

Relations extracted from a Portuguese dictionary: results and first evaluation

Hugo Gonalo Oliveira¹, Diana Santos², Paulo Gomes¹
hroliv@dei.uc.pt, diana.santos@sintef.no, pgomes@dei.uc.pt

¹ CISUC, University of Coimbra, Portugal

² Linguateca, SINTEF ICT, Oslo, Norway

Abstract. This paper presents PAPEL, a lexical resource for Portuguese, consisting of relations between terms, extracted by (semi) automatic means from a general language dictionary. An overview on the construction process is given, the included relations are presented and a quantitative vision is provided together with some examples. Synonymy relations were evaluated using a thesaurus as a golden standard and the other relations were rendered to natural language patterns and searched for in a corpus. The results of the evaluation are shown and discussed.

1 Introduction

In this paper, we present the current situation and the first automatic evaluation of PAPEL³ - Palavras Associadas Porto Editora Linguateca, a set of relations between terms compiled into a lexical resource for Portuguese natural language processing (NLP). PAPEL [1] was constructed (semi) automatically by processing the definitions of a general dictionary of Portuguese [2], developed and owned by a large dictionary publisher, Porto Editora, and extracting relations denoted by textual patterns. The resulting relations were then validated, in the following way: synonymy relations were compared to the relations in a large thesaurus, while the other relations were rendered in natural language and searched for in a text corpora.

While for English, in the last decade, WordNet [3] was established as the standard model of a lexical database, the picture is quite different for other languages. There are attempts to create a similar database for Portuguese, namely Wordnet.BR [4], Wordnet.PT [5] and MultiWordnet.PT⁴ but at the time of writing they were not freely available for download.

Those resources are however the product of time-consuming manual work, so we propose a semi-automatic construction of such a resource. Recent attempts on the automatic extraction of relations in Portuguese deal primarily with the hyponymy relation: Freitas and Quental [6] discuss its extraction from corpora while Costa and Seco [7] focus on user's search behaviour in a web search engine (logs). We are not aware of attempts to extract (semi) automatically semantic

³ <http://www.linguateca.pt/PAPEL>

⁴ <http://mwntp.di.fc.ul.pt>

relations of other types from text written in Portuguese and compile them into one independent resource.

This paper starts by describing background work on knowledge extraction from machine readable dictionaries (MRDs) as well as overview methods for the evaluation of ontologies (Section 2). We then present the approach taken in the construction of PAPEL (Section 3), followed by a thorough description of its current contents, more precisely the relations included, their number and illustrative examples (Section 4). The evaluation attempted is finally described (Section 5), before discussing ideas for further work (Section 6).

2 Background

2.1 MRDs as a source of knowledge

The process of using MRDs in NLP started more than thirty years ago with early works of Calzolari [8], Amsler [9], and Chodorow et al. [10]. MRDs were analysed and, taking advantage of the simple structure of the definitions and of the restricted vocabulary used, procedures were developed to extract and organise lexical information. Following these ideas, Alshawi [11] proposed a specific grammar for parsing the definitions of a particular dictionary, based on the syntactic patterns used, and producing semantic structures. In the 1990s, as illustrated by Montemagni and Vanderwende [12], it became more common to use broad-coverage parsers to extract semantic information from dictionary text, claiming they were better suited to capture the distinguishing features in the definition, although this was not consensual in the community (see e.g. Hearst [13]).

In any case, one of the main reasons to use dictionaries and not (only) running text is because MRDs are the "authorities" on word sense [14]. Dictionaries have thus been exploited for several purposes, such as parsing or word sense disambiguation (WSD), but to our knowledge they have not been converted into an independent resource of its own before MindNet [15], that can therefore be claimed to be a kind of independent (dictionary-based) lexical ontology in a way that previous work was not. More recently, Nichols [16] and O'Hara [17] also worked on the extraction of semantic relations from MRDs.

