

Linked Data Semantic Distance with Global Normalization for evaluating Semantic Similarity in a Taxonomy

Anna Formica* Francesco Taglino**

* National Research Council, Istituto di Analisi dei Sistemi ed Informatica “Antonio Ruberti”, Via dei Taurini 19, I-00185 Rome, Italy (e-mail: anna.formica@iasi.cnr.it)

** National Research Council, Istituto di Analisi dei Sistemi ed Informatica “Antonio Ruberti”, Via dei Taurini 19, I-00185 Rome, Italy (e-mail: francesco.taglino@iasi.cnr.it)

Abstract: In this work, the problem of evaluating semantic similarity in a taxonomy by relying on the notion of *information content* is investigated. In particular, a measure that takes into account not only the *generic sense* of a concept but also its *intended sense* in a given context is considered. Such a measure needs a semantic relatedness approach in order to evaluate the relatedness between the generic sense and the intended sense of a concept. In this work, we show that relying on the *Linked Data Semantic Distance with Global Normalization* leads to higher Spearman’s correlation values with human judgment with respect to the original proposal and previous experiments of the authors.

Keywords:

Semantic Similarity, Information Content, Taxonomy, Semantic Relatedness, Concept Sense, Linked Data Semantic Distance.

1. INTRODUCTION

The information content (IC) approach has been employed in several semantic similarity methods and has been recognized as “the state of the art on semantic similarity”, Batet and Sánchez (2020). It has higher correlation values with human judgment (*HJ*) with respect to other proposals which do not originate from it, Chandrasekaran and Mago (2022). However, it has shown some limitations, for instance, in the presence of pairs of concepts sharing the same “more general” meaning, Jeong et al. (2017).

In Formica and Taglino (2021), a novel approach has been presented that allows semantic similarity to be computed by taking into account not only the ICs of the concepts but also their *context*, i.e., the meanings of the concepts in the given application domain. Indeed, as shown in the mentioned paper, the context (or *perspective*, Lin (1998)) is fundamental in evaluating semantic similarity, and different contexts can lead to different similarity degrees among the same concepts.

It is important to note that also the approach proposed in Lin (1998) is based on the notion of perspective, but it does not allow the evaluation of similarity by addressing a single perspective at a time, and the information-theoretic definition of similarity between concepts is interpreted as “a weighted average of their similarities computed from different perspectives”.

In this work, analogously to Formica and Taglino (2021), we distinguish the notion of concept *generic sense*, i.e., the sense of the concept that is not related to any specific

context, from the concept *intended sense*, i.e., the meaning of the concept in a specific context. In order to compute semantic similarity between concepts, an important activity consists in evaluating the semantic relatedness, Taieb et al. (2020), between the generic sense and the intended sense of a given concept. Note that, in this paper, we address the problem of evaluating both semantic similarity and semantic relatedness by focusing on structured knowledge and, in particular, we rely on ISA taxonomies in the former case and knowledge graphs in the latter case. For this reason, Natural Language Processing (NLP) methods are not addressed, in line with Formica and Taglino (2023a).

In this work, in place of the *CombIC* measure, Schumacher and Ponzetto (2014), adopted in Formica and Taglino (2021), and with respect to the experimentation presented in Formica and Taglino (2023b), we rely on the family of methods proposed in Piao and Breslin (2016), here referred to as *Linked Data Semantic Distance with Global Normalization (LDSDGN)*, and in particular on the *LDSDGN_γ* strategy. Analogously to *CombIC*, this is a knowledge-based approach for computing semantic relatedness of concepts in *Resource Description Framework (RDF)* knowledge graphs.

Essentially, with respect to the *combIC* proposal that leverages *triple weights*, the *LDSDGN_γ* measure focuses on *triple patterns*, and aims at verifying the existence of specific configurations of paths in the RDF knowledge graph, Formica and Taglino (2023a).

In the new experiment presented in this paper, besides *combIC* and *LDSDGN_γ*, the comparison also involves the

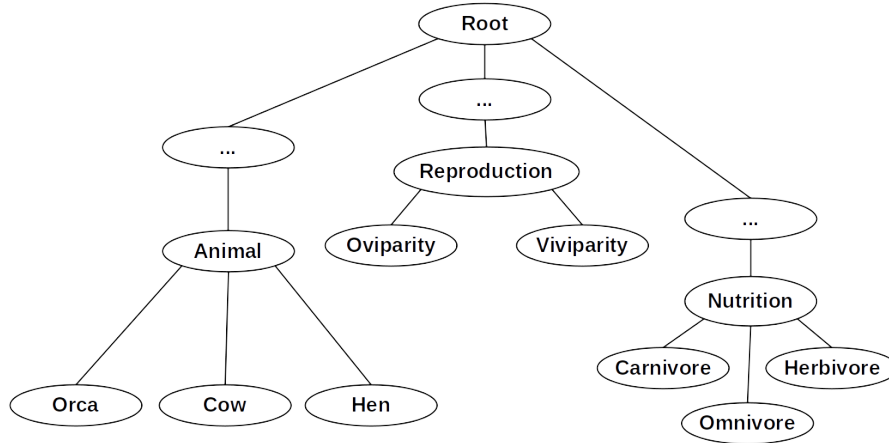


Fig. 1. A simple taxonomy including concept senses

following knowledge-based relatedness measures: $LDSD_{cw}$, which belongs to the *Linked Data Semantic Distance* ($LDSD$) family, Passant (2010), $ASRMP_m^a$, El Vaigh et al. (2020), *Wikipedia Link-based Measure* (WLM), Milne and Witten (2008), *Exclusivity-based measure* ($ExclM$), Hulpus et al. (2015), and *Linked Open Data Description Overlap* ($LODOverlap$), Zhou et al. (2011). Overall, according to the new experimental results, the $LDSDGN_\gamma$ strategy shows the best average Spearman's correlation.

The paper is organized as follows. In Section 2 the similarity measure relying on the intended senses of concepts is presented, and in Section 3 the $LDSDGN$ family of semantic relatedness measures is recalled. In Section 4 the new experiment is illustrated, followed by the related work in Section 5. The conclusion is given in Section 6.

2. THE SEMANTIC SIMILARITY MEASURE

Consider a set of concepts C of a tree-shaped ISA taxonomy (taxonomy for short), and a function p :

$$p: C \rightarrow [0,1]$$

such that, for any $c \in C$, $p(c)$ is the *probability* of the concept c computed on the basis of the relative concept frequency, $freq(c)$, evaluated from large collections of multidisciplinary texts, such as the Brown Corpus of American English. In particular, the probability of a concept c is defined as $p(c) = freq(c)/N$, where N is the total number of concepts in the corpus. Hence, the information content of a concept c , indicated as $IC(c)$, is computed as:

$$IC(c) = -\log p(c)$$

which means that, intuitively, as the probability of a concept increases its informativeness decreases. Given two concepts $c_i, c_j \in C$, the concept semantic similarity proposed in Lin (1998), $sim_L(c_i, c_j)$, is defined as follows:

$$sim_L(c_i, c_j) = \frac{2 \times IC(lcs(c_i, c_j))}{IC(c_i) + IC(c_j)} \quad (1)$$

where lcs is the *least common subsumer* (or the most informative subsumer in Resnik (1995)) of the concepts c_i, c_j in the taxonomy.

As extensively discussed in Formica and Taglino (2021), the approaches relying on the method mentioned above do not consider the semantic similarity of the meanings of the concepts according to a given context. For this reason, an enrichment of it has been proposed, by characterizing the meanings of the compared concepts with respect to a given application domain as recalled below.

