final observation is that Sim_{TaxME} is robust. By robustness we mean here the fact that the number of incorrect correspondences is high only for very low values of Sim_{TaxME} and decreases very sharply as soon as these values increase. We need robustness as it shows dependencies

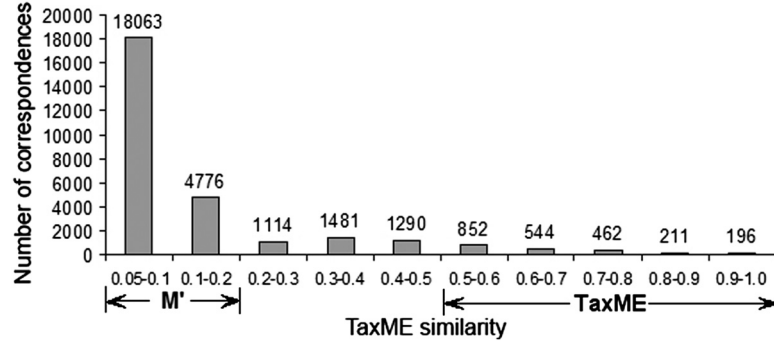


Figure 4 Distribution of the correspondences according to the TaxME similarity measure

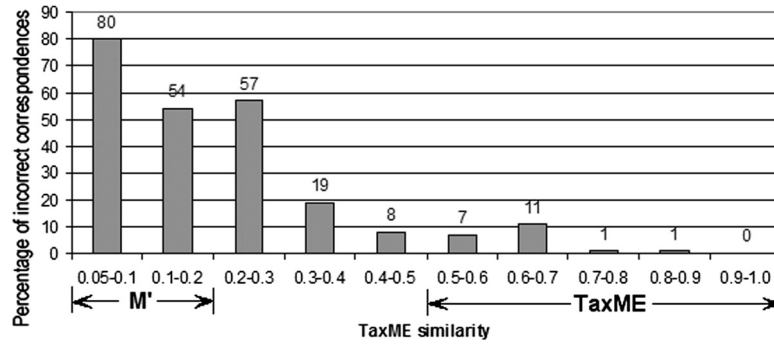


Figure 5 Distribution of the incorrect correspondences

between the values of Sim_{TaxME} and the human-observed similarity. We randomly selected 100 correspondences in nine intervals of range 0.1 and in one interval of range 0.05 and manually evaluated their correctness. This resulted in a reasonable amount of manual work (2 months by one person); we analyzed around one thousand correspondences. The results are presented in Figure 5 and show that Sim_{TaxME} is robust, namely:

- It is stable with a small percentage of the incorrect correspondences for the $[0.3 \ 1]$ range;
- The number of incorrect correspondences becomes substantial for very small values of Sim_{TaxME} , namely those less than 0.1.

4 A dataset for evaluating Precision

Following Equation (6), in order to evaluate Precision, we need to know the false positives (FP). This in turn, as from Figure 2, requires to know the reference alignment (H). However, computing H in the case of a large-scale matching task often requires an unfeasible human effort. We cannot use the reference alignment composed from positive correspondences, that is, $TaxME$. In this case, as shown in Figure 6, FP cannot be computed, because $FP_{unknown} = S \cap (H - TaxME)$, marked as gray area, is not known.

Our proposal here is to construct a reference alignment for the evaluation of both Recall and Precision, let us call it $TaxME2$, defined as follows:

$$TaxME2 = TaxME \cup N_{T2} \quad (12)$$

N_{T2} is an incomplete reference alignment containing *only* negative correspondences, that is, $N_{T2} \subset M - H$ (see Figure 6). $TaxME2$ must be a good representative of M . Thus, N_{T2} must be big

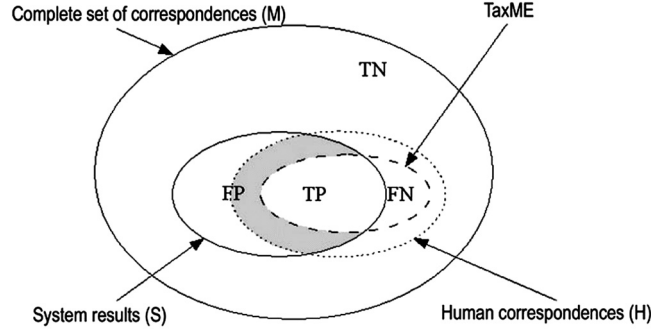


Figure 6 Alignment comparison using TaxME. *TP*, *FN* and *FP* stand for true positives, false negatives, and false positives, respectively

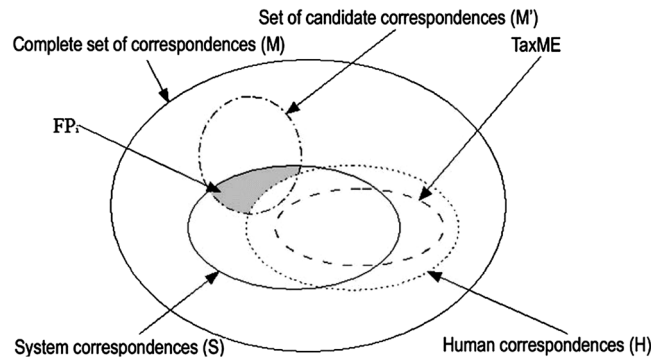


Figure 7 Alignment sets in TaxME2. Gray area stands for FP_i , which is a set of *FP* produced by the i th matching system on M'

enough in order to be the source of meaningful results. Therefore, we require N_{T2} to be at least the same size as TaxME, namely $|N_{T2}| \geq |\text{TaxME}|$.

N_{T2} is computed from the complete alignment set M in two macro steps. Let us first introduce them briefly and then discuss them in detail.

- *Step 1: Candidate correspondence selection.* The goal of this step is to select a set M' , such that $M' \subseteq M$, and which contains a big number of ‘hard’ negative correspondences.
- *Step 2: Negative correspondence selection.* The goal of this step is to filter all the positive correspondences from M' . In order to achieve this goal, M' is first pruned to the size that allows manual evaluation of the correspondences. Finally, the negative correspondences are manually selected from the remaining set of correspondences.