2.2 Evaluation of ontologies

When it comes to ontology evaluation, and although the information retrieval precision and recall measures are increasingly being used for evaluation in NLP (see e.g. Santos [18]), for semantic lexicons it is hard to have an independent golden standard for what should be there in the first place. The knowledge that should be represented is not clear. If we compare it with semantic data extracted from text, we have to remember that different interpretations and different meanings are often possible [19].

For domain ontologies, Brank et al. [20] divide evaluation approaches into four groups: performed by human subjects; comparison with a golden standard;

as for coverage, comparison with a collection of documents about a domain covered by the ontology; accomplishment of some task that uses the ontology.

Although the most reliable in the end, human evaluation does not take advantage of computer programs and relies heavily on time consuming work from at least one domain specialist. In order to make human evaluation easier, Navigli et al. [21] generated natural language descriptions of concepts, based on a grammar with distinct generation rules for each type of semantic relation.

Of course, the ontology can be compared with some other resource (e.g. another ontology) that is known to be correct, usually because it was created by specialists. But if this may be OK to validate a particular automatic method, it is obviously of little practical interest, because one expects to be creating new ontologies, not recreating existing ones. So, while the approach of compiling a human resource is commonly followed in joint evaluations, for example ReRelEM [22], which evaluated system's capabilities to recognise semantic relations between named entities, it can only encompass a few examples.

The third approach consists of finding how adequate a particular ontology is for representing the knowledge contained in a collection of documents, as in Brewster et al.'s [19] measurement of the fit between an ontology and a corpus. After identifying salient terms in a domain corpus and looking for them in an ontology for the same domain, the fit is proportional to the number of terms found in both corpus and ontology. The problem is that we cannot define a clear set of salient terms for general language, so this method cannot be applied to a lexical ontology that is supposed to describe the former.

The last approach, external or task-based evaluation, performs (indirect) evaluation by assessing the performance of an application which uses the ontology to do some task. Porzel and Malaka [23] proposed this approach aiming at evaluating ontologies with respect to the fit of the vocabulary, the fit of the taxonomy and the adequacy of non-taxonomic semantic relations.

These methods were used to evaluate domain ontologies, but if we consider dictionaries or lexical ontologies, evaluation is not a common practice, possibly because most of these resources are manually created by specialists. Ide and Verónis [24] are very critical of this fact and produced work for assessing the quality and usefulness of information extracted from MRDs. They concluded that the obtained the structures obtained by applying *Chodorow-like* procedures were incomplete and had several other problems but, if they merged the results extracted from several MRDs, the amount of problems decreased drastically.

Among the few attempts for evaluating information automatically extracted from MRDs, Richardson et al. [25] hand-checked a random sample of 250 semantic relations automatically extracted from a dictionary (and later included in MindNet), relying on common statistical techniques to estimate the representativeness of the accuracy for all the relations extracted. For MindNet [26], an (incomplete) evaluation of the quality of the semantic relations is mentioned, but its authors do not go very far in the description of the evaluation process. One comment made is that the quality varies according to the relation type. Likewise, to evaluate an ontology extracted automatically from a MRD, Nichols

et al. [16] used Wordnet and GoiTaiKei [27] as a golden resource. In this process, they noticed that some relations were only in one of the two golden resources, which might indicate that they are both incomplete.

3 The approach for building PAPEL

The construction process followed four stages: (i) creation of the extraction grammars; (ii) extraction of the relations; (iii) manual result inspection and, finally, (iv) relations adjustment.

3.1 Extraction grammars

Inspired by Alshawi’s work [11], specific grammars to parse the dictionary definitions were manually created. These grammars aim at the extraction of specifically predefined relations (described in Section 4) and are based on a previous empirical analysis of the structure of the definitions and of the vocabulary used. Table 1 shows some of the patterns that are included in the grammars together with the relations they are associated with.

Pattern	Associated relation
tipo género classe forma de	Hypernymy
parte membro de	Meronymy
que causal provoca origina	Causation
usado utilizado para	Purpose
<i>a word or an enumeration of words</i>	Synonymy

Table 1. Examples of patterns used in the grammars.

```

PARTE{
  nome:nome * PARTE_DE:INCLUI;
  nome:adj * PARTE_DE_ALGO_COM_PROPRIEDADE:PROPRIEDADE_DE_ALGO_QUE_INCLUI;
  adj:nome * PROPRIEDADE_DE_ALGO_PARTE_DE:INCLUI_ALGO_COM_PROPRIEDADE;
}

```

Fig. 1. Examples of the description of the meronymy relations group.