Suppose sim is an information content-based similarity measure (e.g., sim_L), and D_k is an application domain, the semantic similarity of the concepts $c_i, c_j \in C$ in D_k , indicated as $sim^{D_k}(c_i, c_j)$, is defined as follows:

$$sim^{D_k}(c_i, c_j) = sim(c_i, c_j) \times (1 - \omega_k) + sim(\mathcal{S}_{D_k}(c_i), \mathcal{S}_{D_k}(c_j)) \times \omega_k \quad (2)$$

where ω_k is a weight, $0 \leq \omega_k \leq 1$, defined by the domain expert according to D_k , and $\mathcal{S}_{D_k}$ is a function, referred to as the *intended sense* function. Such a function associates a concept with its meaning according to D_k as follows:

$$\mathcal{S}_{D_k}: C \rightarrow C$$

and $\mathcal{S}_{D_k}(c) = s$ if $s \in C$ is the *intended sense* of c in D_k , $\mathcal{S}_{D_k}(c) = c$ otherwise.

Note that the weight ω_k , depending on D_k , allows a balance between the roles of the generic senses and the intended senses of the concepts, according to the relevance they have in the domain D_k .

For instance, consider the taxonomy shown in Figure 1, and an application domain, say D_1 , where it is important to characterize animals on the basis of their reproductive mode, which can be either *Viviparity* (i.e., the development of the embryo occurs inside the body of the mother), or *Oviparity* (i.e., the embryo grows inside an egg that is external to the body of the mother). Then, let $\mathcal{S}_{D_1}$ be a function associating a concept representing an animal (e.g., *Orca*, *Cow*, and *Hen*) with its intended sense in the domain D_1 , as follows:

$$\mathcal{S}_{D_1}(Orca) = Viviparity$$

$$\mathcal{S}_{D_1}(Cow) = Viviparity$$

$$\mathcal{S}_{D_1}(Hen) = Oviparity$$

Suppose to evaluate the similarity between *Orca* and *Cow*, and successively between *Orca* and *Hen* according to Eq. 2, by relying on the taxonomy in Figure 1. The *lcs* between *Orca* and *Cow* and between *Orca* and *Hen* coincide with the concept *Animal*, whereas with regard to the *lcs* of their intended senses, the following holds:

$$\begin{aligned} lcs(\mathcal{S}_{D_1}(Orca), \mathcal{S}_{D_1}(Cow)) &= lcs(Viviparity, Viviparity) \\ &= Viviparity \end{aligned}$$

$$\begin{aligned} lcs(\mathcal{S}_{D_1}(Orca), \mathcal{S}_{D_1}(Hen)) &= lcs(Viviparity, Oviparity) \\ &= Reproduction \end{aligned}$$

i.e., the *lcs* of the intended senses of *Orca* and *Cow*, which is *Viviparity*, is more informative than *Reproduction*, which is the *lcs* of the intended senses of *Orca* and *Hen*.

Therefore, assuming that siblings have the same information content, we expect that:

$$sim(Orca, Cow) > sim(Orca, Hen)$$

since *Orca* and *Cow* are both viviparous, whereas *Hen* is oviparous.

Now, suppose we have another scenario, say D_2 , requiring animals to be classified according to the kind of nutrition they follow. Then an animal can belong to one of the following classes, *Carnivore*, *Herbivore*, and *Omnivore*, and the function $\mathcal{S}_{D_2}$ is defined as follows:

$$\mathcal{S}_{D_2}(Orca) = Omnivore$$

$$\mathcal{S}_{D_2}(Cow) = Herbivore$$

$$\mathcal{S}_{D_2}(Hen) = Omnivore$$

In this second scenario, we expect the following:

$$sim(Orca, Cow) < sim(Orca, Hen)$$

because *Orca* and *Hen* are both omnivorous, whereas *Cow* is herbivorous. Indeed, according to the taxonomy shown in Figure 1:

$$\begin{aligned} lcs(\mathcal{S}_{D_2}(Orca), \mathcal{S}_{D_2}(Cow)) &= lcs(Omnivore, Herbivore) \\ &= Nutrition \end{aligned}$$

$$\begin{aligned} lcs(\mathcal{S}_{D_2}(Orca), \mathcal{S}_{D_2}(Hen)) &= lcs(Omnivore, Omnivore) \\ &= Omnivore \end{aligned}$$

and *Nutrition* is a concept less informative than *Omnivore* since it is its parent. This very simple example shows that different senses associated with concepts can lead to different similarity scores, subject to the definition of the weight ω_k according to D_k .

3. EVALUATING THE RELATEDNESS OF CONCEPT INTENDED SENSES

In the previous section we have seen that the weight ω_k depends on the domain D_k , and is established by domain experts on the basis of the roles of the generic senses and intended senses of the concepts in D_k . However in

our experiment, since finding several experts in different application domains is not feasible, in order to evaluate this proposal we relied on semantic relatedness methods, Taieb et al. (2020). In particular, we focus on knowledge-based semantic relatedness methods that depend on the availability of a formal knowledge base, such as a knowledge graph, an ontology, or a taxonomy.

In this paper, we consider the methods for evaluating semantic relatedness of resources in RDF¹ knowledge graphs, where concepts are represented by the nodes of the graph. RDF is a family of specifications designed as a standard model for data interchange on the Web. It is used for the conceptual description or modeling of information on Web resources, each identified by a Uniform Resource Identifier (URI). RDF is based upon the idea of making statements about resources by means of expressions in the form of triples following a *subject–predicate–object* pattern. The subject denotes the resource that is being described, and the predicate expresses a relation between the subject and the object, which can be a resource or a literal (e.g., a string, or a number).

Let $R = \{r_1, r_2, \dots, r_n\}$ be a finite set of URIs each representing a resource, and $L = \{l_1, l_2, \dots, l_m\}$ a finite set of literals, an RDF triple has the form:

$$\langle s, p, o \rangle$$

where $s \in R$ is the subject, $p \in R$ is the predicate, and $o \in R \cup L$ is the object. An RDF knowledge graph is a set of RDF triples, where subjects and objects are nodes, and predicates are directed arcs (also called links, edges or arrows).

3.1 LDSGD with Global Normalization (LDSGDGN)

In this paper we address the *Linked Data Semantic Distance with Global Normalization (LDSGDGN)* family of methods, Piao and Breslin (2016), which is an evolution of the *Linked Data Semantic Distance (LDSGD)* measures presented in Passant (2010). Such a family is based on the following assumptions.

Given two resources r_a and r_b , with $r_a \neq r_b$, a distance measure d should satisfy the following three axioms:

- (i) Equal self-distance, i.e., $d(r_a, r_a) = d(r_b, r_b) = 0$.
- (ii) Symmetry, i.e., $d(r_a, r_b) = d(r_b, r_a)$.
- (iii) Minimality, i.e., $d(r_a, r_a) < d(r_a, r_b)$.

Since the measures introduced by Passant do not satisfy these three axioms, in order to meet them, in Piao and Breslin (2016) the authors propose the *LDSGDGN* family containing the *LDSGDGN_α*, *LDSGDGN_β*, and *LDSGDGN_γ* distances, that are recalled below.

LDSGDGN_α. Given a RDF knowledge graph $\mathcal{G}$, let C_d be a function that computes the number of direct links between two resources in the graph as follows.

Given two resources r_a, r_b , and the predicate p_j , then $C_d(p_j, r_a, r_b) = 1$ if in the graph there exists a triple $\langle r_a, p_j, r_b \rangle$, otherwise $C_d(p_j, r_a, r_b) = 0$.

