

AD-764 652

PREFERENCE SEMANTICS

Yorick Wilks

Stanford University

Prepared for:

Advanced Research Projects Agency

July 1973

DISTRIBUTED BY:

NTIS

National Technical Information Service
U. S. DEPARTMENT OF COMMERCE
5285 Port Royal Road, Springfield " " 22151

**BEST
AVAILABLE COPY**

STAN-CS-73-377

PREFERENCE SEMANTICS

BY

YORICK WILKS

SUPPORTED BY

ADVANCED RESEARCH PROJECTS AGENCY

ARPA ORDER NO. 457

JULY 1973

COMPUTER SCIENCE DEPARTMENT

School of Humanities and Sciences

STANFORD UNIVERSITY

DDC
RECEIVED
AUG 15 1973
RECEIVED
E



DOCUMENT CONTROL DATA - R & D

(Security classification of title, body of abstract and indexing annotation must be entered when the overall report is classified)

1. ORIGINATING ACTIVITY (Corporate author) Stanford University Computer Science Department Stanford, California 94305		2a. REPORT SECURITY CLASSIFICATION Unclassified	
		2b. GROUP	
3. REPORT TITLE PREFERENCE SEMANTICS			
4. DESCRIPTIVE NOTES (Type of report and inclusive dates) technical report, July 1973			
5. AUTHOR(S) (First name, middle initial, last name) Yorick Wilks			
6. REPORT DATE July 1973		7a. TOTAL NO. OF PAGES 20 24	7b. NO. OF REFS 18
8a. CONTRACT OR GRANT NO. SD-183		9a. ORIGINATOR'S REPORT NUMBER(S) STAN-CS-73-377	
b. PROJECT NO.		AIM-206	
c.		9b. OTHER REPORT NO(S) (Any other numbers that may be assigned this report)	
d.			
10. DISTRIBUTION STATEMENT Releasable without limitations on dissemination.			
11. SUPPLEMENTARY NOTES		12. SPONSORING MILITARY ACTIVITY	
13. ABSTRACT Preference semantics [PS] is a set of formal procedures for representing the meaning structure of natural language, with a view to embodying that structure within a system that can be said to understand, rather than within what I would call the "derivational paradigm", of transformational grammar [TG] and generative semantics [GS], which seeks to determine the well-formedness, or otherwise, of sentences. I outline a system of preference semantics that does this: for each phrase or clause of a complex sentence, the system builds up a network of lexical trees with the aid of structured items called templates and, at the next level, it structures those networks with higher level items called paraplates and common-sense inference rules. At each stage the system directs itself towards the correct network by always opting for the most "semantically dense" one it can construct. I suggest that this opting for the "greatest semantic density" can be seen as an interpretation of Joos' "Semantic Axiom Number 1". I argue that the analysis of quite simple examples requires the use of inductive rules of inference which cannot, theoretically cannot, be accommodated within the derivational paradigm. I contrast this derivational paradigm of language processing with the artificial intelligence [AI] paradigm.			

JULY 1973

COMPUTER SCIENCE DEPARTMENT
REPORT NO. CS-377

PREFERENCE SEMANTICS

by

Yorick Wilks

ABSTRACT:

Preference semantics [PS] is a set of formal procedures for representing the meaning structure of natural language, with a view to embodying that structure within a system that can be said to understand, rather than within what I would call the "derivational paradigm", of transformational grammar [TG] and generative semantics [GS], which seeks to determine the well-formedness, or otherwise, of sentences. I outline a system of preference semantics that does this: for each phrase or clause of a complex sentence, the system builds up a network of lexical trees with the aid of structured items called templates and, at the next level, it structures those networks with higher level items called paraplates and common-sense inference rules. At each stage the system directs itself towards the correct network by always opting for the most "semantically dense" one it can construct. I suggest that this opting for the "greatest semantic density" can be seen as an interpretation of Joos' "Semantic Axiom Number 1". I argue that the analysis of quite simple examples requires the use of inductive rules of inference which cannot, theoretically cannot, be accommodated within the derivational paradigm. I contrast this derivational paradigm of language processing with the artificial intelligence [AI] paradigm.

