

ALGORITHM FOR AUTOMATIC INTERPRETATION OF NOUN SEQUENCES

Lucy Vanderwende

Microsoft Research
lucyv@microsoft.com

ABSTRACT

This paper describes an algorithm for automatically interpreting noun sequences in unrestricted text. This system uses broad-coverage semantic information which has been acquired automatically by analyzing the definitions in an on-line dictionary. Previously, computational studies of noun sequences made use of hand-coded semantic information, and they applied the analysis rules sequentially. In contrast, the task of analyzing noun sequences in unrestricted text strongly favors an algorithm according to which the rules are applied in parallel and the best interpretation is determined by weights associated with rule applications.

1. INTRODUCTION

The interpretation of noun sequences (henceforth NSs, and also known as noun compounds or complex nominals) has long been a topic of research in natural language processing (NLP) (Finin, 1980; Sparck Jones, 1983; Leonard, 1984; Isabelle, 1984; Lehnert, 1988; and Riloff, 1989). The challenge in analyzing NSs derives from the semantic nature of the problem: their interpretation is, at best, only partially recoverable from a syntactic or a morphological analysis of NSs. To arrive at an interpretation of *plum sauce* which specifies that *plum* is the Ingredient of *sauce*, or of *knowledge representation*, specifying that *knowledge* is the Object of *representation*, requires semantic information for both the first noun (the modifier) and the second noun (the head).

In this paper, we are concerned with interpreting NSs which are composed of two nouns, in absence of the context in which the NS appears; this scope is similar to most of the studies mentioned above. The algorithm for interpreting a sequence of two nouns is intended to be basic to the algorithm for interpreting sequences of more than two nouns: each pair of NSs will be interpreted in turn, and the best interpretation forms a constituent which can modify, or be modified by, another noun or NS (see also Finin, 1980). There is no doubt that

context, both intra- and inter-sentential, plays a role in determining the correct interpretation of a NS, since the most plausible interpretation in isolation might not be the most plausible in context. It is, however, a premise of the present system that, whatever the context is, the interpretation of a NS is always available in the list of possible interpretations. A NS that is already listed in an on-line dictionary needs no interpretation because the meaning can be derived from its definition.

In the studies of NSs mentioned above, the systems for interpreting NSs have relied on hand-coded semantic information, which is limited to a specific domain by the sheer effort involved in creating such a semantic knowledge base. The level of detail made possible by hand-coding has led to the development of two main algorithms for the automatic interpretation of NSs: concept dependent and sequential rule application.

The concept dependent algorithm (Finin, 1980) requires each lexical item to contain an index to the rule(s) which should be applied when that item is part of a NS; it has the advantage that only those rules are applied for which the conditions are met and each noun potentially suggests a unique interpretation. Whenever the result of the analysis is a set of possible interpretations, the most plausible one is determined on the basis of the weight which is associated with a role fitting procedure. The disadvantage of this approach is that this level of lexical information cannot be acquired automatically, and so this approach cannot be used to process unrestricted text.

The algorithm for sequential rule application (Leonard, 1984) focuses on the process of determining which interpretation is the most plausible; the fixed set of rules are applied in a fixed order and the first rule for which the conditions are met results in the most plausible interpretation. This algorithm has the advantage that no weights are associated with the rules. The disadvantage of this approach is that the degree to which the rules are satisfied cannot be expressed, and so, in some cases, the most plausible

interpretation of an NS will not be produced. Also, section 4 will show that this algorithm is suitable only when the sense of each noun is a given, a situation which is not true for processing unrestricted text.

This paper introduces an algorithm which is specifically designed for analyzing NSs in unrestricted text. The task of processing unrestricted text has two consequences: firstly, hand-coded semantic information, and therefore a concept dependent algorithm, is no longer feasible; and secondly, the intended sense of each noun can no longer be taken as a given. The algorithm described here, therefore, relies on semantic information which has been extracted automatically from an on-line dictionary (see Montemagni and Vanderwende, 1992; Dolan et al., 1993). This algorithm manipulates a set of general rules, each of which has an associated weight, and a general procedure for matching words. The result of this algorithm is an ordered set of interpretations and partial sense-disambiguation of the nouns by taking note of which noun senses were most relevant in each of the possible interpretations.