¹ <http://www.w3.org/TR/2014/REC-rdf11-concepts-20140225/>

Furthermore, the total number of resources that can be reached from r_a by means of the predicate p_j is indicated by $C_d(p_j, r_a)$ (analogously, $C_d(p_j, r_b)$). Similarly, $C_{io}(p_j, r_a)$ and $C_{ii}(p_j, r_a)$ are the total number of resources indirectly linked to r_a via outgoing and incoming links labeled with the predicate p_j , respectively.

On the basis of the assumption that resources are more related if there is a great number of them linked via a given predicate p_k , the C_{io} (C_{ii}) function defined by Passant has been generalized by using the function C'_{io} (C'_{ii}) as follows: $C'_{io}(p_k, r_a, r_b)$ ($C'_{ii}(p_k, r_a, r_b)$) computes the total number of resources linked to r_a and r_b via an outgoing (incoming) predicate p_k .

Therefore, given the resources r_a, r_b , the first distance is the $LDS\mathcal{D}GN_\alpha(r_a, r_b)$ measure defined in Eq. 3:

$$LDS\mathcal{D}_\alpha(r_a, r_b) = \frac{1}{1 + f_1 + f_2} \quad (3)$$

where:

$$f_1 = \sum_{p_j \in U} \frac{C_d(p_j, r_a, r_b)}{1 + \log(C_d(p_j, r_a))} + \sum_{p_j \in V} \frac{C_d(p_j, r_b, r_a)}{1 + \log(C_d(p_j, r_b))}$$

$$f_2 = \sum_{p_j \in W} \frac{C'_{io}(p_j, r_a, r_b)}{1 + \log(C_{io}(p_j, r_a))} + \sum_{p_j \in Z} \frac{C'_{ii}(p_j, r_a, r_b)}{1 + \log(C_{ii}(p_j, r_a))}$$

and $U, V, W, Z \subseteq R$ are the sets of the predicates p_j in the graph $\mathcal{G}$ such that $C_d(p_j, r_a, r_b) = 1$, $C_d(p_j, r_b, r_a) = 1$, $C'_{io}(p_j, r_a, r_b) > 0$, and $C'_{ii}(p_j, r_a, r_b) > 0$, respectively.

As an example, consider the resources r_a , and r_b in the graph $\mathcal{G}$ of Figure 2. In order to apply the $LDS\mathcal{D}GN_\alpha$ measure to this pair, the sets U, V, W , and Z in Eq. 3 are defined as follows.

$U = \{p_1\}$ because $C_d(p_j, r_a, r_b) = 1$ only if $j = 1$, and $C_d(p_1, r_a) = 2$ since r_b and r_4 are all the resources that can be reached from r_a through the predicate p_1 . $W = \{p_2\}$ because $C'_{io}(p_j, r_a, r_b) = 1$ only for $j = 2$, and $C_{io}(p_2, r_a) = 2$ since r_b and r_2 are indirectly linked to r_a via outgoing links labeled with p_2 , thanks to the following pairs of triples:

$$(\langle r_a, p_2, r_1 \rangle, \langle r_b, p_2, r_1 \rangle);$$

$$(\langle r_a, p_2, r_1 \rangle, \langle r_2, p_2, r_1 \rangle).$$

Finally, V and Z are both equal to the empty set.

$LDS\mathcal{D}GN_\beta$. With respect to $LDS\mathcal{D}GN_\alpha$, in the following $LDS\mathcal{D}GN_\beta$ measure the last two addenda at the denominator are modified by addressing the averages between $C_{io}(p_j, r_a)$, $C_{io}(p_j, r_b)$, and $C_{ii}(p_j, r_a)$, $C_{ii}(p_j, r_b)$ respectively, in order to achieve symmetry.

Given the resources r_a , and r_b , $LDS\mathcal{D}GN_\beta(r_a, r_b)$ is defined according to the following Eq. 4:

$$LDS\mathcal{D}_\beta(r_a, r_b) = \frac{1}{1 + f_1 + f_2} \quad (4)$$

where:

$$f_1 = \sum_{p_j \in U} \frac{C_d(p_j, r_a, r_b)}{1 + \log(C_d(p_j, r_a))} + \sum_{p_j \in V} \frac{C_d(p_j, r_b, r_a)}{1 + \log(C_d(p_j, r_b))}$$

$$f_2 = \sum_{p_j \in W} \frac{C'_{io}(p_j, r_a, r_b)}{1 + \log(\frac{C_{io}(p_j, r_a) + C_{io}(p_j, r_b)}{2})}$$

$$+ \sum_{p_j \in Z} \frac{C'_{ii}(p_j, r_a, r_b)}{1 + \log(\frac{C_{ii}(p_j, r_a) + C_{ii}(p_j, r_b)}{2})}$$

and the sets U, V, W , and Z are the same sets considered by the $LDS\mathcal{D}GN_\alpha$ measure.

For example, when applying the $LDS\mathcal{D}GN_\beta$ measure to the resources r_a , and r_b in Figure 2, the further function to be computed with respect to $LDS\mathcal{D}GN_\alpha$ is $C_{io}(p_2, r_b)$, which is equal to 3. Indeed, the resources r_a, r_2 , and r_9 are the only resources that can be reached from r_b via outgoing links labelled with the predicate p_2 , thanks to the following pairs of triples:

$$(\langle r_b, p_2, r_1 \rangle, \langle r_a, p_2, r_1 \rangle);$$

$$(\langle r_b, p_2, r_1 \rangle, \langle r_2, p_2, r_1 \rangle);$$

$$(\langle r_b, p_2, r_8 \rangle, \langle r_9, p_2, r_8 \rangle).$$

$LDS\mathcal{D}GN_\gamma$. In Piao and Breslin (2016), the authors propose a further measure, namely $LDS\mathcal{D}GN_\gamma$, relying on a *global normalization* notion that essentially considers the importance of a path between two resources according to the number of its occurrences in the whole graph $\mathcal{G}$. In the following, let $C_{dp}(p_j)$ be the global occurrences of the link p_j between two resources in $\mathcal{G}$.

Furthermore, $C_{io}(p_k, r_j, r_a, r_b) = 1$ if there exists a resource r_j such that $\langle r_a, p_k, r_j \rangle$ and $\langle r_b, p_k, r_j \rangle$, and $C_{ii}(p_k, r_j, r_a, r_b) = 1$ if there exists a resource r_j such that $\langle r_j, p_k, r_a \rangle$ and $\langle r_j, p_k, r_b \rangle$. The normalizations of $C_{io}(p_k, r_j, r_a, r_b)$ and $C_{ii}(p_k, r_j, r_a, r_b)$ are carried out by using the global occurrences $C_{iop}(p_k, r_j)$ and $C_{iip}(p_k, r_j)$ of r_j as follows. $C_{iop}(p_k, r_j)$ returns the global occurrences of r_j in the undirected paths $[\langle r_n, p_k, r_j \rangle, \langle r_s, p_k, r_j \rangle]$, for any resources r_n, r_s in the graph $\mathcal{G}$ and, analogously, $C_{iip}(p_k, r_j)$ computes the global occurrences of r_j in the undirected paths $[\langle r_j, p_k, r_n \rangle, \langle r_j, p_k, r_s \rangle]$, for any resources r_n, r_s in $\mathcal{G}$.