4.1 Candidate correspondence selection

The candidate set of correspondences M' is selected from M , see Figure 7. The goal of this step is to ensure that M' contains a *big number* of ‘hard’ negative correspondences. Intuitively a ‘hard’ negative correspondence is the correspondence with relatively high value of similarity measure, which is incorrect according to manual annotation. Given the robustness of $\text{Sim}_{\text{TaxME}}$, we have decided to exploit it as a similarity measure for M' construction. Let us revisit Figures 4 and 5. However, a big enough number of the negative correspondences can be obtained *only* for the values of $\text{Sim}_{\text{TaxME}}$ in the $[0 \ 0.2]$ range. As a pragmatic decision, we have selected M' as correspondences having $\text{Sim}_{\text{TaxME}}$ values in the $[0.05 \ 0.2]$ range. As from Figure 4, this allowed us to obtain $18.063 + 4.776 = 22.836$ candidate correspondences.

4.2 Negative correspondence selection

The negative correspondence selection step is devoted to the computation of N_{T2} . The process is organized in two phases as follows:

Matching system selection. The goal of this phase is to select a set of matching systems whose results are exploited for constructing N_{T2} . The set of the selected systems should be heterogeneous, that is, the selected systems should make different mistakes. Thus, the selected systems have to be the representatives of the different classes of the existing matching approaches. This also prevents N_{T2} from being biased towards a particular class of matching solutions.

As the result of this phase, based on the classifications of matching approaches in Giunchiglia and Shvaiko (2003) and Shvaiko and Euzenat (2005), we have selected three different matching systems, namely COMA (Do & Rahm, 2002), Similarity Flooding (SF) (Melnik *et al.*, 2002), and S-Match (SM) (Giunchiglia *et al.*, 2004); see also Section 7 for a brief comparison of these. Notice that these systems were used in the versions reported according to the above mentioned references.

Computation of negative correspondences. The goal of this phase is to compute N_{T2} , exploiting the results obtained by running the selected matching systems on M' . In particular, N_{T2} is computed based on the false positives as $N_{T2} = \cup_i FP_i$, where FP_i stands for the false positives produced by running the i th matching system on M' . The result of this exercise is depicted in Figure 7, where the gray area stands for FP_i . This construction schema ensures that N_{T2} will be hard for the existing systems, given that the set of matching systems evaluated on M' is representative and heterogeneous. An implicit constraint is that the number of false positives produced by each of the systems should be comparable. This prevents the existence of a bias towards a particular class of matching solutions. Notice that the computation of false positives requires the human annotation of the system results.

As the result of this phase, we have executed COMA, SF, and SM on M' . We manually evaluated the correspondences found by the systems and selected the false positives from them. Notice that we did not distinguish among different semantic relations while evaluating the matching quality. Therefore, for example, the correspondence $A \sqsubseteq B$ produced by SM and $A_1 = B_1$ produced by COMA were considered as true positives if $A = B$ and $A_1 \sqsubseteq B_1$ are so based on human judgment. Finally, we computed N_{T2} as the union of the false positives produced by the matching systems.

Table 2 provides a quantitative description of the content of N_{T2} . As from the first row of Table 2, the total number of annotated correspondences was $2.553 + 2.163 + 2.151 = 6.867$. Notice that this is six orders of magnitude less than the number of correspondences to be considered in the case of the complete reference alignment, which involves considering about 10^{10} correspondences (see Section 2.2). Notice also that the number of correspondences per system is balanced.

Figure 8 shows how the false positives produced by the systems are partitioned. In particular, there are no false positives found by SM, COMA, and SF, or even by SM and COMA together. There are small intersections between the false positives produced by SM and SF (0.1%) or by COMA and SF (2.3%). These results justify our assumption that all three systems belong to the distinct classes.

The final result is that N_{T2} consists of 2.374 correspondences. Notice that the size of N_{T2} is not equal to the sum of the false positives reported in the second row of Table 2, since, as from

Table 2 Total number of correspondences and the size of false positives as computed by COMA, SF, and SM on M'

	COMA	SF	SM
Found (S)	2.553	2.163	2.151
Incorrect (FP)	870	776	781

S = the number of correspondences produced by a matching system;
 FP = false positives.

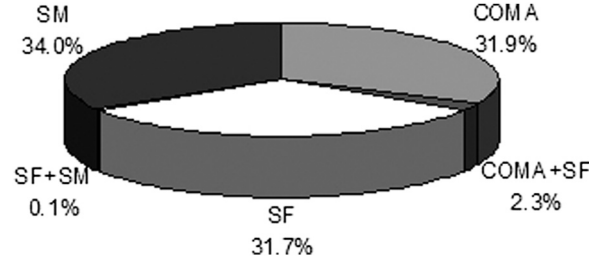


Figure 8 Partition of the false positives as computed by COMA, SF and SM on M'

Figure 8, there are intersections among these sets. The union of N_{T2} with TaxME resulted in the TaxME2 reference alignment. This, in turn, allowed for the evaluation of both Recall and Precision of $2.265 + 2.374 = 4.639$ correspondences.

5 The TaxME2 dataset properties

The dataset developed provides an incomplete reference alignment since it contains only part of the correspondences in H (see, for e.g., Figure 6). The key difference between Figures 2 and 6 is the fact that a complete reference alignment (the area inside the dotted circle in Figure 6) is simulated by exploiting an incomplete one (the area inside the dashed circle in Figure 6). However, if we assume that TaxME is a good representative of H , we can use the definitions of Recall and Precision as defined, respectively, in Equations (7) and (6) for their estimate. In order to ensure that this assumption holds, a set of requirements has to be satisfied:

1. *Correctness*, namely the fact that $\text{TaxME} \subset H$, modulo annotation errors.
2. *Complexity*, namely the fact that state-of-the-art matching systems experience difficulties when run on TaxME2.
3. *Incrementality*, namely the fact that TaxME2 allows for the incremental discovery of the weaknesses of the tested systems.
4. *Discrimination capability*, namely the fact that different sets of correspondences taken from TaxME2 are hard for the different systems.
5. *Monotonicity*, namely the fact that the matching quality measures, calculated on the subsets of the dataset, do not differ substantially from the measures calculated on the whole dataset.