We will begin by reviewing the classification schema for NSs described in Vanderwende (1993) and the type of general rules which this algorithm is designed to handle. The matching procedure will be described; by introducing a separate matching procedure, the rules in Vanderwende (1993) can be organized in such a way as to make the algorithm more efficient. We will then show the algorithm for rule application in detail. This algorithm differs

from Leonard (1984) by applying all of the rules before determining which interpretation is the most plausible (effectively, a parallel rule application), rather than determining the best interpretations by the order in which the rules are applied (a serial rule application). In section 4, we will provide examples which illustrate that a parallel algorithm is required when processing unrestricted, undisambiguated text. Finally, the results of applying this algorithm to a training and a test corpus of NSs will be presented, along with a discussion of these results and further directions for research in NS analysis.

1.1 NS interpretations

Table 1 shows a classification schema for NSs (Vanderwende, 1993) which accounts for most of the NS classes studied previously in theoretical linguistics (Downing, 1977; Jespersen, 1954; Lees, 1960; and Levi, 1978). The relation which holds between the nouns in a NS has conventionally been given names such as Purpose or Location. The classification schema that is used in this system has been formulated as *wh*-questions. A NS 'can be classified *according to which wh-question the modifier (first noun) best answers*' (Vanderwende, 1993). Deciding how a NS should be classified is not at all clear and we need criteria for judging whether a NS has been classified appropriately. The formulation of NS classes as *wh*-questions is intended to provide at least one criterion for judging NS classification; other criteria are provided in Vanderwende (1993).

Table 1. Classification schema for NSs

Relation	Conventional Name	Example
Who/what?	Subject	press report
Whom/what?	Object	accident report
Where?	Locative	field mouse
When?	Time	night attack
Whose?	Possessive	family estate
What is it part of?	Whole-Part	duck foot
What are its parts?	Part-Whole	daisy chain
What kind of?	Equative	flounder fish
How?	Instrument	paraffin cooker
What for?	Purpose	bird sanctuary
Made of what?	Material	alligator shoe
What does it cause?	Causes	disease germ
What causes it?	Caused-by	drug death

1.2 General rules for NS analysis

Each general rule can be considered to be a description of the configuration of semantic and syntactic attributes which provide evidence for a particular NS interpretation, i.e., a NS classification. Exactly how these rules are applied is the topic of this paper. Typically, the general rules correspond in a many-to-one relation to the number of classes in the classification schema because more than one combination of semantic attributes can identify the NS as a member of a particular class. This is illustrated in Table 2, which presents two of the rules for establishing a 'What for?' interpretation.

The first rule (H1) tests whether the definition of the head contains a PURPOSE or INSTRUMENT-FOR attribute which matches (i.e., has the same lemma as) the modifier. When this rule is applied to the NS *bird sanctuary*, the rule finds that a PURPOSE attribute has been identified automatically in the definition of the head: *sanctuary* (L n,3) 'an area for birds or other kinds of animals where they may not be hunted and their animal enemies are controlled'. (All examples are from the Longman Dictionary of Contemporary English, Longman Group, 1978.) The values of this PURPOSE attribute are *bird* and *animal*. The rule H1 verifies that the definition of *sanctuary* contains a PURPOSE attribute, and that one of its values, namely *bird*, has the same lemma as the modifier, the first noun.

The second rule (H2) tests a different configuration, namely, whether the definition of the head contains a LOCATION-OF attribute which matches the modifier; the example *bird cage* will be presented in section 2.

These rules are in a notation modified from Vanderwende (1993, pp.166-7). Firstly, the rules have been divided into those that test attributes on the head, as rules H1 and H2 do, and those that test attributes on the modifier. Secondly, associated with each rule is a weight. Unlike in Vanderwende (1993), this rule weight is only part

of the final score of a rule application; the final score of a rule application is composed of both the rule weight and the weight returned from the matching procedure, which will be described in the next section.

2. THE MATCHING PROCEDURE

Matching is a general procedure which returns a weight to reflect how closely related two words are, in this case how related the value of an attribute is to a given lemma. The weight returned by the matching procedure is added to the weight of the rule to arrive at the score of the rule as a whole. In the best case, the matching procedure finds that the lemma is the same as the value of the attribute being tested. We saw above that in the NS *bird sanctuary*, the modifier *bird* has the same lemma as the value of a PURPOSE attribute which can be identified in the definition of the head, *sanctuary*. The weight associated with such an exact match is 0.5. Applying rule H1 in Table 2 to the NS *bird sanctuary* has an overall score of 1; the match weight 0.5 added to the rule weight 0.5.