According to the above assumptions, given r_a and r_b , $LDS\mathcal{D}GN_\gamma(r_a, r_b)$ is defined in Eq. 5:

$$LDS\mathcal{D}_\gamma(r_a, r_b) = \frac{1}{1 + f_1 + f_2} \quad (5)$$

where:

$$f_1 = \sum_{p_j \in U} \frac{C_d(p_j, r_a, r_b)}{1 + \log(C_{dp}(p_j))} + \sum_{p_j \in V} \frac{C_d(p_j, r_b, r_a)}{1 + \log(C_{dp}(p_j))}$$

$$f_2 = \sum_{(p_k, r_j) \in W} \frac{C_{io}(p_k, r_j, r_a, r_b)}{1 + \log(C_{iop}(p_k, r_j))} + \sum_{(p_k, r_j) \in Z} \frac{C_{ii}(p_k, r_j, r_a, r_b)}{1 + \log(C_{iip}(p_k, r_j))}$$

and $U \subseteq R$ is the set of the predicates p_j in the graph $\mathcal{G}$ such that $C_d(p_j, r_a, r_b) = 1$, $V \subseteq R$ is the set of the predicates p_j in $\mathcal{G}$ such that $C_d(p_j, r_b, r_a) = 1$, $W \subseteq R \times R$ is the set of pairs (p_k, r_j) such that $C_{io}(p_k, r_j, r_a, r_b) = 1$, and $Z \subseteq R \times R$ is the set of pairs (p_k, r_j) such that $C_{ii}(p_k, r_j, r_a, r_b) = 1$.

When applying the $LDS\mathcal{D}GN_\gamma$ measure to r_a and r_b in the graph in Figure 2, we have the following.

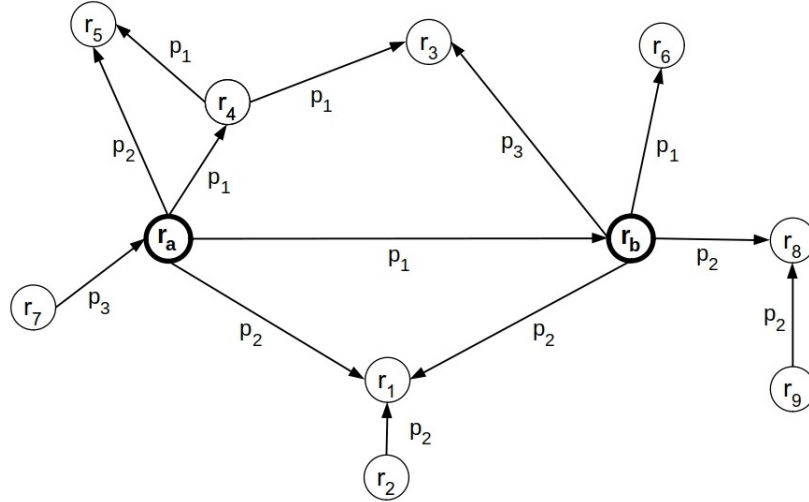


Fig. 2. An example of the knowledge graph $\mathcal{G}$

$U = \{p_1\}$, and V is the empty set. Then, we have $C_{dp}(p_1) = 5$ because of the presence of the following triples $\langle r_a, p_1, r_b \rangle$, $\langle r_a, p_1, r_4 \rangle$, $\langle r_b, p_1, r_6 \rangle$, $\langle r_4, p_1, r_5 \rangle$, and $\langle r_4, p_1, r_3 \rangle$. $W = \{(p_2, r_1)\}$ because $C_{io}(p_k, r_j, r_a, r_b) = 1$ only when $k = 2$, and $j = 1$, whereas Z is the empty set.

Then we have $C_{iop}(p_2, r_1) = 3$ because of the following three pairs of triples:

$$\begin{aligned} &(\langle r_a, p_2, r_1 \rangle, \langle r_b, p_2, r_1 \rangle); \\ &(\langle r_a, p_2, r_1 \rangle, \langle r_2, p_2, r_1 \rangle); \\ &(\langle r_b, p_2, r_1 \rangle, \langle r_2, p_2, r_1 \rangle). \end{aligned}$$

According to the results presented in Piao and Breslin (2016), $LDSDGN_\gamma$ performs better than the $LDSDGN_\alpha$ and $LDSDGN_\beta$ measures, and this is the measure we have considered in our experiment.

4. EXPERIMENTAL RESULTS

In line with Formica and Taglino (2021), the measure addressed in this paper relies on a novel approach for which the experimentation requires, besides the dataset composed of a set of pairs of concepts, further pairs of concepts representing the concept senses. Therefore, in order to compare the new experimental results with the ones of the original proposal, the *M&C* dataset, Miller and Charles (1991), has been addressed and, for each pair of concepts of this dataset, all the pairs of concepts of the same dataset have been considered as possible contexts.

With regard to the methods, the same six information content-based approaches discussed in Formica and Taglino (2021) have been analyzed, that are sim_R , Resnik (1995), sim_L , Lin (1998), $sim_{J\&C}$, Jiang and Conrath (1997), $sim_{P\&S}$, Pirrò (2009), sim_A , Adhikari et al. (2018), and $sim_{A\&M}$, Adhikari et al. (2016). The $sim_{W\&P}$, Wu and Palmer (1994), has also been addressed, as representative of the edge-counting approach, Rada et al. (1989).

Consider the *M&C* dataset, and the same 28 pairs of concepts belonging to it in order to associate each pair of the dataset with 28 possible application domains D_k , $k = 1 \dots 28$, in the following referred to as contexts (therefore

we have $28 \times 28 = 784$ similarity scores). For instance, for the pair of concepts (*boy*, *lad*), the 28 contexts are:

$$\mathcal{S}_{D_1}(\text{boy}) = \text{car}$$

$$\mathcal{S}_{D_1}(\text{lad}) = \text{automobile}$$

$$\mathcal{S}_{D_2}(\text{boy}) = \text{gem}$$

$$\mathcal{S}_{D_2}(\text{lad}) = \text{gewel}$$

...

$$\mathcal{S}_{D_{28}}(\text{boy}) = \text{rooster}$$

$$\mathcal{S}_{D_{28}}(\text{lad}) = \text{voyage}.$$

As mentioned in Section 2, in general, the intended senses of concepts are supposed to be estimated by domain experts, together with the related weight ω_k in the given context D_k . In the experiment presented in Formica and Taglino (2021), in order to quantify such a weight, which represents the relevance of a pair of senses with respect to the pair of contrasted concepts, we used the method proposed in Schuhmacher and Ponzetto (2014) whereas, as already mentioned, in this experiment we rely on the $LDSDGN_\gamma$ measure. In particular, given a pair of concepts c_i , c_j and a context D_k , we assume that ω_k is defined as follows:

$$\omega_k = (r_1 + r_2)/2 \quad (6)$$

where $r_1 = \text{rel}(c_i, \mathcal{S}_{D_k}(c_i))$ and $r_2 = \text{rel}(c_j, \mathcal{S}_{D_k}(c_j))$, and rel is the relatedness degree computed according to the best strategy presented in Piao and Breslin (2016).