In the next section, we argue that the above mentioned properties hold for the dataset developed. Notice that these five properties, in general, are essential (though not exhaustive, but good enough) for any plausible ontology matching evaluation dataset.

6 Evaluation

The evaluation was designed in order to assess the major dataset properties discussed previously. Overall, we exploit the results of around two dozen matching systems. Specifically, the results for the following seven matching systems: oMAP, CMS, Dublin20, Falcon, FOAM, OLA, and ctxMatch2, were taken from OAEI-2005 (see Ashpole *et al.* (2005) and Euzenat *et al.* (2005)). In turn, the results for the following seven matching systems: HMatch, Falcon, AUTOMS, RiMOM, OCM, COMA++, and Prior were taken from OAEI-2006 (see Euzenat *et al.* (2006) and Shvaiko *et al.* (2006)). The results for the following nine matching systems: Falcon, ASMOV, DSSim, Lily, OLA2, OntoDNA, Prior+, RiMOM, and X-SOM were taken from OAEI-2007 (see Euzenat *et al.* (2007) and Shvaiko *et al.* (2007a)). Only the Falcon system participated in all three evaluations, while the RiMOM and Prior teams participated only in 2006 and 2007; the OLA team, in turn, participated in 2005 and 2007. Finally, we also use the results of the matching systems exploited during the dataset construction process (see Section 4), namely COMA (Do & Rahm, 2002), SF (Melnik *et al.*, 2002), and SM (Giunchiglia *et al.*, 2004). For the systems, we used the default

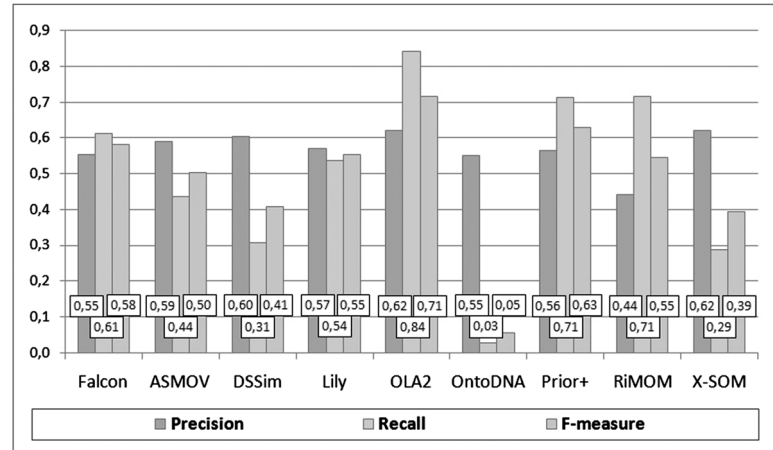


Figure 9 OAEI-2007: matching quality results

settings or, if applicable, the settings provided by the authors for the OAEI-2005, OAEI-2006, OAEI-2007 evaluations (Ashpole *et al.* (2005), Shvaiko *et al.* (2006) and Shvaiko *et al.* (2007a)).

In the rest of this section, we provide the evaluation of five key dataset properties, namely: correctness (Section 6.1), complexity (Section 6.2), incrementality (Section 6.3), discrimination capability (Section 6.4) and monotonicity (Section 6.5). For the sake of the presentation, we do not report the complete results⁶ but only for the selected systems, whose results we found the most interesting.

6.1 Correctness

We have manually analyzed correctness of the correspondences provided by TaxME2. Notice that part of the dataset for evaluating Precision is 100% correct by construction (modulo annotation errors). In turn, manually checking the part of the dataset for evaluating Recall revealed only around 3% of the incorrect correspondences. Taking into account the notion of idiosyncratic classification (Goren-Bar & Kuflik, 2005), namely the fact that human annotators on the sufficiently big and complex datasets tend to have resemblance of up to 20% in comparison with their own results, such a mismatch can be considered as marginal.

6.2 Complexity

Precision, Recall, and F-measure of the system results in OAEI-2007 are shown in Figure 9. The highest F-measure of 0.71 was demonstrated by the OLA2 system. The highest Precision of 0.62 was demonstrated by both the OLA2 system and X-SOM. The highest Recall of 0.84 was demonstrated by the OLA2 system. In turn, the average Precision, Recall, and F-measure were 0.57, 0.5, and 0.49, respectively.

Figure 9 indicates the complexity of TaxME2. The dataset is quite hard for the state-of-the-art matching systems. For example, the best F-measure in 2007, of 0.71, is significantly lower than the results demonstrated by the same systems on the other OAEI datasets, such as the benchmarks (Euzenat *et al.*, 2007), where the best F-measure was 0.97. The other interesting observation is that the systems exploited during the dataset construction process (e.g., the F-measure of COMA was 0.45) demonstrate comparable performance with the other systems that participated in the OAEI campaigns (see Figure 9).

⁶ The complete results can be found as follows:

OAEI-2005: <http://oaei.ontologymatching.org/2005/results/>

OAEI-2006: <http://oaei.ontologymatching.org/2006/results/>

OAEI-2007: <http://oaei.ontologymatching.org/2007/results/>

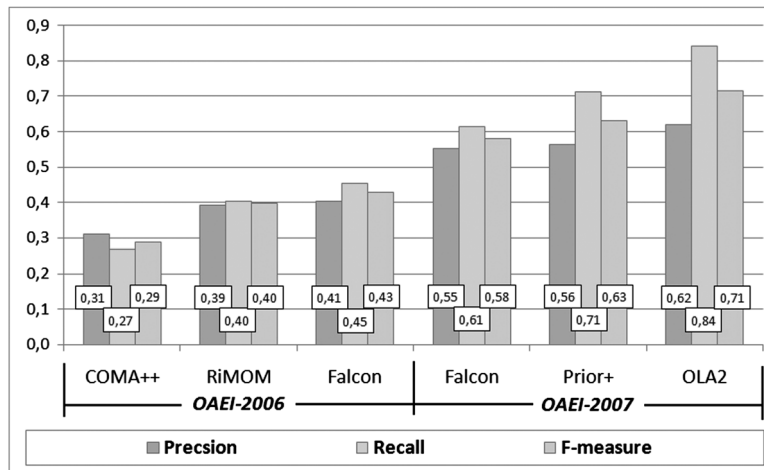


Figure 10 Comparison of matching quality results in 2006 and 2007

6.3 Incrementality

In 2007, the systems demonstrated substantially higher quality results than in previous two years. In particular, the average F-measure of the systems increased from approximately 29% in 2006 to 49% in 2007. The average Precision of the systems increased from approximately 35% in 2006 to 57% in 2007. The average Recall of the systems increased from approximately 22% in 2005 to 26% in 2006, and to 50% in 2007. Notice that in 2005 this dataset allowed for estimating only Recall (as described in Section 3), therefore in the above observations there are no values of Precision and F-measure for 2005.