This paper was presented at the Colloquium on the Formal Semantics of Natural Language, Cambridge, April 1973, and will appear in the Proceedings, (ed. E. Keenan), Cambridge University Press, 1974.

The views and conclusions contained in this document are those of the author and should not be interpreted as representing necessarily the official policies, either expressed or implied, of the Advanced Research Projects Agency or the U. S. Government.

This research was supported by the Advanced Research Projects Agency, Department of Defense (SD 183), USA.

Reproduced in the USA. Available from the National Technical Information Service, Springfield, Virginia, 22151.

Introduction

In this paper I want to oppose a method of semantic analysis to the contemporary paradigm. By that I mean the transformational grammar(TG)-generative semantics(GS) one, rather than recent developments in modal logic and set theory. It seems to me that attacking the claims of the latter about natural language may be fun, but it is not a pressing matter in the way that criticising GS is. For GS has gone so far in the right direction, towards a system for understanding natural language adequately, that perhaps with one more tiny tap the whole carapace of the "derivational paradigm" might burst.

What I intend by that phrase is the picture of language imported into linguistics from proof theory by Chomsky. Both TG and GS claim to devote themselves to the production of bodies of rules that would perform repeated derivations, and so pass from some initial symbol to an ultimate surface string, that is to say a well-formed sentence. The field of all possible derivations with such a body of rules is taken to define the class of well-formed sentences of the language in question: those that can be produced by derivation in that way are "well formed", those that cannot are not. This description is not, in its essentials, in dispute.

I have argued elsewhere(Wilks 1971) that there are good abstract reasons for thinking that this sorting cannot, even in principle, be done: at least not if the task is taken to be one of sorting the meaningful sentences of a language from the meaningless ones. The reason is simply that, given any disputed utterance, we could not know formally of it that it was not meaningful, because speakers have the ability to embed odd-looking utterances in stories so as to make them meaningful in the context of use. However, even if this gigantic sorting task could be done, it has no connexion whatever with Lakoff's recently expressed (1973) desideratum for GS, as opposed to TG, that it should take "into account the fact that language is used by human beings to communicate in a social context". And no generative linguist to my knowledge, whether of the TG or GS persuasion, has ever unambiguously rejected Chomsky's original sorting-by-derivation as the central task of linguistic theory. In this paper I want to argue that there at least two sorts of example, quite simple examples, that cannot be analysed adequately within the derivational paradigm. To do so, I describe a non-standard system of semantic analysis that can deal with such examples. I describe the system in a rough and ready way, with nothing like an adequate justification, or motivation as the fashionable word is, of its primitive elements and assumptions. Linguists who dislike non-standard systems, and are prone to consider them "unmotivated" per se should skip immediately to the discussion section so as not to miss the substantive point of this paper.

The preference semantics(PS) system I shall describe is at present functioning as part of an analysis and generation system for natural language programmed on a computer (see Wilks 1973 a&b): one with no independent syntax base, everything being handled through the strong semantic representation described. This, I argue, provides an additional argument for its adequacy in handling natural language, over and above the mere labelling of examples. I assume, too, that such a system cannot be dismissed as "mere performance": partly because, as I shall show, it explicates real competences of human understanders inadequately treated in TG/GS systems; and in part because the intellectual weight of the "competence-performance" distinction is insufficient to dismiss systems that merely differ from the conventional TG/GS paradigm.

A. outline of preference semantics.

A fragmented text is to be represented by an interlingual structure consisting of TEMPLATES bound together by PARAPHRASES and Common Sense[CS] INFERENCES. These three items consist of FORMULAS (and predicates and functions ranging over them and over sub-formulas), which in turn consist of ELEMENTS.