It is important to recall that in order to compute the 28 tables, one for each pair of the *M&C* dataset, each table containing 28 possible contexts for that pair, a disambiguation step has to be performed. In fact, it is well-known that in Wikipedia, and consequently in DBpedia, terms are addressed with the possible meanings they have, i.e., a term is associated with multiple senses due to, for instance, semantic change or shift, Armaselu et al. (2022), Montariol et al. (2021). For this reason, in the experiment, the disambiguation is necessary in order to address senses

Table 1. Average Spearman's correlations in the 28 contexts according to the experiment presented in Formica and Taglino (2021) based on *CombIC*

$concept_1, concept_2$	sim_R	$sim_{W\&P}$	sim_L	$sim_{J\&C}$	$sim_{P\&S}$	sim_A	$sim_{A\&M}$
car, automobile	0.81	0.80	0.84	0.85	0.85	0.84	0.86
gem, jewel	0.80	0.77	0.83	0.82	0.84	0.87	0.87
journey, voyage	0.89	0.87	0.90	0.86	0.92	0.92	0.92
boy, lad	0.85	0.81	0.91	0.85	0.92	0.85	0.85
coast, shore	0.94	0.85	0.93	0.90	0.92	0.93	0.91
asylum, madhouse	0.83	0.89	0.91	0.85	0.93	0.87	0.85
magician, wizard	0.92	0.86	0.89	0.90	0.91	0.87	0.87
midday, noon	0.87	0.89	0.91	0.91	0.92	0.87	0.87
furnace, stove	0.65	0.38	0.63	0.50	0.67	0.66	0.66
food, fruit	0.85	0.58	0.84	0.81	0.89	0.83	0.81
bird, cock	0.84	0.82	0.88	0.82	0.86	0.85	0.84
bird, crane	0.83	0.83	0.88	0.85	0.87	0.86	0.85
tool, implement	0.80	0.78	0.84	0.77	0.84	0.82	0.81
brother, monk	0.82	0.71	0.83	0.83	0.93	0.82	0.81
crane, implement	0.69	0.67	0.73	0.74	0.79	0.73	0.75
lad, brother	0.71	0.68	0.80	0.66	0.69	0.75	0.73
journey, car	0.69	0.58	0.71	0.78	0.77	0.71	0.72
monk, oracle	0.77	0.59	0.65	0.67	0.70	0.76	0.75
food, rooster	0.85	0.64	0.85	0.86	0.89	0.86	0.87
coast, hill	0.62	0.72	0.70	0.69	0.65	0.70	0.71
forest, graveyard	0.81	0.79	0.84	0.84	0.88	0.83	0.83
monk, slave	0.70	0.59	0.62	0.62	0.67	0.63	0.63
coast, forest	0.82	0.78	0.86	0.88	0.92	0.86	0.82
lad, wizard	0.64	0.59	0.64	0.54	0.72	0.59	0.59
chord, smile	0.58	0.76	0.78	0.99	0.99	0.76	0.76
glass, magician	0.85	0.85	0.85	0.91	0.89	0.81	0.77
noon, string	0.92	0.88	0.92	0.94	0.86	0.94	0.93
rooster, voyage	0.79	0.73	0.82	0.86	0.90	0.83	0.82
<i>Avg Correl.</i>	0.79	0.74	0.81	0.80	0.84	0.81	0.80

in line with the *HJ* evaluation in the *M&C* experiment. For instance, *crane* in Wikipedia has two main senses, that are *bird* and *machine*, therefore when paired for instance with *implement*, it is disambiguated by using the sense *machine*.

Note that in the experiment, in order to address significant contexts, the ones associated with weight $\omega_k = 0$, for a given k , or both the values $r_1, r_2 = 0$ have been ignored, and a threshold for *HJ* equal to 0.5 (on a scale from 0 to 4) has been introduced.

In Table 1, the average correlations for all the 28 pairs according to Spearman are given, where the relatedness degrees have been computed by leveraging the semantic relatedness measure presented in Schuhmacher and Ponzetto (2014). In Table 2 the same values have been obtained by relying on the best semantic relatedness strategy proposed in Piao and Breslin (2016).

In Table 3 the above average Spearman's correlations are compared to the ones obtained according to the following methods: *LDSD_{cw}*, Passant (2010), *ASRMP_m^a*, El Vaigh et al. (2020), *WLM*, Milne and Witten (2008), *ExclM*, Hulpus et al. (2015), and *LODO_{Overlap}*, Zhou et al. (2011).

Before recalling all the mentioned approaches, below an example is given in order to clarify how the values of the above tables have been obtained.

Consider the pair of concepts (*journey, voyage*), the Resnik's similarity sim_R , and the context D_5 represented by the pair (*coast, shore*), corresponding to the fifth row

of the Table 1 (or Table 2). Therefore, in order to evaluate Eq. 2, we have:

$$c_i = \textit{journey},$$

$$c_j = \textit{voyage},$$

$$sim = sim_R,$$

$$\mathcal{S}_{D_5}(\textit{journey}) = \textit{coast},$$

$$\mathcal{S}_{D_5}(\textit{voyage}) = \textit{shore}.$$

Now, we focus on the weight ω_5 of the context D_5 , which represents the relevance of the pair of senses (*coast, shore*) with respect to the pair of contrasted concepts (*journey, voyage*). In Table 1 the *CombIC* relatedness measure presented in Schuhmacher and Ponzetto (2014) has been used, whereas in Table 2 the values have been obtained by relying on the *LDSDGN _{γ}* function defined in Eq. 5 presented in Piao and Breslin (2016). Therefore, in Table 2, in order to quantify ω_5 in Eq. 6, we have:

$$r_1 = LDSDGN_{\gamma}(\textit{journey}, \mathcal{S}_{D_5}(\textit{journey})) =$$

$$LDSDGN_{\gamma}(\textit{journey}, \textit{coast}))$$

$$r_2 = LDSDGN_{\gamma}(\textit{voyage}, \mathcal{S}_{D_5}(\textit{voyage})) =$$

$$LDSDGN_{\gamma}(\textit{voyage}, \textit{shore}))$$

As a result, for instance, in Table 2 the value 0.83 in correspondence to the row (*journey, voyage*) and the column sim_R has been obtained by computing the average of the

Table 2. Average Spearman's correlations in the 28 contexts according to the new experiment based on $LDSDGN_\gamma$

$concept_1, concept_2$	sim_R	$sim_{W\&P}$	sim_L	$sim_{J\&C}$	$sim_{P\&S}$	sim_A	$sim_{A\&M}$
car,automobile	0.83	0.82	0.84	0.83	0.88	0.82	0.82
gem,jewel	0.83	0.82	0.84	0.83	0.88	0.82	0.82
journey,voyage	0.83	0.81	0.83	0.82	0.88	0.82	0.82
boy,lad	0.83	0.80	0.83	0.82	0.88	0.82	0.82
coast,shore	0.83	0.80	0.83	0.82	0.88	0.82	0.81
asylum,madhouse	0.83	0.81	0.83	0.82	0.88	0.82	0.82
magician,wizard	0.83	0.81	0.84	0.83	0.88	0.82	0.82
midday,noon	0.83	0.81	0.84	0.83	0.88	0.82	0.82
furnace,stove	0.83	0.77	0.82	0.82	0.88	0.82	0.81
food,fruit	0.80	0.76	0.83	0.82	0.88	0.81	0.81
bird,cock	0.83	0.80	0.83	0.82	0.88	0.82	0.81
bird,crane	0.83	0.79	0.83	0.82	0.88	0.82	0.81
tool,implement	0.83	0.80	0.83	0.82	0.88	0.82	0.81
brother,monk	0.82	0.78	0.81	0.82	0.88	0.82	0.82
crane,implement	0.83	0.79	0.83	0.82	0.88	0.82	0.81
lad,brother	0.82	0.79	0.83	0.82	0.87	0.82	0.81
journey,car	0.78	0.75	0.80	0.82	0.87	0.81	0.80
monk,oracle	0.82	0.77	0.81	0.82	0.87	0.82	0.81
food,rooster	0.79	0.79	0.81	0.82	0.87	0.81	0.80
coast,hill	0.83	0.79	0.83	0.82	0.88	0.82	0.81
forest,graveyard	0.78	0.75	0.80	0.82	0.87	0.81	0.81
monk,slave	0.82	0.79	0.82	0.82	0.88	0.82	0.81
coast,forest	0.80	0.76	0.81	0.82	0.87	0.81	0.81
lad,wizard	0.82	0.79	0.83	0.82	0.87	0.82	0.81
chord,smile	0.83	0.78	0.83	0.82	0.87	0.82	0.82
glass,magician	0.80	0.76	0.81	0.82	0.87	0.82	0.82
noon,string	0.79	0.75	0.80	0.82	0.87	0.81	0.81
rooster,voyage	0.79	0.75	0.80	0.82	0.87	0.81	0.81
Avg Correl.	0.82	0.78	0.82	0.82	0.87	0.82	0.81