Each grammar is made to process definitions of words belonging to only one of the four open grammatical categories (nouns, verbs, adjectives, adverbs). The relations to be extracted are defined according to their group, their name, the name of their inverse relation and, since cross-categorical relations can also be extracted, the grammatical category of its arguments (see the example in Figure 1). Most of the relation types have a grammar for both the type defined as direct and the inverse type.

3.2 Relation extraction proper

In the extraction stage, a chart parser uses the grammars to process the dictionary definitions. Every time a definition suits the grammar rules, the parser generates a derivation tree. Although it is possible to get more than one derivation for the same grammar/definition pair, currently only the derivation with less unidentified nodes is chosen. The same derivation tree can be used to extract several relations, provided they were identified by the grammar (see the example in Figure 2).

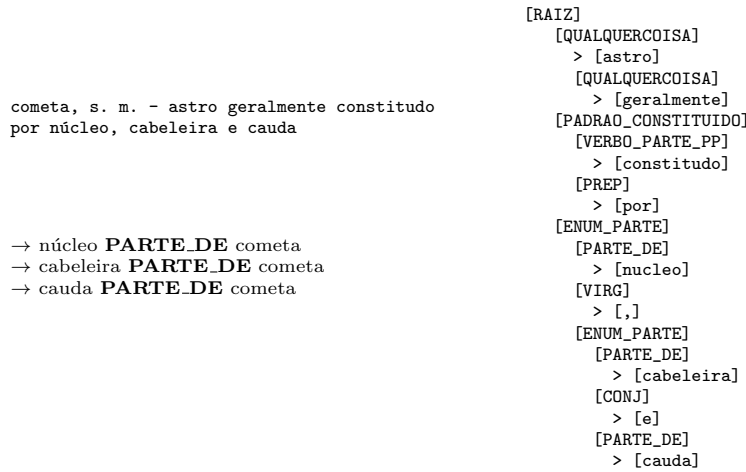


Fig. 2. Derivation for the definition of *cometa*.

3.3 Manual result inspection

The extraction results are inspected in order to identify systematic problems, and with the two previous steps form a loop that can be repeated at will.

Results from different extraction runs can be automatically compared to guarantee that newer results are better than older ones. After this procedure, we go back to the first stage, in which newer versions of the grammars are created, hopefully with some of the identified problems corrected.

3.4 Relations adjustment

After several loops of processing, the construction enters a new stage, where the relations with inadequate arguments (i.e. arguments whose grammatical category does not agree with the relation name) are either corrected or discarded. In order to simplify the relation set, all relations are translated into the type defined as direct. This stipulation was made based on what seemed more natural to the grammar writer, and not on frequency considerations. For example,

manga INCLUI *punho* is translated to *punho* PARTE.DE *manga* and *dor* RESULTADO.DE *distensão* becomes *distensão* CAUSADOR.DE *dor*.

Given that the grammars have little grammatical information (introduced manually) and each dictionary entry only contains the grammatical category of the word being defined, in some cases we get relations with arguments that do not belong to the correct category. So, the grammatical category of each argument is verified, with the help of the grammatical information in the dictionary and, when the argument is not defined in the dictionary, with the help of the Jspell [28] morphological analyser. If the arguments of a relation are not adequate but there is a relation type that belongs to the same group and suits the categories of the arguments, the relation type is replaced, otherwise the relation is discarded. For example, the relation *loucura* ACCAO.QUE.CAUSA *desvario* becomes *loucura* CAUSADOR.DE *desvario*, because both arguments are nouns. During this verification, if an argument is not in the lemma form, it is changed to it, again with the help of Jspell.