ELEMENTS are sixty primitive semantic units used to express the semantic entities, states, qualities and actions about which humans speak and write. The elements fall into five classes, which can be illustrated as follows (elements in upper case):

(a) entities: MAN(human being), STUFF(substances), THING(physical object), PART(parts of things), FOLK(human groups), ACT(acts), STATE(states of existence), BEAST(animals), etc. (b) actions: FORCE(compels), CAUSE(causes to happen), FLOW(moving as liquids do), PICK(choosing), BE(exists) etc. (c) type indicators: KIND(being a quality), HOW(being a type of action) etc. (d) sorts: CONT(being a container), GOOD(being morally acceptable), THRU(being an aperture) etc. (e) cases: TO(direction), SOUR(source), GOAL(goal or end), LOCA(location), SUBJ(actor or agent), OBJE(patient of action), IN(containment), POSS(possessed by) etc.

FORMULAS are constructed from elements and right and left brackets. They express the senses of English words; one formula to each sense. The formulas are binarily bracketed lists of whatever depth is necessary to express the word sense. They are written and interpreted with, in each pair at whatever level it comes, a dependence of left side on corresponding right, and thus the right most element of the whole formula is its principal element, called its HEAD. Formulas can be thought of, and written out, as binary trees of semantic primitives, and in that form they are not unlike the lexical decomposition trees of Lakoff and McCawley.

Consider the action "drink" and its relation to the formula(or binary tree):

```
(((*ANI SUBJ) (((FLOW STUFF) OBJE) ((*ANI IN) (((THIS(*ANI (THRU PART))) TO) (BE CAUSE))))))
```

*ANI here is simply the name of a class of elements, those expressing animate entities namely, MAN, BEAST and FOLK(human groups). In order to keep a small usable list of semantic elements, and to avoid arbitrary extensions of the list, many notions are coded by conventional sub-formulas: so, for example, (FLOW STUFF) is used to indicate liquids, and (THRU PART) is used to indicate apertures.

Let us now decompose the formula for "drink". It is to be read as an action, preferably done by animate things (*ANI SUBJ) to liquids((FLOW STUFF)OBJE), of causing(CAUSE being the head of the formula) the liquid to be in the animate thing(*ANI IN) and via (TO indicating the direction case) a particular aperture of the animate thing; the mouth of course. It is hard to indicate a notion as specific as "mouth" with such general concepts. But it would be simply irresponsible, I think, to suggest adding MOUTH as a semantic primitive, as do semantic systems that simply add an awkward lexeme as a new "primitive". Lastly, the THIS indicates that the part is a specific part of the subject.

The notion of preference is the important one here: SUBJ case displays the preferred agents of actions, and OBJE case the preferred objects, or patients. We cannot enter such preferences as stipulations, as many linguistic systems do, such as that of Katz and Postal (1964) with "selection restrictions", where, if such a restriction is violated, the result is "no reading". For we can be said to drink in the atmosphere, and cars are said to drink gasoline. It is proper to prefer the normal, but it would be absurd, in an intelligent understanding system, not to accept the abnormal if it is described. Not only everyday metaphor, but the description of the simplest fictions, require it.

Just as elements are to be explained by seeing how they functioned within formulas, so formulas, one level higher, are to be explained by describing how they function within TEMPLATES, the third kind of semantic item in the system. The notion of a template is intended to correspond to an intuitive one of message: one not reducible merely to unstructured associations of word-senses.

A template consists of a network of formulas grounded on a basic actor-action-object triple of formulas. This basic formula triple is located during initial analysis in frames or formulas, one formula for each fragment word in each frame, by means of a device called a bare template. A bare template is simply a triple of elements which are the heads of three formulas in actor-action-object form.

For example: "Small men sometimes father big sons", when represented by a string of formulas, will give the two sequences of heads:

KIND MAN HOW MAN KIND MAN

and

KIND MAN HOW CAUSE KIND MAN.