Table 3. Averages of the average Spearman's correlations with different relatedness methods

	sim_R	$sim_{W\&P}$	sim_L	$sim_{J\&C}$	$sim_{P\&S}$	sim_A	$sim_{A\&M}$	Avg.
<i>CombIC</i>	0.79	0.74	0.81	0.80	0.84	0.81	0.80	0.80
<i>LDSD_{cw}</i>	0.77	0.72	0.80	0.80	0.85	0.81	0.80	0.79
<i>ASRMP_m^a</i>	0.78	0.72	0.78	0.77	0.81	0.77	0.76	0.77
<i>WLM</i>	0.74	0.71	0.78	0.80	0.83	0.78	0.76	0.77
<i>ExclM</i>	0.77	0.72	0.77	0.74	0.78	0.73	0.74	0.75
<i>LODOverlap</i>	0.78	0.73	0.80	0.79	0.83	0.80	0.79	0.79
<i>LDSDGN_γ</i>	0.82	0.78	0.82	0.82	0.87	0.82	0.81	0.82

related sim_R (Eq. 2) for all the 28 contexts (pairs) of the *M&C* datasets.

In Table 3, the average correlations are shown, in particular, in the case of *CombIC* and *LDSDGN_γ* the last rows of Table 1 and Table 2 are given, respectively. In addition, the corresponding values obtained by using the other relatedness measures, namely *LDSD_{cw}*, *ASRMP_m^a*, *WLM*, *ExclM*, and *LODOverlap*, are shown.

Below the mentioned relatedness methods are briefly recalled.

The *CombIC* method relies on the IC notion, and in particular focuses on the computation of the ICs of the predicates and the objects of the triples of the knowledge graph.

LDSD_{cw} belongs to the *LDSD* family proposed by Passant which is at the basis of the above illustrated *LDSDGN_γ* measure. As mentioned in Subsection 3.1, the *LDSD* family does not satisfy the three axioms of Equal self-distance, Symmetry, and Minimality, and the

LDSDGN measures originate from it in order to meet these requirements.

According to *ASRMP_m^a*, the key idea is to evaluate all the directed paths connecting the resources to be compared, of length equal to m .

The *WLM* method is inspired by the *Normalized Google Distance* measure. It relies on the information gathered by the nodes that are adjacent to the compared resources and, in particular, on those with incoming links to them.

ExclM focuses on a given number of undirected paths and, in particular, it is based on a path weighting function that allows the selection of the undirected paths between the compared resources with the greatest weights.

Analogously to *WLM*, in the case of *LODOverlap* only the nodes that are adjacent to the compared resources are addressed, although the method also considers those with outgoing links from the compared resources.

According to the experimental results, *LDSDGN_γ* shows an average increment of the average Spearman's correla-

tion with HJ of about 0.02. Hence, the global normalization notion on which $LDSDGN_\gamma$ is based, which allows a path to be considered on the basis of the number of its occurrences in the whole knowledge graph, provides the best strategy.

As a result, the combination of the measure presented in Section 2 with the approach proposed in Piao and Breslin (2016) provides the best Spearman's correlation values in order to evaluate semantic similarity in a taxonomy by addressing the concept intended senses in a given context.

The data concerning the new experiment are available at Taglino (2022).

5. RELATED WORK

Evaluating semantic similarity is a fundamental topic in Computer Science in different application domains, De Nicola et al. (2023), such as health, Abdelrahman and Kayed (2015), network security, Weller-Fahy et al. (2015), and bibliometrics, Avram et al. (2012), just to mention a few examples. In particular, semantic similarity methods based on the IC approach have been employed in different research areas, such as Natural Language Processing, Majumder et al. (2016), Semantic Web, Formica et al. (2013), Meymandpour and Davis (2016), Formal Concept Analysis, Formica (2019), Wang et al. (2020), Geographical Information Systems, Formica and Pourabbas (2009), Formica et al. (2020), Schwering (2008), and e-learning, Barbagallo and Formica (2017).

The IC approach has shown higher correlation values with human judgment (HJ) with respect to other proposals which do not originate from it, Adhikari et al. (2018), Batet and Sánchez (2020), Chandrasekaran and Mago (2022), Formica et al. (2013). However, one objection to the early IC-based measures relies on the use of large-scale corpora, Adhikari et al. (2018), Banu et al. (2015), Batet and Sánchez (2020), Jeong et al. (2017), Zhang et al. (2018). In fact, evaluating the IC on the basis of statistical information taken from textual corpora requires a huge amount of manual effort at the level of both design and maintenance of the corpus. For this reason, in the literature, an evolution of the IC notion has been extensively investigated, referred to as *intrinsic information content* (IIC), although there is a lack of a statistically significant difference between the performances of the IIC models and the corpus-based ones, Lastra-Díaz and García-Serrano (2015). In particular, the IIC is evaluated independently of textual corpora, and in accordance with the intrinsic structure of the taxonomy, i.e., on the basis of the number of hyponyms and/or hypernyms of the concepts (see for instance Adhikari et al. (2016), Adhikari et al. (2018)). However, in this paper we do not present a new IC or IIC computing model, and this proposal is independent of it. In fact, although the IIC approaches show high accuracy in the similarity evaluation, they do not involve concept meaning and, in particular, the related similarity measures do not address the intended senses of concepts according to a given application domain.

OntoLex-Lemon, McCrae et al. (2017), allows the modeling and publishing of lexical resources on the Semantic Web, Khan et al. (2022). In particular, it provides support

for enriching ontologies with linguistic information, by modeling, for instance, the different meanings a word can denote. Furthermore, in recent evolutions, Chiarcos et al. (2022a), Chiarcos et al. (2022b), it is used to organize contextualized embeddings. However, as in the intention of the authors, it is a modeling instrument that does not allow semantic similarity to be computed.

Note that the semantic similarity measure proposed in Jeong et al. (2017) originates from the need to overcome the limitations we highlighted in Section 2 in WordNet. However, the authors propose a solution by associating the different kinds of relationships (e.g., ISA and PartOf) with different weights, which is again a proposal independent of the concept intended senses. We emphasize that in this paper we limited our attention to semantic similarity in ISA taxonomies, rather than considering more general knowledge graphs such as WordNet (and related proposals, as for instance Al Tarouti and Kalita (2016), and Chiru et al. (2021)). This is due to the novelty of this approach that, to our knowledge, has been experimented only with ISA taxonomies in Formica and Taglino (2021), and in Formica and Taglino (2023b), and the contribution of this paper is a further improvement in this direction.

In particular, although in Formica and Taglino (2021) the Spearman's correlation has not been explicitly mentioned because we focused on the Pearson's correlation, in this paper the $LDSDGN_\gamma$ shows the best performances with respect to the *CombIC* when considering the Spearman's correlation (see Table 3). Furthermore, with respect to the work presented in Formica and Taglino (2023b), we enriched the experimentation by adding the method $LDSDGN_\gamma$ which outperforms the other approaches (see Table 6 in the mentioned paper and Table 3 of this work).

The notion of sense has been addressed in Resnik (1999), where semantic similarity is used to identify and select the appropriate sense of a concept when it appears in a group of related terms. However, this paper addresses word sense disambiguation in the field of computational linguistics, where semantic similarity is not the objective of the work but is used in order to associate a noun with the right sense in a given context. On the contrary, we use the concept intended senses to improve the computation of semantic similarity.