4 A closer look at PAPEL

The set of relation types in PAPEL was chosen after reviewing the relations described in the literature and included in similar resources like WordNet [3] or MindNet [29]. We also took into consideration potential relation types that could be extracted from the most frequent patterns used in the definitions of the dictionary. The relations extracted are divided into eight main groups and have specific names according to the group and the grammatical category of the arguments.

After automatically correcting some of the relations based on the grammatical category of their arguments and removing duplicate relations we got the final set, comprising slightly more than 200,000 relations. Table 2 presents the specific types of relations, the grammatical categories of their attributes and quantifies each type of relation in PAPEL along with examples of extracted relations.

As it can be seen, synonymy and hypernymy are the most frequent relations, and we note that this set can still be further augmented if relations are combined and rules are applied in order to give rise to new implicit relations, as we intend to experiment later. This can be done in a similar fashion to what was done in the ReReLEM [22] task where, before comparing the golden collection and the participant run, they both had their relation set expanded with the application of inverse and transitive rules⁵.

5 Evaluation of PAPEL

As we stated in Section 2, the evaluation of lexical ontologies is not that common. In order to evaluate PAPEL, we had a look at the approaches used to evaluate domain ontologies, discussed in the same section.

⁵ These programs have been made publicly available by the HAREM [30] organisers.

Group	Name	Args.	Qnt.	Examples
Synonymy	SINONIMO.DE	same	80,432	(flexível, moldável)
Hypernymy	HIPERONIMO.DE	n,n	63,455	(planta, salva)
Meronymy	PARTE.DE	n,n	14,453	(cauda, cometa)
	PARTE.DE_ALGO_COM_PROP	n,adj	3,715	(tampa, coberto)
	PROP.DE_ALGO_PARTE.DE	adj,n	962	(celular, célula)
	CAUSADOR.DE	n,n	1,125	(fricção, assadura)
Cause	CAUSADOR.DE_ALGO_COM_PROP	n,adj	16	(paixão, passionai)
	PROP.DE_ALGO_CAUSADOR.DE	adj,n	5,15	(reactivo, reacção)
	ACCAO.QUE_CAUSA	v,n	6,424	(limpar, purgação)
	CAUSADOR.DA_ACCAO	n,v	39	(gases, fumigar)
	PRODUTOR.DE	n,n	932	(romãzeira, romã)
Producer	PRODUTOR.DE_ALGO_COM_PROP	n,adj	31	(sublimação, sublimado)
	PROP.DE_ALGO_PRODUTOR.DE	adj,n	348	(fotógeno, luz)
	FINALIDADE.DE	n,n	2,095	(defesa, armadura)
Purpose	FINALIDADE.DE_ALGO_COM_PROP	n,adj	23	(reprodução, reprodutor)
	ACCAO.FINALIDADE.DE	v,n	5,640	(fazer_rir, comédia)
	ACC.FINALIDADE.DE_ALGO_COM_PROP	v,adj	255	(corrigir, correcional)
	MANEIRA_POR_MEIO.DE	adv,n	1,433	(timidamente, timidez)
	LOCAL.ORIGEM.DE	n,n	768	(Japão, japonês)
Property	PROP.DE_ALGO_REFERENTE_A	adj,n	3,700	(dinâmico, movimento)
	PROP.DO_QUE	adj,v	17,028	(familiar, ser_conhecido)

Table 2. The relations of PAPEL.

We had so far not the time to do human evaluation, nor access to similar resources. Fortunately, TeP (Thesaurus Eletrônico para o Português do Brasil) [31] is not only available through a web interface⁶, but its knowledge base can be fully downloaded, so we used it as the golden standard for the evaluation of synonymy relations. Of course we are fully aware of lexical differences between the two varieties of Portuguese [32], but we believe that the common core is still and by far largest.

When it comes to the other relations, we decided to follow another validation approach, which can be presented as a combination of task-based evaluation and text corpora validation. We developed a "deconstruction" procedure based on grammars for translating the relations into natural language patterns, which were then searched for in a corpus, in order to identify whether the corpus lent the relations some support.