A comparison of the results in 2006 and 2007 for the top three systems of each of the years based on the highest values of the F-measure indicator is shown in Figure 10. The key observation here is that quality of the best F-measure result of 2006 demonstrated by Falcon has almost doubled (increased by ~ 1.7 times) in 2007 by OLA2. The best Precision result of 2006 demonstrated by Falcon was increased by ~ 1.5 times in 2007 by both OLA2 and X-SOM. Finally, for what concerns Recall, the best result of 2005 demonstrated by OLA was increased by ~ 1.4 times in 2006 by Falcon and further increased by ~ 1.8 times in 2007 by OLA2. Thus, the OLA team managed to improve by ~ 2.6 times its Recall result of 2005 in 2007. Similar improvements were also demonstrated by the systems exploited during the dataset construction process, for example, Recall of SM was 0.29 in 2005 and improved by employing an iterative version of the semantic matching algorithm in 2006 (Giunchiglia *et al.*, 2006) up to 0.46.

The summative observation here is that the ontology matching community has made a substantial progress in 2007. Thus, the system designers managed to recognize major weaknesses in their approaches and provided plausible improvements. Notice that these improvements are generic, since the OAEI evaluation of this dataset was blind in 2005, 2006, and 2007, that is, participants did not know the reference alignment. Moreover, they had to run their systems in a fixed configuration for all the OAEI test cases they decided to address. As Figure 10 indicates, quality of the results is almost doubled from 2006 to 2007. This suggests that the systems experience fewer difficulties on the test case, although there still exists large room for further improvements.

6.4 Discrimination capability

Partitions of positive and negative correspondences according to the system results in OAEI-2007 are presented in Figures 11 and 12, respectively.

Figure 11 shows that the systems managed to discover all the positive correspondences (the ‘Nobody’ category is 0%). Only 15% of the positive correspondences were found by almost all (eight) matching systems. Figure 12 shows that almost all (eight) systems found 11% of the

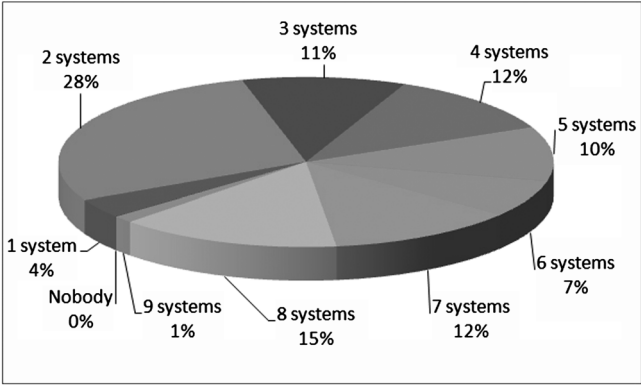


Figure 11 Partition of the system results on the positive correspondences

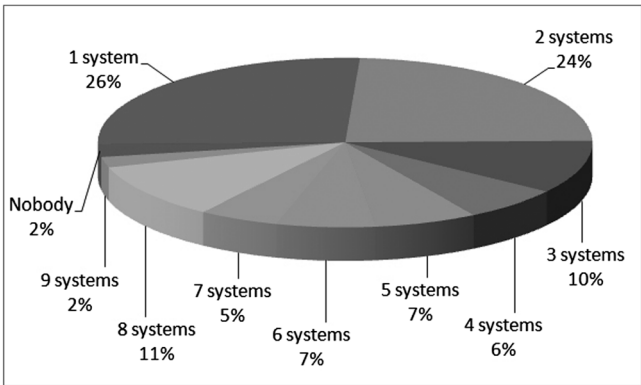


Figure 12 Partition of the system results on the negative correspondences

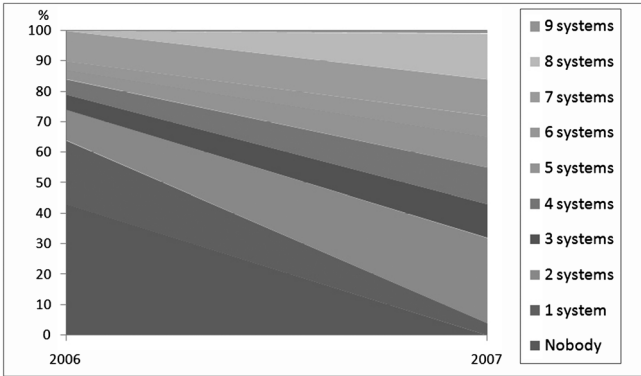


Figure 13 Comparison of partitions of the system results on the positive correspondences in 2006 and 2007

negative correspondences, that is, mistakenly returned them as positive. The last two observations suggest that the discrimination capability of the dataset is still high.

Let us now compare partitions of the system results in 2006 and 2007 on the positive and negative correspondences (see Figures 13 and 14, respectively).

Figure 13 shows that 43% of the positive correspondences have not been found by any of the matching systems in 2006, while in 2007 all the positive correspondences have been collectively found; see also how the selected regions (e.g., for two systems) consequently enlarge from 2006 to 2007.