6. CONCLUSION AND FUTURE WORK

In this paper, the problem of evaluating the semantic similarity of concepts organized according to a taxonomy has been addressed. In particular, the experiment presented in Formica and Taglino (2021) has been considered and, in order to evaluate the relatedness of the generic sense of a concept with its intended sense, a new experiment has been performed by relying on the $LDSDGN_\gamma$ measure, Piao and Breslin (2016), rather than the *CombIC* relatedness measure based on the IC approach presented in Schuhmacher and Ponzetto (2014). According to the new experiment, $LDSDGN_\gamma$ shows the best average Spearman's correlation against HJ with respect to the original proposal of the authors, Formica and Taglino (2021), and also with respect to the experiment presented in Formica and Taglino (2023b). To the best of our knowledge, this is a novel approach that is worth further investigation.

For future work, we would like to analyze the use of the *SemSim* semantic similarity method, Formica et al. (2013), in order to support the disambiguation step which is a fundamental aspect at the basis of both semantic similarity and the more general notion of semantic relatedness.

ACKNOWLEDGEMENT

We gratefully acknowledge the partial support of the PNRR MUR project PE0000013-FAIR.

REFERENCES

- Abdelrahman, A.M.B. and Kayed, A. (2015). A survey on semantic similarity measures between concepts in health domain. *American Journal of Computational Mathematics*, 05, 204–214.
- Adhikari, A., Dutta, B., Dutta, A., Mondal, D., and Singh, S. (2018). An intrinsic information content-based semantic similarity measure considering the disjoint common subsumers of concepts of an ontology. *J. Assoc. Inf. Sci. Technol.*, 69(8), 1023–1034. doi:10.1002/asi.24021.
- Adhikari, A., Singh, S., Mondal, D., Dutta, B., and Dutta, A. (2016). A novel information theoretic framework for finding semantic similarity in wordnet. *CoRR*, abs/1607.05422.
- Al Tarouti, F. and Kalita, J. (2016). Enhancing automatic Wordnet construction using word embeddings. In *Proc. of the Workshop on Multilingual and Cross-lingual Methods in NLP*, 30–34. Association for Computational Linguistics, San Diego, California. doi:10.18653/v1/W16-1204.
- Armaselu, F., Apostol, E.S., Khan, A.F., Liebeskind, C., McGillivray, B., Truică, C.O., Utku, A., Valūnaitė Oleškevičienė, G., and van Erp, M. (2022). LL (O) D and NLP perspectives on semantic change for humanities research. *Semantic Web*, 13(6), 1051–1080. doi:10.3233/SW-222848.
- Avram, S., Caragea, D., and Dumitrach, I. (2012). A new approach to bibliometrics based on semantic similarity of scientific papers. *Journal of Control Engineering and Applied Informatics*, 13, 35–42.
- Banu, A., Fatima, S.S., and Khan, K.U.R. (2015). Information content based semantic similarity measure for concepts subsumed by multiple concepts. *Int. J. Web Appl.*, 7, 85–94.
- Barbagallo, A. and Formica, A. (2017). ELSE: an ontology-based system integrating semantic search and e-learning technologies. *Interactive Learning Environments*, 25(5), 650–666. doi:10.1080/10494820.2016.1172240.
- Batet, M. and Sánchez, D. (2020). Leveraging synonymy and polysemy to improve semantic similarity assessments based on intrinsic information content. *Artif. Intell. Rev.*, 53(3), 2023–2041. doi:10.1007/s10462-019-09725-4.
- Chandrasekaran, D. and Mago, V. (2022). Evolution of semantic similarity - A survey. *ACM Comput. Surv.*, 54(2), 41:1–41:37. doi:10.1145/3440755.
- Chiarcos, C., Apostol, E.S., Kabashi, B., and Truică, C.O. (2022a). Modelling frequency, attestation, and corpus-based information with OntoLex-FrAC. In *Proc. of the 29th Int. Conf. on Computational Linguistics*, 4018–4027. Int. Committee on Computational Linguistics, Gyeongju, Republic of Korea.
- Chiarcos, C., Gkirtzou, K., Ionov, M., Kabashi, B., Khan, F., and Truică, C.O. (2022b). Modelling collocations in OntoLex-FrAC. In *Proc. of Globalex Workshop on Linked Lexicography within the 13th Language Resources and Evaluation Conference*, 10–18. European Language Resources Association, Marseille, France.
- Chiru, C.G., Truică, C.O., Apostol, E.S., and Ionescu, A. (2021). Improving wordnet using word embeddings. In *2021 23rd Int. Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNASC)*, 121–128. doi:10.1109/SYNASC54541.2021.00030.
- De Nicola, A., Formica, A., Missikoff, M., Pourabbas, E., and Taglino, F. (2023). A parametric similarity method: Comparative experiments based on semantically annotated large datasets. *Journal of Web Semantics*, 76, 100773. doi:org/10.1016/j.websem.2023.100773.
- El Vaigh, C.B., Goasdoué, F., Gravier, G., and Sébillot, P. (2020). A novel path-based entity relatedness measure for efficient collective entity linking. In J.Z. Pan, V.A.M. Tamma, C. d’Amato, K. Janowicz, B. Fu, A. Polleres, O. Seneviratne, and L. Kagal (eds.), *The Semantic Web - ISWC 2020 - 19th Int. Semantic Web Conference, Athens, Greece, November 2-6, 2020, Proceedings, Part I*, volume 12506 of *LNCS*, 164–182. Springer. doi:10.1007/978-3-030-62419-4_10.
- Formica, A. (2019). Similarity reasoning in formal concept analysis: from one- to many-valued contexts. *Knowl. Inf. Syst.*, 60(2), 715–739. doi:10.1007/s10115-018-1252-4.
- Formica, A., Mazzei, M., Pourabbas, E., and Rafanelli, M. (2020). Approximate query answering based on topological neighborhood and semantic similarity in openstreetmap. *IEEE Access*, 8, 87011–87030. doi:10.1109/ACCESS.2020.2992202.
- Formica, A., Missikoff, M., Pourabbas, E., and Taglino, F. (2013). Semantic search for matching user requests with profiled enterprises. *Comput. Ind.*, 64(3), 191–202. doi:10.1016/j.compind.2012.09.007.
- Formica, A. and Pourabbas, E. (2009). Content based similarity of geographic classes organized as partition hierarchies. *Knowl. Inf. Syst.*, 20(2), 221–241. doi:10.1007/s10115-008-0177-8.
- Formica, A. and Taglino, F. (2021). An enriched information-theoretic definition of semantic similarity in a taxonomy. *IEEE Access*, 9, 100583–100593. doi:10.1109/ACCESS.2021.3096598.
- Formica, A. and Taglino, F. (2023a). Semantic relatedness in DBpedia: A comparative and experimental assessment. *Information Sciences*, 621, 474–505. doi:https://doi.org/10.1016/j.ins.2022.11.025.
- Formica, A. and Taglino, F. (2023b). Semantic similarity in a taxonomy by evaluating the relatedness of concept senses with the linked data semantic distance. *TLDKS LIII*, LNCS 13840, 1–24. doi:org/10.1007/978-3-662-66863-4_3.
- Hulpus, I., Prangnawarat, N., and Hayes, C. (2015). Path-based semantic relatedness on linked data and its use to word and entity disambiguation. In M. Arenas, Ó. Corcho, E. Simperl, M. Strohmaier, M. d’Aquino, K. Srinivas, P. Groth, M. Dumontier, J. Heflin, K. Thirunarayan,