When an exact match cannot be found between the lemma and the attribute value, the matching procedure can investigate a match given semantic information for each of the senses of the lemma. (Only in the worst case would this be equivalent to applying each rule to each combination of modifier and head senses.) Of course the HYPERNYM attribute will be useful to find a match. Applying rule H1 to the NS *owl sanctuary*, a match is found between the PURPOSE attribute in the definition of *sanctuary* and the modifier *owl*, because the definition of *owl* (L n,1): 'any of several types of night bird with large eyes, supposed to be very wise', identifies *bird* (one of the values of *sanctuary*'s PURPOSE attribute) as the HYPERNYM of *owl*. Whenever the HYPERNYM attribute is used, the weight returned by the matching procedure is only 0.4.

Table 2. Rules for a 'What for?' interpretation

SENS class	Rule name	Modifier attributes	Head attributes	Example	Weight
What for?	H1	match	PURPOSE, INSTRUMENT-FOR	<i>water heater, bird sanctuary</i>	0.5
	H2	match	LOCATION-OF	<i>bird cage, refugee camp</i>	0.5

Other semantic attributes are also relevant for finding a match. Fig. 1 shows graphically how the attribute HAS-PART can be used to establish a match. One of the ‘Who/What?’ rules tests whether any of the verbal senses of the head has a BY-MEANS-OF attribute which matches the modifier. In the verb definition of *scratch* (L v,1): ‘to rub and tear or mark (a surface) with something pointed or rough, as with claws or fingernails’, a BY-MEANS-OF attribute can be identified with *claw* and *fingernail* as its values, neither of which match the modifier noun *cat*. Now the matching procedure investigates the senses of *cat* attempting to find a match. The definition of *cat* (L n,1): ‘a small animal with soft fur and sharp teeth and claws (nails), ...’ identifies *claw* (one of *scratch*’s BY-MEANS-OF attributes) as one of the values of HAS-PART, thus establishing the match shown in Fig. 1. The weight associated with a match using HAS-PART, PART-OF, HAS-MATERIAL, or MATERIAL-OF is 0.3.

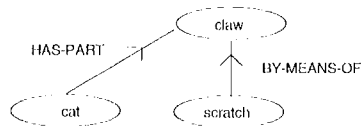


Fig. 1. ‘Who/what?’ interpretation for *cat scratch* with *cat* (L n,1) and *scratch* (L v,1)

Fig. 2 shows how also the attributes HAS-OBJECT and HAS-SUBJECT can be used; this type of match is required when a rule calls for a match between a lemma (which is a noun) and an attribute which typically has a verb as its value, since we can expect no link between a noun and a verb according to hypernymy or any part relation. In the definition of *cage* (L n,1): ‘a framework of wires or bars in which animals or birds may be kept or carried’, a LOCATION-OF attribute can be identified, with as its value the verbs *keep* and *carry* and a nested HAS-OBJECT attribute, with *animal* and *bird* as its value; it is the HAS-OBJECT attribute which can match the modifier noun *bird*. A match using the HAS-OBJECT or HAS-SUBJECT attribute carries a weight of 0.2.

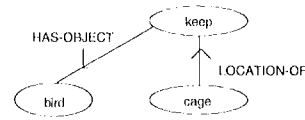


Fig. 2. ‘What for?’ interpretation for *bird cage* with *cage* (L n,1)

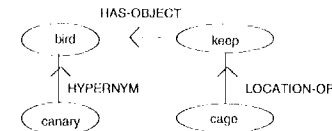


Fig. 2. ‘What for?’ interpretation for *canary cage* with *canary* (L n,1) and *cage* (L n,1)

In Vanderwende (1993), the rules themselves specified how to find the indirect matches described above. By separating the matching information from the information relevant to each rule, the matching can be applied more consistently; but equally important, the rules specify only those semantic attributes that indicate a specific interpretation.