5.1 Evaluation of synonymy

TeP is an electronic thesaurus manually created for Brazilian Portuguese, comprising 19,888 synsets and 44,678 lexical units. Similarly to WordNet, each synset is a list of words that can have the same meaning.

The evaluation process should compare our set of synonymy relations with the relations implicitly defined by every synset. But since there were terms in PAPEL that did not appear in TeP and vice-versa, we first discarded all relations with at least one argument not either in TeP or in PAPEL and were left with about 68% of PAPEL and 35% of the 405,026 TeP relations⁷ that we used for

⁶ <http://www.nilc.icmc.usp.br/tep2/>

⁷ To convert Tep, all elements of a synset were considered to be involved in a synonymy relation with all other elements of the same synset.

validation. After comparing both sets, 50% of the relations of PAPEL were found in TeP and 39% of the relations in TeP were in PAPEL.

As we said back in Section 4, the relations of PAPEL include only relations that were found in the dictionary and have not been the target of any kind of combination to infer new relations. If they were, we would have more relations to submit to the evaluation process. So, we applied the transitivity rule to our set: each pair of relations was combined and every time a pair had one common argument and a different one, a new synonymy relation was inferred (A SINONIMO_DE B $\wedge$ B SINONIMO_DE C $\rightarrow$ A SINONIMO_DE C).

After applying the transitivity rule, our 80,432 synonymy relations became 689,073. It was applied only once, otherwise the set would be much larger but would also have more inconsistencies than it already had after the first expansion. This happens because in PAPEL the key structure is the term and we do not handle polysemy or even homonymy, which means that two homographs will be treated as the same word. Transitivity in these conditions can give rise to clearly incorrect relations, such as: *queda* SINONIMO_DE *ruína* $\wedge$ *queda* SINONIMO_DE *habilidad* $\rightarrow$ *ruína* SINONIMO_DE *habilidad*. After comparing the expanded set with TeP, as expected the number of correct cases attested in TeP dropped to just 14% but, on the other hand, almost all synonyms in TeP were attested in our resource: 90%.

5.2 Validation of other relations between nouns

To assess the correctness of the non-synonymy relations we searched for natural language renderings of those relations in a corpus. This procedure was partly inspired by Etzioni et al. [33], who search for hypernymy patterns in the web to evaluate if a named entity is an instance of a specific class.

For this purpose we used CETEMPúblico, an annotated corpus provided by Linguatca with text from the Público newspaper published between 1991 and 1998, amounting to approximately 180 million words [34]. Despite being available for download, we used the AC/DC project [35] interface to query the corpus.

Although we started trying out queries in order to check all relations in our resource, we soon realised that for some of the relations it would be extremely improbable to find them in any (naturally occurring) text. For example, it is unlikely to find patterns to validate most of our cross-categorical relations, which seem to be precisely characteristic of the "dictionary genre": *liquidar* ACCAO_QUE_CAUSA *liquidação*, *fósforo* PARTE_DE_ALGO_COM_PROPRIEDADE *fosforoso*. It is not likely to find both arguments of each one of these relations in the same sentence.

As consequence, in the validation procedure followed, we dealt only with the noun to noun relations, and also with the cases attested in the corpus. That is, before starting the validation, all the relations including at least one argument that is not present in CETEMPúblico were discarded, by using the frequency lists of words and lemmas also provided by Linguatca.

We also had to select two random samples of the two most represented relations in our relation set, because they were just too many to validate in a

short time: so a random sample of 3,145 hypernymy relations (8%) and of 2,343 meronymy relations (63%) were selected. For the remaining relations, whose results are shown in Table 3, we used the complete noun to noun relation sets.

Relation	Relations w/ args in CETEMPúblico	%	Sample	%	Hits	%
Hypernymy	40,079	63%	3,145	8%	560	18%
Meronymy	3,746	35%	2,343	63%	521	22%
Causation	557	50%	557	100%	20	4%
Producer	414	44%	414	100%	12	3%
Purpose	1,718	59%	1,718	100%	173	10%

Table 3. Results of the validation of non-synonymy noun to noun relations.