(CAUSE is the head of the verbal sense of "father"; "to father" is analyzed as "to cause to have life".)

The first sequence has no underlying template; however, in the second we find MAN CAUSE MAN which is a legitimate bare template. Thus we have disambiguated "father", at the same time as picking up a sequence of three formulas which is the core of the template for the sentence. It must be emphasised here that the template is the sequence of formulas, and not to be confused with the triple of elements(heads) used to locate it.

It is a hypothesis of this work that we can build up a finite but useful inventory of bare templates adequate for the analysis of ordinary language: a list of the messages that people want to convey at some fairly high level of generality (for template matching is not in any sense phrase-matching at the surface level). The bare templates are an attempt to explicate a notion of a non-atomistic linguistic pattern: to be located whole in texts in the way that human beings appear to when they read or listen.

Let me avoid all further questions of analysis in order to illustrate the central processes of expansion and preference by considering the sentence :

[1] "The big policeman interrogated the crook".

Let us take the following formulas for the four main word senses:

(1) "policeman": ((FOLK SOUR) (((NOTGOOD MAN)OBJE)PICK) (SUBJ MAN)))

i. e. a person who selects bad persons out of the body of people(FOLK) The case marker SUBJ is the dependent in the last element pair, indicating that the normal "top first" order for subject-entities in formulas has been violated, and necessarily so if the head is also to be the last element in linear order.

(2) "big": ((*PHYSOB POSS) (MUCH KIND))

i. e. a property preferably possessed by physical objects(substances are not big)

(3) "interrogates": ((MAN SUBJ)((MAN OBJE)(TELL FORCE))) i.
e. forcing to tell something, done preferably by humans, to humans.

(4a) "crook": (((NOTGOOD ACT)OBJE)DO)((SUBJ MAN))

i. e. a man who does bad acts. And we have to remember here that we are ignoring other senses of "crook" at the moment, such as the shepherd's

(4b) crook: ((((((THIS BEAST)OBJE)FORCE) (SUBJ MAN))POSS) (LINE THING))

i. e. a long straight object possessed by a man who controls a particular kind of animal.

The analysis algorithm will have seen the sentence as a frame of formulas, one for each of its words, and will look only at the heads of the formulas. Given that MAN FORCE MAN is in the inventory of bare templates, then one scan of a frame of formulas (containing formula (4a) for "crook"), will have picked up the sequence of formulas labelled above 1 3 4a, in that order. Again when a frame containing formula (4b), the shepherd's sense of "crook", is scanned, the sequence of formulas 1 3 4b will also be selected as a possible initial structure for the sentence, since MAN FORCE THING is also a proper bare template sequence.

We now have two possible template representations for the sentence after the initial match; both a triple of formulas in actor-action-object form. Next, the templates are expanded, if possible. This process consists of extending the simple networks we have so far; both by attaching other formulas into the network, and strengthening the bonds between those already in the template, if that is possible. Qualifier formulas can be attached where appropriate, and so the formula 2 (for "big") is tied to that for "policeman" in both templates. But now comes a crucial difference between the two representations, one which will resolve the sense of "crook".

The expansion algorithm looks into the formulas expressing preferences and sees if any of the preferences are satisfied: as we saw formula 2 for "big" prefers to qualify physical objects. A policeman is such an object and that additional dependency is marked in both templates; similarly for the preference of "interrogate" for human actors, in both representations. The difference comes with preferred objects: only the formula 4a for human crooks can satisfy that preference, the formula 4b, for shepherd's crooks, cannot. Hence the former template network is denser by one dependency, and is preferred over the latter in all subsequent processing: its connectivity (using numbers for the corresponding formulas, ignoring the "the"s, and using one arrow for each dependency established) is:

2→1→3→4a

and so that becomes the template for this sentence. The other possible template was connected as follows:

2→1→3→4b

and it is now discarded.