3. ALGORITHM FOR APPLYING RULES

The algorithm controls how the set of general rules will be applied in order to interpret NSs in unrestricted text. Given that a separate procedure for matching exists, the rules are naturally formulated as conditions, in the form of a semantic attribute(s) to be satisfied, on either the modifier or head, but not necessarily on both at the same time. This allows the rules to be divided into groups: modifier-based, head-based, and deverbal-head based. NSs with a deverbal head require additional conditions in the rules; if deverbal-head based rules were applied on par with the head-based rules, the deverbal-head rules would apply far too often, leading to spurious interpretations, because in English nouns and verbs are often homographs.

The algorithm for interpreting NSs has four steps:

1. apply the head-based rules to each of the noun senses of the head and the lemma of the modifier
2. apply the modifier-based rules to each of the noun senses of the modifier and the lemma of the head
3. if no interpretation has received a weight above a certain threshold, then apply the deverbal-head rules to each of the verb senses of the head and the lemma of the modifier
4. order the possible interpretations by comparing the weights assigned by the rule applications and return the list in order of likelihood

The semantic attributes which are found in the head-based conditions are: LOCATED-AT, PART-OF, HAS-PART, HYPERNYM, BY-MEANS-OF, PURPOSE, INSTRUMENT-FOR, LOCATION-OF, TIME, MADE-OF, ROLE, CAUSES and CAUSED-BY. The semantic attributes which are found in the modifier-based conditions are: SUBJECT-OF, OBJECT-OF, LOCATION-OF, TIME-OF, HAS-PART, PART-OF, HYPERNYM, MATERIAL-OF, CAUSES and CAUSED-BY. The semantic attributes in the deverbal-head based conditions are: HAS-SUBJECT, BY-MEANS-OF, and HAS-OBJECT.

In Vanderwende (1993), it was suggested that each rule is applied to each combination of head sense and modifier sense. If the modifier has three noun senses and the head has four noun senses, then each of the 34 general rules would apply to each of the (3x4) possible combinations, for a total of 408 rule applications. With the current algorithm, if the modifier has three noun senses and the head has four noun senses, then first the eleven modifier rules apply (3x11), then the sixteen head rules apply (4x16), and if the head can be analyzed as a deverbal noun, then also the seven deverbal-head rules apply (4x7), for a total of 125 rule applications. Only after all of the rules have applied are the possible interpretations ordered according to their scores.

It may seem that we have made the task of interpreting NSs artificially difficult by taking into consideration each noun sense in the modifier and head; one might argue that it is reasonable to assume that these nouns could be sense-disambiguated before NS analysis. We are

not aware of any study which describes sense-disambiguation of the nouns in a NS. On the contrary, Braden-Harder (1992) suggests that the results of disambiguation can be improved when relations such as verb-object, purpose, and location, are available; these relations are the result of our NS analysis, not the input.

4. PARALLEL VERSUS SERIAL RULE APPLICATION

As we have seen above, the overall score for each possible interpretation is a combination of the weight of a rule and the weight returned by the matching procedure. A rule with a relatively high weight may have a low score overall if the match weight is very low, and a rule with a relatively low weight could have a high overall score if the match weight is particularly high. It is therefore impossible to order the rules a priori according to their weight.

In Leonard (1984), the most plausible interpretation is determined by the order in which the rules are applied. By ordering the 'search for a material modifier' ahead of the 'search for a related verb', the interpretations of both *silver pen* and *ink pen* will be the same, given that both *silver* and *ink* are materials. In fact, only *silver pen* is correctly analyzed by the 'search for a material modifier' rule, while the correct interpretation of *ink pen* would have used the 'search for a related verb'.

The problem with rule ordering is compounded when more than one sense of each noun is considered. In Leonard's lexicon, *pen*[1] is the writing implement and *pen*[2] is the enclosure for keeping animals in. By ordering a 'search for a related verb' ahead of a 'search for a locative', the interpretation of the NS *bull pen* is incorrect: '*a pen*[1] *that a bull or bulls writes something with*'. Less likely is the correct locative interpretation '*a pen*[2] *for or containing a bull or bulls*'.