As one can see, around 20% of the hypernymy and meronymy relations seem to be supported by the text in the corpus. When it comes to other relations, the percentage of hits is smaller. We believe anyway that the evidence is good since it is obvious that a 200 million-word corpus of a newspaper genre has to contain much less general knowledge than a general dictionary. Besides, we only used a small set of patterns – often similar to the ones used in the extraction grammars – while there is a huge amount of possibilities to represent each relation in unrestricted text, full of modifiers and anaphoric references, contrary to simple and structured dictionary definitions.

To give a more concrete feeling about the validation results, Table 4 contains some correct relations that were supported by the corpus, but also relations that seem to be supported but are incorrect. On the other hand, there were correct relations that were not found in CETEMPúblico, for instance: *fruto* HIPERONIMO.DE *alperce*, *algoritmia* PARTE.DE *matemática*, *ausência* CAUSADOR.DE *saudade* or *aquecimento* FINALIDADE.DE *salamandra*.

6 Conclusions and further work

A new publicly-available lexical resource for Portuguese was presented in detail in this paper along with some first attempt to evaluate it. We expect it to be useful for all kinds of NLP applications, from automatic generation of text to intelligent search and writing aids, as well as to more theoretical studies of the semantics of the Portuguese language.

Although the evaluation results are still preliminary, we intend to use other sources of information. For instance, validate the relations by searching for indicative sentences in the whole web, through general search engines. However we will have to deal with even more patterns for expressing the relations because these engines do not have lemma or other linguistic annotations. We also plan to use other corpus resources, for example those already annotated with synonyms.

It should be emphasised that we do not see this resource as a final one, but as an important seed for further enrichment by the whole community. Not

Relation	Correct	Support
<i>língua</i> HIPERONIMO_DE <i>italiano</i>	Yes	As línguas latinas, como o italiano ou o português, tornam-se mais fáceis por causa das vogais.
<i>arbusto</i> PARTE_DE <i>floresta</i>	Yes	A floresta é um conjunto de árvores, arbustos e ervas de várias qualidades e tamanhos.
<i>cólera</i> CAUSADOR_DE <i>diarreia</i>	Yes	A cólera provoca fortes diarreias e vômitos e pode levar à desidratação e, consequentemente, à morte em poucas horas.
<i>oliveira</i> PRODUTOR_DE <i>azeitona</i>	Yes	Também a quantidade e tamanho das azeitonas produzidas por uma oliveira biológica é inferior, já que não são utilizados compostos de azoto que ajudam a planta a crescer.
<i>recrutamento</i> FINALIDADE_DE <i>inspecção</i>	Yes	Menos de metade dos jovens entre os 20 e os 22 anos apresentaram-se às inspecções para recrutamento , revelou o ministro da Defesa.
<i>músico</i> PARTE_DE <i>música</i>	No	... um espectáculo baseado na obra "Cantos de Maldoror", de Lautréamont, com música composta pelo músico inglês Steven Severin...
<i>fim</i> FINALIDADE_DE <i>sempre</i>	No	Sicília aponta sempre para o fim do dia, para o fim da luz.

Table 4. Results of the validation of non-synonymy noun to noun relations.

only we suppose there can be other sources for (semi) automatically enriching it – according to Ide and Verónis [24], the information in dictionaries is often inconsistent and incomplete, so, in order to minimize that, a lexical ontology should be the result of a merge of several sources – but we intend to go on devising ways to increase coverage, linguistic correctness and validation.

Acknowledgments

We would like to thank the R&D group of Porto Editora for making their dictionary available for this research. The project PAPEL is supported by the Linguateca, jointly funded by the Portuguese Government, the European Union (FEDER and FSE), under contract ref. POSC/339/1.3/C/NAC, and by UMIC and FCCN.