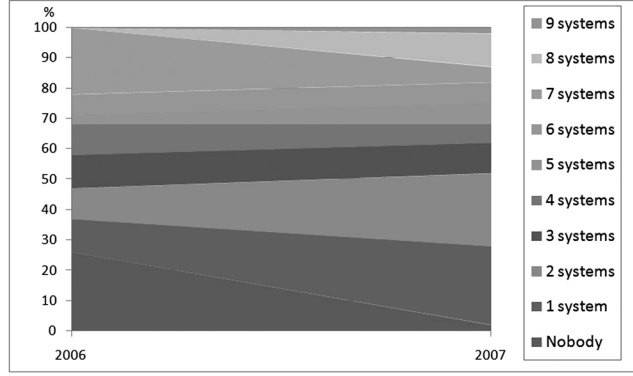


Figure 14 Comparison of partitions of the system results on the negative correspondences in 2006 and 2007

Figure 14 shows that in 2006, in general, the systems have correctly not returned 26% of the negative correspondences, while in 2007, this indicator decreased to 2%. In turn in 2006, 22% of the negative correspondences were mistakenly found by all (seven) the matching systems, while in 2007, this indicator decreased to 5%. An interpretation of these observations could be that systems keep trying various combinations of both ‘brave’ and ‘cautious’ strategies in discovering correspondences with a convergence towards better quality, since average Precision increased from 2006 to 2007.

Finally, as partitions of the positive and negative correspondences indicate the dataset retains good discrimination capability, that is, different sets of correspondences are still hard for the different systems.

6.5 Monotonicity

The dataset is said to demonstrate a monotonous behavior if the matching quality measures calculated on its subsets of gradually increasing size converge to the values obtained on the whole dataset. This property illustrates the fact that the dataset, as a whole, and its parts are not biased to the particular matching solution(s). It also shows how the gradual increase in the dataset size influence the results of the matching systems in terms of the matching quality, and gives a clue of whether a further increase in the dataset size may significantly influence the values of the matching quality measures.

In order to evaluate the monotonicity property, we randomly sampled 50, 100, 200, 500, and 1000 correspondences from TaxME2. For example, in order to obtain a 100 correspondences sample, 50 correspondences were randomly selected and added to the previously selected 50 correspondences sample. Then the matching quality measures for the matching systems were calculated on the samples of various sizes. An error was computed as follows:

$$Error_{Measure} = \frac{Measure_{sample} - Measure_{dataset}}{Measure_{dataset}} \quad (13)$$

$Measure_{sample}$ stands for a matching quality measure calculated on the sample. $Measure_{dataset}$ denotes a matching quality measure calculated on the whole dataset. Thus, for example, if a matching system had on TaxME2 dataset Precision of 0.2, and on the 100 correspondences sample randomly selected from the dataset a Precision of 0.21, the error of the system is as follows:

$$Error_{Precision} = \frac{|0.21 - 0.2|}{0.2} = 0.05$$

$Error_{Precision}$, $Error_{Recall}$, and $Error_{F-measure}$ for the various sample sizes averaged for 10 sample selections are summarized in Figures 15, 16, and 17, respectively. The results of the selected systems were taken from OAEI-2005 (Ashpole *et al.*, 2005; Euzenat *et al.*, 2005) and OAEI-2006 (Euzenat *et al.*, 2006; Shvaiko *et al.*, 2006).

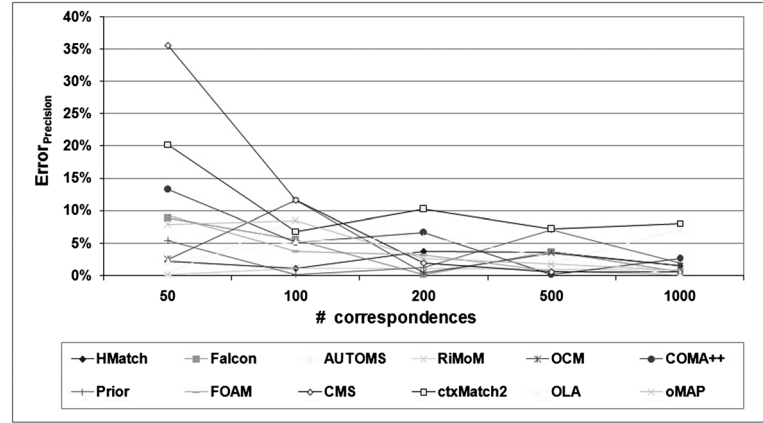


Figure 15 $Error_{Precision}$ depending on the sample size

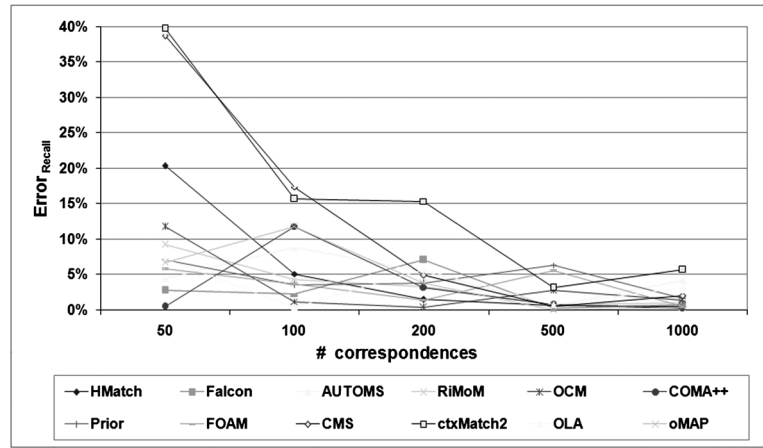


Figure 16 $Error_{Recall}$ depending on the sample size

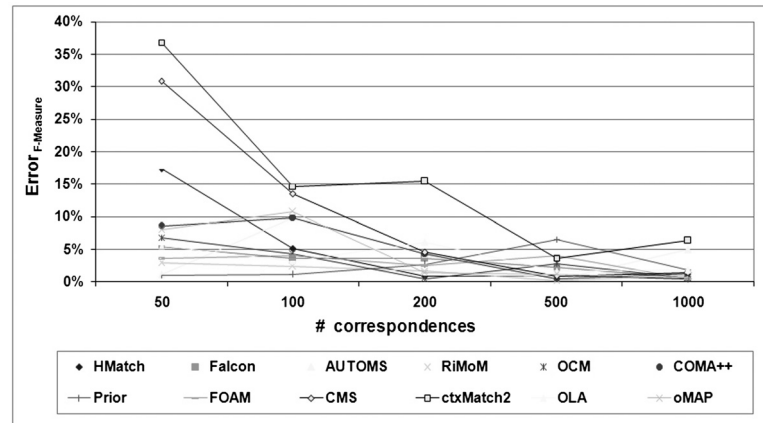


Figure 17 $Error_{F-measure}$ depending on the sample size

Notice that in Figures 15, 16, and 17, the errors drop very quickly as the sample size increases. Therefore, the matching quality measures obtained on the randomly selected samples of the various sizes converge quickly to their values on the whole dataset. In particular, given 500 correspondences sample, $Error_{Precision}$ and $Error_{Recall}$ is less than 10% what, given the results

depicted in Figure 9, corresponds to 0.02–0.05 difference in absolute values. Considering this difference as marginal we conclude that TaxME2 is monotonous.