In our system, the most likely interpretations of *bull pen* are ordered correctly because, for the locative interpretation, we find meaningful matches in the definitions of *bull* and *pen*: the definition of *pen* (L n,1): 'a small piece of land enclosed by a fence, used esp. for keeping animals in', identifies a PURPOSE attribute, with the verb *keep* and a nested HAS-OBJECT *animal* as its values. The HAS-OBJECT *animal* can be matched with the modifier lemma *bull*, because one of the HYPERNYMs of *bull* (L n,2) is *animal*. For the related verb interpretation,

however, we find no match between the typical subjects the verb related to *pen*, namely *write*, and the modifier *bull*; a ‘Who/What?’ interpretation is only possible because *bull* is an animate, and, by default, animates can be the subject of a verb.

We must conclude that what is important is the degree to which there is a match between the values of these attributes and the lemma, and not merely the presence or absence of semantic attributes. Only after all of the rules have been applied can the most plausible interpretation be determined.

5. TEST, RESULTS AND DISCUSSION

The results that are under discussion were obtained on the basis of semantic information which was automatically extracted from Longman Dictionary of Contemporary English (LDOCE) as described in Montemagni and Vanderwende (1992)¹. The semantic information has not been altered in any way from its automatically derived form, and so there are still errors: for the 94,000 attribute clusters extracted from nearly 75,000 single noun and verb definitions in LDOCE, we estimate the accuracy to be 78%, with a margin of error of $\pm 5\%$ (see Richardson et al., 1993).

A training corpus of 100 NSs was collected from the examples of NSs in the previous literature, to ensure that all known classes of NSs are handled in this system. These results were expected to be good because these NSs were used to develop the rules and their weights. The system successfully identified the most likely interpretation for 79 of the 100 NSs (79%). Of the remaining 21 NSs, the most plausible interpretation was among the possible interpretations 8 times, (8 %), and no interpretation at all was given for 4 NSs (4 %).

The test corpus consisted of 97 NSs from the tagged version of the Brown corpus (Francis and Kucera, 1989), to ensure the adequacy of applying this approach to unrestricted text; the results for an expanded test corpus will be reported in Vanderwende (in preparation). The system currently identified successfully the most likely interpretation for 51 of the 97 NSs (52%). Of the remaining 46 NSs, the most likely interpretation was presented second for 21 NSs

(22 %); when first and second interpretations were considered, the system was successful approximately 74% of the time. A wrong or no interpretation was given for 25 NSs (26 %). Upon examination of these results, several areas for improvement are suggested. First is to improve the semantic information: Dolan et al. (1993) describes a network of semantic information, given not only the definition of the lexical entry but also all of the other definitions which have a labeled relation to that entry.

Secondly, while the NS classification proposed in Vanderwende (1993) proves to be adequate for analyzing NSs in unrestricted text, an additional ‘What about?’ class, suggested in Levi (1978), may be justified. In the current classification schema, NSs such as *cigarette war* and *history conference* have been considered ‘Whom/what?’ NSs given the verbs that are associated with the head, *fight* and *confer/talk about* respectively. In unrestricted text, similar NSs are quite frequent, for example *university policy*, *prevention program*, *care plan*, but the definitions of the heads do not always specify a related verb. The head definitions for *policy*, *program* and *plan*, however, do allow a HAS-TOPIC semantic feature to be identified, and this HAS-TOPIC can be used to establish a ‘What about?’ interpretation.

Applying this algorithm to previously unseen text also produced a very promising result: the verbs that are associated with nouns in their definitions (i.e., *role nominals* in Levin, 1980) are being used often and correctly to produce NS interpretations. While some rules had been developed to handle obvious cases in the training corpus, how often the conditions on these rules would be met could not be predicted. In fact, such NS interpretations are frequent. For example, the NS *wine cellar* is analyzed as a ‘What for?’ relation with a high score, and the system provides as a paraphrase: *cellar which is for storing wine*, given the definition of *cellar* (L n,1): ‘an underground room, usu. used for storing goods; basement’. This result is promising for two reasons: first, by analyzing the definitions (and later also the example sentences) in an on-line dictionary, we now have access to a non-handcoded source of semantic information which includes the verbs and their relation to the nouns, essential for determining role nominals. Second, the related verbs are used to construct the paraphrases of a NS, and doing so makes a general interpretation such as ‘What for?’ more

¹Although LDOCE includes some semantic information in the form of box codes and subject codes, these were not used in this system. This approach is designed to work with semantic information from any dictionary.

specific, e.g., a *service office* is not an *office for service*, but an *office for providing service*, and a *vegetable market* is not a *market for vegetables*, but a *market for buying and selling vegetables*. Enhancing the general interpretations with the related verb(s) approximates at least in part Downing's observation that the types of relations that can hold between the nouns in a NS are possibly infinite (Downing, 1977).