References

1. Gonçalo Oliveira, H., Gomes, P., Santos, D., Seco, N.: PAPEL: a dictionary-based lexical ontology for Portuguese. In Teixeira, A., de Lima, V.L.S., de Oliveira, L.C., Quaresma, P., eds.: Proc. Computational Processing of the Portuguese Language, 8th Intl. Conference (PROPOR). Volume Vol. 5190., Springer Verlag (2008) 31–40
2. : Dicionário PRO da Língua Portuguesa. Porto Editora, Porto (2005)
3. Fellbaum, C., ed.: WordNet: An Electronic Lexical Database (Language, Speech, and Communication). The MIT Press (May 1998)
4. Dias-da-Silva, B.C., Oliveira, M., Moraes, H.: Groundwork for the Development of the Brazilian Portuguese Wordnet. In Mamede, N., Ranchhod, E., eds.: Proc. Advances in Natural Language Processing: 3rd Intl. Conference. Lecture Notes in Artificial Intelligence, Berlin/Heidelberg, Springer (23-26 de Junho 2002) 189–196

5. Marrafa, P.: Portuguese wordnet: general architecture and internal semantic relations. *Documentação de Estudos em Linguística Teórica e Aplicada (DELTA)* **18** (2002) 131–146
6. Freitas, C., Quental, V.: Subsídios para a elaboração automática de taxonomias. In: XXVII Congresso da SBC - V Workshop em Tecnologia da Informação e da Linguagem Humana (TIL). (2007) 1585–1594
7. Costa, R.P., Seco, N.: Hyponymy extraction and web search behavior analysis based on query reformulation. In: Proc. 11th Ibero-American Conference on Artificial Intelligence. LNAI, Springer Verlag (2008) 332–341
8. Calzolari, N., Pecchia, L., Zampolli, A.: Working on the italian machine dictionary: a semantic approach. In: Proc. 5th conference on Computational linguistics, Morristown, NJ, USA, Association for Computational Linguistics (1973) 49–52
9. Amsler, R.A.: A taxonomy for english nouns and verbs. In: Proc. 19th annual meeting on Association for Computational Linguistics, Morristown, NJ, USA, Association for Computational Linguistics (1981) 133–138
10. Chodorow, M.S., Byrd, R.J., Heidorn, G.E.: Extracting semantic hierarchies from a large on-line dictionary. In: Proc. 23rd annual meeting on Association for Computational Linguistics, Morristown, NJ, USA, Association for Computational Linguistics (1985) 299–304
11. Alshawi, H.: Analysing the dictionary definitions. *Computational lexicography for natural language processing* (1989) 153–169
12. Montemagni, S., Vanderwende, L.: Structural patterns vs. string patterns for extracting semantic information from dictionaries. In: Proc. 14th conference on Computational linguistics, Morristown, NJ, USA, Association for Computational Linguistics (1992) 546–552
13. Hearst, M.A.: Automatic acquisition of hyponyms from large text corpora. In: Proc. 14th conference on Computational linguistics, Morristown, NJ, USA, Association for Computational Linguistics (1992) 539–545
14. Kilgarriff, A.: I don’t believe in word senses: Computing and the Humanities **31**(2) (1997) 91–113
15. Richardson, S.D., Dolan, W.B., Vanderwende, L.: Mindnet: Acquiring and structuring semantic information from text. In: COLING-ACL. (1998) 1098–1102
16. Nichols, E., Bond, F., Flickinger, D.: Robust ontology acquisition from machine-readable dictionaries. In Kaelbling, L.P., Saffiotti, A., eds.: Proc. 19th Intl. Joint Conference on Artificial Intelligence (IJCAI), Professional Book Center (2005) 1111–1116
17. O’Hara, T.P.: Empirical Acquisition of Conceptual Distinctions via Dictionary Definitions. PhD thesis, NMSU CS (August 2005)
18. Santos, D.: Evaluation in natural language processing (6-17 August 2007)
19. Brewster, C., Alani, H., Dasmahapatra, S., Wilks, Y.: Data-driven ontology evaluation. In: Proc. Language Resources and Evaluation Conference (LREC), Lisbon, Portugal, European Language Resources Association (2004) 164–168
20. Brank, J., Grobelnik, M., Mladenic, D.: A survey of ontology evaluation techniques. In: Proc. Conference on Data Mining and Data Warehouses (SiKDD). (2005)
21. Navigli, R., Velardi, P., Cucchiarelli, A., Neri, F.: Quantitative and qualitative evaluation of the ontolearn ontology learning system. In: Proc. 20th Intl. conference on Computational Linguistics, Morristown, NJ, USA, Association for Computational Linguistics (2004)