- and S. Staab (eds.), *The Semantic Web - ISWC 2015 - 14th Int. Semantic Web Conference, Bethlehem, PA, USA, October 11-15, 2015, Proceedings, Part I*, volume 9366 of *LNCS*, 442–457. Springer. doi:10.1007/978-3-319-25007-6_26.
- Jeong, S., Yim, J.H., Lee, H.J., and Sohn, M.M. (2017). Semantic similarity calculation method using information contents-based edge weighting. *J. Internet Serv. Inf. Secur.*, 7(1), 40–53. doi:10.22667/JISIS.2017.02.31.040.
- Jiang, J.J. and Conrath, D.W. (1997). Semantic similarity based on corpus statistics and lexical taxonomy. In K. Chen, C. Huang, and R. Sproat (eds.), *Proc. of the 10th Research on Computational Linguistics Int. Conf., ROCLING 1997, Taipei, Taiwan, August 1997*, 19–33. The Association for Computational Linguistics and Chinese Language Processing (ACLCLP).
- Khan, A.F., Chiacaros, C., Declerck, T., Gifu, D., García, E.G., Gracia, J., Ionov, M., Labropoulou, P., Mambri, F., McCrae, J.P., Pagé-Perron, É., Passarotti, M., Muñoz, S.R., and Truica, C. (2022). When linguistics meets web technologies. Recent advances in modelling linguistic linked data. *Semantic Web*, 13(6), 987–1050. doi:10.3233/SW-222859.
- Lastra-Díaz, J.J. and García-Serrano, A. (2015). A new family of information content models with an experimental survey on wordnet. *Knowl. Based Syst.*, 89, 509–526. doi:10.1016/j.knsys.2015.08.019.
- Lin, D. (1998). An information-theoretic definition of similarity. In J.W. Shavlik (ed.), *Proc. of the Fifteenth Int. Conference on Machine Learning (ICML 1998)*, Madison, Wisconsin, USA, July 24–27, 1998, 296–304. Morgan Kaufmann.
- Majumder, G., Pakray, P., and Gelbukh, A.F. (2016). Measuring semantic textual similarity of sentences using modified information content and lexical taxonomy. *Int. J. Comput. Linguistics Appl.*, 7(2), 65–85.
- McCrae, J.P., Gil, J.B., Gràcia, J., Bitelaar, P., and Cimi-ano, P. (2017). The ontolx-lemon model: Development and applications.
- Meymandpour, R. and Davis, J.G. (2016). A semantic similarity measure for linked data: An information content-based approach. *Knowl. Based Syst.*, 109, 276–293. doi:10.1016/j.knsys.2016.07.012.
- Miller, G.A. and Charles, W.G. (1991). Contextual correlates of semantic similarity. *Language and Cognitive Processes*, 6(1), 1–28.
- Milne, D. and Witten, I.H. (2008). An effective, low-cost measure of semantic relatedness obtained from wikipedia links. In *Proc. of AAAI Workshop on Wikipedia and Artificial Intelligence: an Evolving Synergy*, 25–30. AAAI Press.
- Montariol, S., Martinc, M., and Pivovarov, L. (2021). Scalable and interpretable semantic change detection. In K. Toutanova, A. Rumshisky, L. Zettlemoyer, D. Hakkani-Tür, I. Beltagy, S. Bethard, R. Cotterell, T. Chakraborty, and Y. Zhou (eds.), *Proc. of the 2021 Conf. of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL-HLT 2021, Online, June 6–11, 2021*, 4642–4652. Association for Computational Linguistics. doi:10.18653/v1/2021.naacl-main.369.
- Passant, A. (2010). Measuring semantic distance on linking data and using it for resources recommendations. In *Linked Data Meets Artificial Intelligence, Papers from the 2010 AAAI Spring Symposium, Technical Report SS-10-07, Stanford, California, USA, March 22–24, 2010*. AAAI.
- Piao, G. and Breslin, J.G. (2016). Measuring semantic distance for linked open data-enabled recommender systems. In S. Ossowski (ed.), *Proc. of the 31st Annual ACM Symposium on Applied Computing, Pisa, Italy, April 4–8, 2016*, 315–320. ACM. doi:10.1145/2851613.2851839.
- Pirró, G. (2009). A semantic similarity metric combining features and intrinsic information content. *Data Knowl. Eng.*, 68(11), 1289–1308. doi:10.1016/j.datak.2009.06.008.
- Rada, R., Mili, H., Bicknell, E., and Blettner, M. (1989). Development and application of a metric on semantic nets. *IEEE Trans. Syst. Man Cybern.*, 19(1), 17–30. doi:10.1109/21.24528.
- Resnik, P. (1995). Using information content to evaluate semantic similarity in a taxonomy. In *Proc. of the Fourteenth Int. Joint Conf. on Artificial Intelligence, IJCAI 95, Montréal Québec, Canada, August 20–25 1995, 2 Volumes*, 448–453. Morgan Kaufmann.
- Resnik, P. (1999). Semantic similarity in a taxonomy: An information-based measure and its application to problems of ambiguity in natural language. *J. Artif. Intell. Res.*, 11, 95–130. doi:10.1613/jair.514.
- Schuhmacher, M. and Ponzetto, S.P. (2014). Knowledge-based graph document modeling. In B. Carterette, F. Diaz, C. Castillo, and D. Metzler (eds.), *Seventh ACM Int. Conf. on Web Search and Data Mining, WSDM 2014, New York, NY, USA, February 24–28, 2014*, 543–552. ACM. doi:10.1145/2556195.2556250.
- Schwering, A. (2008). Approaches to semantic similarity measurement for geo-spatial data: A survey. *Trans. GIS*, 12(1), 5–29. doi:10.1111/j.1467-9671.2008.01084.x.
- Taglino, F. (2022). Semantic Similarity with Concept Senses: new Experiment. Mendeley Data. URL <https://data.mendeley.com/datasets/v2bwh7z8kj/1>. V1.
- Taieb, M.A.H., Zesch, T., and Aouicha, M.B. (2020). A survey of semantic relatedness evaluation datasets and procedures. *Artif. Intell. Rev.*, 53(6), 4407–4448. doi:10.1007/s10462-019-09796-3.
- Wang, F., Wang, N., Cai, S., and Zhang, W. (2020). A similarity measure in formal concept analysis containing general semantic information and domain information. *IEEE Access*, 8, 75303–75312. doi:10.1109/ACCESS.2020.2988689.
- Weller-Fahy, D.J., Borghetti, B.J., and Sodemann, A.A. (2015). A survey of distance and similarity measures used within network intrusion anomaly detection. *IEEE Commun. Surv. Tutorials*, 17(1), 70–91. doi:10.1109/COMST.2014.2336610.
- Wu, Z. and Palmer, M.S. (1994). Verb semantics and lexical selection. In J. Pustejovsky (ed.), *32nd Annual Meeting of the Association for Computational Linguistics, 27–30 June 1994, New Mexico State University, Las Cruces, New Mexico, USA, Proceedings*, 133–138. Morgan Kaufmann Publishers / ACL. doi:10.3115/981732.981751.
- Zhang, X., Sun, S., and Zhang, K. (2018). An information content-based approach for measuring concept semantic

- similarity in wordnet. *Wirel. Pers. Commun.*, 103(1), 117–132. doi:10.1007/s11277-018-5429-7.
- Zhou, W., Wang, H., Chao, J., Zhang, W., and Yu, Y. (2011). LODDO: using linked open data description overlap to measure semantic relatedness between named entities. In J.Z. Pan, H. Chen, H. Kim, J. Li, Z. Wu, I. Horrocks, R. Mizoguchi, and Z. Wu (eds.), *The Semantic Web - Joint Int. Semantic Technology Conf., JIST 2011, Hangzhou, China, December 4-7, 2011.*, volume 7185 of *LNCIS*, 268–283. Springer.