Thus, the sub-formulas that express preferences both express the meaning of the corresponding word sense, and can also be interpreted as implicit procedures for the construction of correct templates. This preference for the greatest semantic density works well, and can be seen as an expression of what Joos(1971) calls "semantic axiom number 1", that the right meaning is the least meaning, or what Scriven(1972)\$\$\$ has called "the trick, in meaning analysis, of creating redundancies in the input". This uniform principle works over both the areas that are conventionally distinguished in linguistics as syntax and semantics. There is no such distinction in this system, since all manipulations are of formulas and templates, constructed out of elements of a single type.

As a simple example of linguistic syntax, done by preference simply to illustrate the general principle, let us take the sentence :

[2] "John gave Mary the book",

onto which the matching routine will have matched two templates with heads as follows, since it has no reason so far to prefer one to the other:

John gave Mary the book

MAN-GIVE-----THING

MAN-GIVE-MAN

The expansion routine now seeks for dependencies between formulas, in addition to those between the three formulas constituting the template itself. In the case of the first, a GIVE action can be expanded by any substantive formula to its immediate right which is not already part of the template, (which is to say that indirect object formulas can depend on the corresponding action formula.)Again "book" is qualified by an article, which fact is not noticed by the second template. So then, by expanding the first template we have established in the following dependencies at the surface level, where the dependency arrows "→" correspond to relations established between formulas for the words they link.

\$\$\$

It has been pointed out to me that Quillian's work(1968) does contain a preferential metric, and Dr.J.Hurford(private comm.) has told me that he has been forced to some preference principle in studying derivations of number structures from a linguistic base.

John → gave ← book ← the
 ↑
 Mary

Where, for the present purpose, we are omitting any indication by arrow of the preference of "give" for a human agent, because it occurs equally in both expansions. Now, if we try to expand the second template by the same method, we find we cannot, because the

formula for "Mary" cannot be made dependent on the one for "give", since in that template "Mary" has already been seen, wrongly of course, as a direct object of giving, and it cannot be an indirect object as well. So then, the template with heads MAN GIVE MAN cannot be expanded to yield any dependency arcs connecting formulas to the template; whereas the template with heads MAN GIVE THING yields two dependency arcs on expansion, and so corresponds to the preferred representation. This method can yield virtually all the results of a conventional grammar, while using only relations between semantic elements.

The limitation of the illustrative examples, so far, has been that they are the usual short example sentences, whereas what we actually have here is a general system for application to paragraph length texts. I will now sketch in, for two sorts of case, how the system deals with non-sentential text fragments with a general template format.

In the actual implementation of the system, as an analysis system, an input text is initially fragmented, and templates are matched with each fragment of the text. The input routine partitions paragraphs at the occurrence of any of an extensive list of KEY words. The list contains almost all punctuation marks, subordinations, conjunctions and prepositions. In difficult cases, described in detail in (Wilks 1972), fragmentations are made even though a key word is not present, as at the stroke in "John knows/Mary loves him", while in other cases a fragmentation is not made in the presence of a key word, such as "that" in "John loves that woman".

Let us consider the sentence "John is / in the house", fragmented into two parts at the point marked by the stroke. It should be clear that the three part template, of standard agent-act-action form, cannot be matched onto the fragment "John is". In such a case, a degenerate template with heads MAN BE DTHIS is matched onto the two items of this sentence; the last item DTHIS being a dummy object, indicated by the D.

With the second fragment "in the house" a dummy subject DTHIS fills out the form to give a degenerate template with heads DTHIS PBE POINT. The PBE is the same as the head of the formula for "in", since formulas for prepositions are assimilated to those for actions and have the head PDO or PBE. The fact that they originate in a preposition is indicated by the P, so distinguishing them from

straightforward action formulas with heads DO and BE. POINT is the head of the formula for "house", so this bare template triple for the fragment only tells us that "something is at a point in space". At a later stage, after the preliminary assignment of template structures to individual fragments, TIE routines attach the structures for separated fragments back together. In that process the dummies are tied back to their antecedents. So, in "John is in the house", the DTHIS in the MAN BE DTHIS template for the first fragment of the sentence, ties to the whole template for the second fragment, expressing where John is.