7 Related work

In this section, we discuss the related work along the two dimensions, namely ontology matching systems (Section 7.1) and ontology matching evaluation (Section 7.2).

7.1 Ontology matching systems

There exists a line of semi-automated schema matching systems (see, for instance, Bergamaschi *et al.* (1999), Madhavan *et al.* (2001), Modica *et al.* (2001), Melnik *et al.* (2002), Kang and Naughton (2003), Dhamankar *et al.* (2004), Euzenat and Valtchev (2004), Gal *et al.* (2005), He and Chang (2006), Su *et al.* (2006), and Ziegler *et al.* (2006))⁷. A good survey and a classification of matching approaches up to 2001 is provided in Rahm and Bernstein (2001), an extension of its schema-based part and a user-centric classification of matching systems is provided in Shvaiko and Euzenat (2005), while the work in Euzenat and Shvaiko (2007) considers both Rahm and Bernstein (2001) and Shvaiko and Euzenat (2005) as well as some other classifications.

In particular, Shvaiko and Euzenat (2005) discuss how the systems can be distinguished in the matter of considering the correspondences and the matching task, thus representing the end-user perspective. In this respect, the following criteria were proposed: (i) alignments as solutions (these systems consider the matching problem as an optimization problem and the alignment is a solution to it, see, for instance, Similarity Flooding (Melnik *et al.*, 2002); (ii) alignments as theorems (these systems rely on semantics and require the alignments to satisfy it, see, for instance, S-Match (Giunchiglia *et al.*, 2004; Giunchiglia & Yatskevich, 2004; Giunchiglia *et al.*, 2005)); (iii) alignments as likeness clues (these systems produce only reasonable indications to a user for selecting the correspondences, see, for instance, COMA (Do and Rahm, 2002)). This justifies the choice of Similarity Flooding, S-Match, and COMA as representatives of three distinct classes of matching systems used for the TaxME2 dataset construction.

7.2 Ontology matching evaluation

Until very recently, there were no comparative evaluations and it was quite difficult to find two systems evaluated on the same dataset. Early evaluation efforts, such as in Sure *et al.* (2004), focused on artificially produced and quite simple examples rather than on real-world matching tasks. Also, many works (see, for e.g. Madhavan *et al.* (2001) and Do and Rahm (2002)) used for evaluation of the schemas with dozen of nodes, though from real-world applications, such as purchase orders. At the same time, industrial size schemas, such as UNSPSC⁸ and eCl@ss⁹, contain thousands of nodes.

In order to improve the performances of the ontology matching field, through the comparison of algorithms on various test cases, the OAEI¹⁰ has been established. The main goal of OAEI is to be able to compare systems and algorithms on the same basis and to allow anyone for drawing conclusions about the best matching strategies (Shvaiko *et al.*, 2007b). OAEI campaigns have been run in 2005 (Euzenat *et al.*, 2005), 2006 (Euzenat *et al.*, 2006), and 2007 (Euzenat *et al.*, 2007). For example, in OAEI-2007, seven various datasets were used and different evaluation modalities. The most similar and, to the best of our knowledge, unique effort in constructing semi-automatically large reference alignment was in the anatomy test case. It involved a fragment of the National

⁷ See also <http://www.ontologymatching.org/>

⁸ <http://www.unspsc.org>

⁹ <http://www.eclasse.de>

¹⁰ <http://oaei.ontologymatching.org/>

Cancer Institute (NCI) Thesaurus (3.304 classes) describing the human anatomy, published by the NCI¹¹, and the Adult Mouse Anatomical Dictionary¹² (2.744 classes), which has been developed as part of the Mouse Gene Expression Database project. An early work describing the construction of the reference alignment for these test cases can be found in Bodenreider *et al.* (2005), but the final results are still to appear.

Among the manually constructed datasets, it is worth mentioning the work by Ehrig and Sure (2004). This dataset is composed of several hundreds of reference correspondences and was used in the first evaluation event of 2004¹³, which was a predecessor of the OAEI campaigns. The real-world part of systematic tests of OAEI designed in Euzenat *et al.* (2005) contains dozen of them. The reference alignments in these datasets are composed of the positive correspondences, namely the correspondences that hold among the graph structures (for instance, *car* is equivalent to *auto*). All the other correspondences are assumed to be negative (for instance, *car* is not related to *tree*).

Let us now look at several works related to the ideas used during the TaxME2 dataset construction process. Similar to the annotated corpora for information retrieval or information extraction, we need to annotate a corpus of pairwise relationships. Of course, such an approach prevents the opportunity of having large corpora. The number of correspondences between two ontologies is quadratic with respect to the size of ontologies, which makes it hardly possible to manually acquire the reference alignments for large-scale ontologies. Certain heuristics, as used in TaxME2, can help in reducing the search space but the human effort is still too demanding. This approach has been followed by other researchers. For example, in Ichise *et al.* (2003), the interpretation of a node is approximated by a model computed through statistical learning. Of course the accuracy of the interpretation is affected by the error of the learning model. We follow a similar approach but without the statistical approximation. The working hypothesis is that the meaning of two nodes is equivalent if the sets of documents classified under those nodes have a meaningful overlap. Naturally, this can already be considered as an instance-based matching technique in itself. Notice that we use it only for the evaluation of schema-based systems and it cannot be used for the evaluation of instance-based systems, since the dataset does not provide instances.

The monotonicity principle for alignment evaluation was proposed in Gal *et al.* (2005). Notice that it differs significantly from the monotonicity property presented in this paper, since the former is concerned with aggregation of confidence measures of the correspondences whereas the latter applies to the matching quality measures, such as Precision and Recall, calculated for the samples of the dataset of various sizes.