6. CONCLUSIONS

Our goal is to create a system for analyzing NSs automatically on the basis of semantic information extracted automatically from on-line resources; this is a strong requirement for NLP as it moves its focus away from technical sublanguages towards the processing of unrestricted text. We have shown that processing NSs in unrestricted text strongly favors an algorithm which is comprised of a set of general rules and a general procedure for matching two words, each of which have associated weights. This algorithm must apply all of the rules before the most plausible NS interpretation can be determined.

Several directions for further research can be pursued within this approach: a methodology for automatically assigning the weights associated with the rules and the matching procedure, following Richardson (in preparation), and a methodology for incorporating the context of the NS into the analysis of NS interpretations. We are also pursuing the acquisition of semantic information which is not already available, along the same lines as extracting information automatically from on-line dictionaries.

Acknowledgements: I would like to extend my thanks to the members of the Microsoft NLP group: George Heidorn, Karen Jensen, Bill Dolan, Joseph Penetheroudakis, Diana Peterson, and Stephen Richardson.

REFERENCES

- Braden-Harder, L.C. (1991). "Sense Disambiguation using on-line dictionaries: an approach using linguistic structure and multiple sources of information". New York University, NY.
- Dolan, W.B., L. Vanderwende and S.D. Richardson (1993). Automatically deriving structured knowledge bases from on-line dictionaries. *Proceedings of the First Conference of the Pacific Association for Computational Linguistics*, at Simon Fraser University, Vancouver, BC., pp.5-14.
- Downing, P. (1977). On the creation and use of English compound nouns. *Language* 53, pp.810-842.
- Finin, T.W. (1980). "The semantic interpretation of compound nominals". University of Illinois at Urbana-Champaign. University Microfilms International.
- Isabelle, P. (1984). Another look at nominal compounds. *Proceedings of the 22nd Annual Meeting of the Association for Computational Linguistics, COLING-84*, Stanford, CA. pp.509-516.
- Jespersen, O. (1954). *A modern English grammar on historical principles, VI*. George Allen & Unwin Ltd., London, 1909-49; reprinted 1954.
- Lees, R.B. (1960). *The grammar of English nominalizations*. Indiana University, Bloomington, Indiana, (Fourth printing 1966).
- Lehnert, W. (1988). The analysis of nominal compounds, In U. Eco, M. Santambrogio, and P. Violi Ed., *Meaning and Mental Representations*, VS 44/45, Bompiani, Milan.
- Leonard, R. (1984). *The interpretation of English noun sequences on the computer*. North-Holland Linguistic Studies, Elsevier, the Netherlands.
- Levi, J.N. (1978). *The syntax and semantics of complex nominals*. Academic Press, New York.
- Montemagni, S. and L. Vanderwende (1992). Structural Patterns vs. string patterns for extracting semantic information from dictionaries. *Proceedings of COLING92*, Nantes, France, pp.546-552.
- Richardson, S.D., L. Vanderwende and W.B. Dolan (1993). Combining dictionary-based methods for natural language analysis. *Proceedings of TMI-93*, Kyoto, Japan. pp.69-79.
- Riloff, E. (1989). "Understanding gerund-noun pairs". University of Massachusetts, MA.
- Sparck Jones, K. (1983). So what about parsing compound nouns? In K. Sparck Jones and Y. Wilks Ed., *Automatic Natural Language Parsing*, Ellis Horwood, Chichester, England, pp.164-8.
- Vanderwende, L. 1993. SENS: the system for evaluating noun sequences. In K. Jensen, G.E. Heidorn and S.D. Richardson Ed., *Natural Language Processing: the PLNLP approach*, Kluwer Academic Publishers, pp. 161-73.
- Vanderwende, L. (in prep.). "The analysis of noun sequences using semantic information extracted from on-line dictionaries". Georgetown University, Washington, D.C.