It is very important to note that a preference is between alternatives: if the only structure derivable does NOT satisfy a declared preference, then it is accepted anyway. Only in that way can we deal naturally with metaphor.

So, in examples like :

[3] "I heard an earthquake/singing /in the shower"

(with fragmentation as indicated by slashes), as contrasted (given that the fragmentation program is sensitive to "ing" suffixes, with:

[4] "I heard /an earthquake sing/in the shower"

we shall expect, in the first case, to derive the correct representation because of the preference of notions like singing for animate agents. This is done by a simple extension of the density techniques discussed to relations between structures for different fragments (the TIE routines), in this case, by considering alternative connectivities for dummy parts of templates.

Thus, there will be a dummy subject and object template for /singing/, namely DTHIS CAUSE DTHIS, based on the formula

"singing": ((*ANI SUBJ) ((SIGN OBJE) (((MAN SUBJ) SENSE: CAUSE))))

which is to say, an act by an animate agent of causing a human to experience some sign (i. e. the song)

Now the overall density will be greater when the agent DTHIS, in the template for "singing", is tied to a formula for "I" in a preceding template, than when it is tied to one for "earthquake", since only the former satisfies the preference for an animate agent, and so the correct interpretation of the whole utterance is made.

But, and here we come to the point of this example, in the second sentence, with "sing", no such exercise of preference is possible, and the system must accept an interpretation in which the earthquake sings, since only that can be meant.

Other kinds of preference are of objects for certain preferred functions: thus it is expressed in the formula for book that its preferred function is to be read. Also, adjective qualifiers express preferred kinds of entity to qualify. Thus "big" has a formula expressing a preference for qualifying objects, and so in the expansion of a representation for "The big glass is green" we would get a denser, and so preferred, structure for the object, rather than for the substance, glass.

In order to give a rough outline of the system, I have centred upon the stages of analysis within the individual fragment. After the application of the routines described so far, TIE routines are applied again to the expanded templates in a wider context: the same techniques of expansion, dependency and preference are applied between full templates for different fragments of a sentence or paragraph. At that stage, (1) case ties are applied (using the same cases as occur within formulas at a lower level); (2) the equivalence of active and passive forms is established; (3) dummies are attached to "what they stand for" as I indicated with the "earthquake example"; and, importantly, (4) anaphoric ties are settled.

The TIE routines apply PARAPLATES to the template codings, using the same density techniques one level further up, as it were. Paraplates have the general form:

<list of predicates><list of generation items and functions><list of template predicates>

An ordered list of paraplates is attached to English key words. Consider the following three schematic paraplates for "in":

((20BCAS INST GOAL) (PRMARK *DO) IN (into) (FN1 CONT THING) (PCASE *DIRE))

((PRMARK *DO) IN (into (FN1 CONT THING) (PCASE *DIRE))

((20BHEAD NIL) (PRMARK *DO) IN (make part) (PCASE LOCA))

*DIRE is a direction case marker (covering two sub-cases: TO, mentioned above, and FROM), 20BCAS and 20BHEAD are simply predicates that look at both the object (third) formulas of the template in hand, and of the preceding template, i. e. at two objects. 20BHEAD is true iff the two have the same head, and 20BCAS is true iff they contain the same GOAL or INSTRUMENT subformula. The predicates like PRMARK are satisfied iff the representation of the fragment's mark (the text item on which the fragment depends under the corresponding interpretation: "put" in this case) is an action whose head is in the class of elements *DO, a wide class covering the majority of actions including "putting". The lower case words simply explain which sense of "in" is the one appropriate to the paraplate in which it occurs. When the system is functioning as a translator these generation items will in this case be different

French prepositions, to be generated when the corresponding paraplate "fits". The general result after a paraplate has fitted is that two templates have been linked by a correct case tie: the case that is the argument of the "result predicate" PROCASE.