Finally, it is worth noting that the ontology matching evaluation theme has been given a chapter account in Euzenat and Shvaiko (2007), where a more detailed discussion of the available evaluation datasets is given as well as of the other issues related to the evaluation of matching systems, including evaluation methodology (García-Castro & Gómez-Pérez, 2005), evaluation measures (Euzenat, 2007), etc.

8 Conclusions and future work

In this paper, we have presented a large-scale ontology matching evaluation dataset, called TaxME2, which was constructed out of the Google, Yahoo, and Looksmart web directories. The dataset is composed of 4.639 one-to-one matching tasks. It has the reference alignment. This allows for qualitative evaluation of matching systems by estimating both Precision and Recall. Around two-dozen state of the art matching solutions have been evaluated on the dataset during the OAEI campaigns of 2005, 2006, and 2007. The experimental results indicate that the dataset possesses the key properties of (i) correctness, (ii) complexity, (iii) incrementality, (iv) discrimination capability, and (v) monotonicity, thereby justifying the strength of the approach.

¹¹ <http://www.cancer.gov/cancerinfo/terminologyresources/>

¹² http://www.informatics.jax.org/searches/AMA_form.shtml

¹³ <http://www.atl.external.lmco.com/projects/ontology/i3con.html>

The five properties mentioned above are essential (though not exhaustive, but good enough) for any plausible ontology matching evaluation dataset. Specifically, they can be used as a ground for devising the notion of test hardness, which is one of the directions of our future work. Another direction of future work includes investigation of an evaluation dataset construction process in the case of expressive ontologies, such as those available in the biomedical domain. Finally, we are going to elaborate on further automation of the ontology matching evaluation dataset construction process in general. The ultimate goal in this direction is to minimize the human effort while increasing the dataset size.

Acknowledgements

We appreciate support from the Knowledge Web¹⁴ European Network of Excellence (IST-2004-507482) and the OpenKnowledge¹⁵ European STREP (FP6-027253). We are grateful to Jérôme Euzenat for many fruitful discussions on the topic of this paper.

References

- Ashpole, B., Ehrig, M., Euzenat, J. & Stuckenschmidt, H. (eds). 2005. *Proceedings of the Workshop on Integrating Ontologies at the International Conference on Knowledge Capture (K-CAP)*.
- Avesani, P., Giunchiglia, F. & Yatskevich, M. 2005. A large scale taxonomy mapping evaluation. In *Proceedings of the International Semantic Web Conference (ISWC)*, 67–81.
- Bergamaschi, S., Castano, S. & Vincini, M. 1999. Semantic integration of semistructured and structured data sources. *SIGMOD Record* **28**(1), 54–59.
- Bodenreider, O., Hayamizu, T. F., Ringwald, M., De Coronado, S. & Zhang, S. 2005. Of mice and men: aligning mouse and human anatomies. In *Proceedings of the American Medical Informatics Association (AIMA) Annual Symposium*, 61–65.
- Bouquet, P., Serafini, L. & Zanobini, S. 2003. Semantic coordination: a new approach and an application. In *Proceedings of the International Semantic Web Conference (ISWC)*, 130–145.
- Dhamankar, R., Lee, Y., Doan, A.-H., Halevy, A. & Domingos, P. 2004. iMAP: discovering complex semantic matches between database schemas. In *Proceedings of the International Conference on Management of Data (SIGMOD)*, 383–394.
- Do, H.-H. & Rahm, E. 2002. COMA—a system for flexible combination of schema matching approaches. In *Proceedings of the International Conference on Very Large Databases*, 610–621.
- Doan, A.-H. & Halevy, A. 2005. Semantic integration research in the database community: a brief survey. *AI Magazine*, special issue on semantic integration **26**(1), 83–94.
- Doan, A.-H., Madhavan, J., Dhamankar, R., Domingos, P. & Halevy, A. 2003. Learning to map ontologies on the semantic web. *The VLDB Journal* **12**(4), 303–319.
- Ehrig, M., Staab, S. & Sure, Y. 2005. Bootstrapping ontology alignment methods with APFEL. In *Proceedings of the International Semantic Web Conference (ISWC)*, 186–200.
- Ehrig, M. & Sure, Y. 2004. Ontology mapping—an integrated approach. In *Proceedings of the European Semantic Web Symposium (ESWS)*, 76–91.
- Euzenat, J. 2007. Semantic precision and recall for ontology alignment evaluation. In *Proceedings of the International Joint Conference on Artificial Intelligence (IJCAI)*, 348–353.
- Euzenat, J., Isaac, A., Meilicke, C., Shvaiko, P., Stuckenschmidt, H., Šváb, O., Svátek, V., van Hage, W. R. & Yatskevich, M. 2007. Results of the ontology alignment evaluation initiative 2007. In *Proceedings of the International Workshop on Ontology Matching (OM) at the International Semantic Web Conference (ISWC) + Asian Semantic Web Conference (ASWC)*.
- Euzenat, J. & Shvaiko, P. 2007. *Ontology Matching*. Springer-Verlag.
- Euzenat, J., Shvaiko, M. M. P., Stuckenschmidt, H., Šváb, O., Svátek, V., van Hage, W. R. & Yatskevich, M. 2006. Results of the ontology alignment evaluation initiative 2006. In *Proceedings of the International Workshop on Ontology Matching (OM) at the International Semantic Web Conference (ISWC)*.
- Euzenat, J., Stuckenschmidt, H. & Yatskevich, M. 2005. Introduction to the ontology alignment evaluation 2005. In *Proceedings of the Workshop on Integrating Ontologies at the International Conference on Knowledge Capture (K-CAP)*.