22. Freitas, C., Santos, D., Mota, C., Oliveira, H.G., Carvalho, P.: Detection of relations between named entities: report of a shared task. In: *Semantic Evaluations: Recent Achievements and Future Directions (SEW)*, NAACL-HLT Workshop. (junho 2009)
23. Porzel, R., Malaka, R.: A task-based approach for ontology evaluation. In Buiteelaar, P., Handschuh, S., Magnini, B., eds.: *Proc. ECAI 2004 Workshop on Ontology Learning and Population*, Valencia, Spain (2004)
24. Ide, N., Veronis, J.: Knowledge extraction from machine-readable dictionaries: An evaluation. In Steffens, P., ed.: *Machine Translation and the Lexicon*, LNAI, Springer-Verlag (1995)
25. Richardson, S., Vanderwende, L., Dolan, W.: Combining dictionary-based and example-based methods for natural language analysis. In: *Proc. 5th Intl. Conference on Theoretical and Methodological Issues in Machine Translation*, Kyoto, Japan (1993) 69–79
26. Vanderwende, L., Kacmarcik, G., Suzuki, H., Menezes, A.: Mindnet: An automatically-created lexical resource. In: *HLT/EMNLP, The Association for Computational Linguistics* (2005)
27. Ikehara, S., Miyazaki, M., Shirai, S., Yokoo, A., Nakaiwa, H., Ogura, K., Ooyama, Y., Hayashi, Y.: *Goi-Taikei – A Japanese Lexicon*. Iwanami Shoten, Tokyo, Japan (1997)
28. Simões, A.M., Almeida, J.: Jspell.pm – um módulo de análise morfológica para uso em processamento de linguagem natural. In: *Actas do XVII Encontro da Associação Portuguesa de Linguística*, Lisboa (2002) 485–495
29. Richardson, S.D., Dolan, W.B., Vanderwende, L.: Mindnet: acquiring and structuring semantic information from text. In: *Proc. 17th Intl. Conference on Computational linguistics*, Morristown, NJ, USA, Association for Computational Linguistics (1998) 1098–1102
30. Santos, D., Freitas, C., Oliveira, H.G., Carvalho, P.: Second HAREM: new challenges and old wisdom. In et al., A.T., ed.: *Computational Processing of the Portuguese Language, 8th International Conference, Proceedings (PROPOR 2008)*. Volume Vol. 5190., Springer Verlag (2008) 212–215
31. Maziero, E., Pardo, T., Di Felippo, A., Dias-da Silva, B.: A base de dados lexical e a interface web do tep 2.0 - thesaurus eletrônico para o português do brasil. In: *VI Workshop em Tecnologia da Informação e da Linguagem Humana (TIL)*. (2008) 390–392
32. Wittmann, L., Pêgo, T., Santos, D.: Português do Brasil e de Portugal: alguns contrastes. In: *Actas do XI Encontro Nacional da Associação Portuguesa de Linguística*, Lisboa, APL/Colibri (1995) 465–487
33. Etzioni, O., Cafarella, M., Downey, D., Popescu, A.M., Shaked, T., Soderland, S., Weld, D.S., Yates, A.: Unsupervised named-entity extraction from the web: an experimental study. *Artificial Intelligence* **165**(1) (2005) 91–134
34. Santos, D., Rocha, P.: Evaluating CETEMPblico, a free resource for Portuguese. In: *Proceedings of the 39th Annual Meeting of the Association for Computational Linguistics*. (9-11 July 2001) 442–449
35. Santos, D., Bick, E.: Providing Internet access to Portuguese corpora: the AC/DC project. In Gavrilidou, M., Carayannis, G., Markantonatou, S., Piperidis, S., Stainhauer, G., eds.: *Proc. of the 2nd International Conference on Language Resources and Evaluation (LREC)*. (2000) 205–210