¹⁴ <http://www.knowledgeweb.semanticweb.org>

¹⁵ <http://openk.org>

- Euzenat, J. & Valtchev, P. 2004. Similarity-based ontology alignment in OWL-lite. In *Proceedings of the European Conference on Artificial Intelligence (ECAI)*, 333–337.
- Gal, A., Anaby-Tavor, A., Trombetta, A. & Montesi, D. 2005. A framework for modeling and evaluating automatic semantic reconciliation. *The VLDB Journal* **14**, 50–67.
- García-Castro, R. & Gómez-Pérez, A. 2005. Guidelines for benchmarking the performance of ontology management APIs. In *Proceedings of the International Semantic Web Conference (ISWC)*, 277–292.
- Giunchiglia, F., Marchese, M. & Zaihrayeu, I. 2007. Encoding classifications into lightweight ontologies. *Journal on Data Semantics* **VIII**, 57–81.
- Giunchiglia, F. & Shvaiko, P. 2003. Semantic matching. *The Knowledge Engineering Review* **18**(3), 265–280.
- Giunchiglia, F., Shvaiko, P. & Yatskevich, M. 2004. S-Match: an algorithm and an implementation of semantic matching. In *Proceedings of the European Semantic Web Symposium (ESWS)*, 61–75.
- Giunchiglia, F., Shvaiko, P. & Yatskevich, M. 2006. Discovering missing background knowledge in ontology matching. In *Proceedings of the European Conference on Artificial Intelligence (ECAI)*, 382–386.
- Giunchiglia, F. & Yatskevich, M. 2004. Element level semantic matching. In *Proceedings of the Meaning Coordination and Negotiation Workshop at the International Semantic Web Conference (ISWC)*.
- Giunchiglia, F., Yatskevich, M. & Giunchiglia, E. 2005. Efficient semantic matching. In *Proceedings of the European Semantic Web Conference (ESWC)*, 272–289.
- Giunchiglia, F., Yatskevich, M. & Shvaiko, P. 2007. Semantic matching: algorithms and implementation. *Journal on Data Semantics* **IX**, 1–38.
- Goren-Bar, D. & Kuflik, T. 2005. Supporting user-subjective categorization with self-organizing maps and learning vector quantization. *Journal of the American Society for Information Science and Technology* **56**(4), 345–355.
- He, B. & Chang, K. 2006. Automatic complex schema matching across web query interfaces: a correlation mining approach. *ACM Transactions on Database Systems* **31**(1), 1–45.
- Ichise, R., Takeda, H. & Honiden, S. 2003. Integrating multiple internet directories by instance-based learning. In *Proceedings of the International Joint Conference on Artificial Intelligence (IJCAI)*, 22–30.
- Kalfoglou, Y. & Schorlemmer, M. 2003. Ontology mapping: the state of the art. *The Knowledge Engineering Review* **18**(1), 1–31.
- Kang, J. & Naughton, J. 2003. On schema matching with opaque column names and data values. In *Proceedings of the International Conference on Management of Data (SIGMOD)*, 205–216.
- Lambrix, P. & Tan, H. 2006. SAMBO—a system for aligning and merging biomedical ontologies. *Journal of Web Semantics* **4**(3), 196–206.
- Madhavan, J., Bernstein, P. & Rahm, E. 2001. Generic schema matching using Cupid. In *Proceedings of the International Conference on Very Large Data Bases (VLDB)*, 48–58.
- Melnik, S., Garcia-Molina, H. & Rahm, E. 2002. Similarity flooding: a versatile graph matching algorithm. In *Proceedings of the International Conference on Data Engineering (ICDE)*, 117–128.
- Modica, G., Gal, A. & Jamil, H. 2001. The use of machine-generated ontologies in dynamic information seeking. In *Proceedings of the International Conference on Cooperative Information Systems (CoopIS)*, 433–448.
- Noy, N. 2004. Semantic integration: a survey of ontology-based approaches. *SIGMOD Record* **33**(4), 65–70.
- Noy, N. & Musen, M. 2003. The PROMPT suite: interactive tools for ontology merging and mapping. *International Journal of Human-Computer Studies* **59**(6), 983–1024.
- Qu, Y., Hu, W. & Chen, G. 2006. Constructing virtual documents for ontology matching. In *Proceedings of the World Wide Web Conference (WWW)*, 23–31.
- Rahm, E. & Bernstein, P. 2001. A survey of approaches to automatic schema matching. *The VLDB Journal* **10**(4), 334–350.
- van Rijsbergen, C. J. K. 1979. *Information Retrieval*, 2nd edn. Butterworths.
- Shvaiko, P. & Euzenat, J. 2005. A survey of schema-based matching approaches. *Journal on Data Semantics* **IV**, 146–171.
- Shvaiko, P., Euzenat, J., Giunchiglia, F. & He, B. (eds). 2007a. In *Proceedings of the International Workshop on Ontology Matching (OM) at the International Semantic Web Conference (ISWC) + Asian Semantic Web Conference (ASWC)*.
- Shvaiko, P., Euzenat, J., Noy, N., Stuckenschmidt, H., Benjamins, R. & Uschold, M. (eds). 2006. *Proceedings of the International Workshop on Ontology Matching (OM) at the International Semantic Web Conference (ISWC)*.
- Shvaiko, P., Euzenat, J., Stuckenschmidt, H., Mochol, M., Giunchiglia, F., Yatskevich, M., Avesani, P., van Hage, W. R., Švábo, O. & Svátek, V. 2007b. KnowledgeWeb Deliverable 2.2.9: description of alignment evaluation and benchmarking results. <http://exmo.inrialpes.fr/cooperation/kweb/heterogeneity/deli/kweb-229.pdf>
- Su, W., Wang, J. & Lochovsky, F. H. 2006. Holistic schema matching for web query interfaces. In *Proceedings of the International Conference on Extending Database Technology (EDBT)*, 77–94.

- Sure, Y., Corcho, O., Euzenat, J. & Hughes, T. (eds). 2004. *Proceedings of the Workshop on Evaluation of Ontology-based tools (EON) at the International Semantic Web Conference (ISWC)*.
- Tang, J., Li, J., Liang, B., Huang, X., Li, Y. & Wang, K. 2006. Using Bayesian decision for ontology mapping. *Journal of Web Semantics* **4**(1), 243–262.
- Zhang, S. & Bodenreider, O. 2007. Experience in aligning anatomical ontologies. *International Journal on Semantic Web and Information Systems* **3**(2), 1–26.
- Ziegler, P., Kiefer, C., Sturm, C., Dittrich, K. & Bernstein, A. 2006. Detecting similarities in ontologies with the SOQA-SimPack toolkit. In *Proceedings of the International Conference on Extending Database Technology (EDBT)*, 59–76.