Now consider the sentence

[4] "I put the key / in the lock"

, fragmented at the stroke as shown. Let us consider that two templates have been set up for the second fragment: one for "lock" as a fastener, and one for the raising lock on a canal. Both formulas may be expected to refer to the containment case. We apply the first paraplate and find that it fits only for the template with the correct (fastener) sense of "lock", since only there will PROCASE be satisfied, i. e. where the formulas for "lock" and "key" both have a subformula under GOAL indicating that their purpose is to close something. The second paraplate will fit with the template for the canal sense of "lock", but the first is a more extensive fit (indicated by the order of the paraplates, since the higher up the paraplate list, the more non-trivial template functions a paraplate contains) and is preferred. This preference has simultaneously selected both the right template for the second fragment and the correct paraplate linking the two templates for further generation tasks.

If we now take the sentence

[5] "He put the number / in the table"

, with two different templates for the second fragment (corresponding to the list and flat object senses of "table" respectively) we shall find that the intuitively correct template (the list sense) fails both the first paraplate and the second, but fits the third, thus giving us the "make part of" sense of "in", and the right (list) sense of "table", since formulas for "number" and (list) "table" have the same head SIGN, though the formula for (flat, wooden) "table" does not.

Conversely, in the case of

[6] "He put the list / in the table"

, fitting the correct template with the second paraplate will yield "into" sense of "in" (case DIRECTION) and the physical object sense of "table"; and this will be the preferred reading, since the fit (of the incorrect template) with the third paraplate yields the "make part of a list" reading in this case. Here we see the fitting of paraplates, and choosing the densest preferential fit, which is always selecting the highest paraplate on the list that fits, thus determining both word sense ambiguity and the case ambiguity of prepositions at once. Paraplate fitting makes use of the deeper

nested parts (essentially the case relations other than SUBJ and OBJE) of the formulas than does the template matching.

The TIE routines also deal with simple cases of anaphora on a simple preference basis. In cases such as

[7] "I bought the wine, /sat on a rock/and drank it"

, it is easy to see that the last word should be tied by TIE to "wine" and not "rock". This matter is settled by density after considering alternative ties for "it", and seeing which yields the denser representation overall. It will be "wine" in this case since "drink" prefers a liquid object.

In more complex cases of anaphora, that require access to more information than is contained in formulas, templates or paraplates, the system brings down what I referred to earlier as common-sense (CS) inference rules. Cases that require them will be ones like the sentence :

[8] "The soldiers fired at the women and I saw several of them fall"

Simple semantic density considerations in TIE are inadequate here because both soldiers and women can fall equally easily, yet making the choice correctly is vital for a task like translation because the two alternatives lead to differently gendered pronouns in French. In such cases the PS system applies a CS rule, whose form, using variables and sub-formulas, would be $X(((NOTPLEASE (LIFE STATE))OBJE)SENSE) \rightarrow X(NOTUP MOVE)$. For rough expository purposes such a rule is probably better expressed as $X[hurt] \rightarrow X[fall]$, where the words in square parentheses correspond informally to the subformulas in the rule. The rules are applied to "extractions" from the situations to form chains, and a rule only ultimately applies if it can function in the shortest, most-preferred, chain.

The way the CS inferences work is roughly as follows: they are called in at present only when TIE is unable to resolve outstanding anaphoras, as in the present example. A process of extraction is then done and it is to these extractions, and the relevant templates, that the CS rules subsequently apply. The extractions are quasi-inferences from the deep case structure of formulas. So for example, if we were extracting from the template for "John drank the water", unpicking the formula "or "water" given earlier would extract that some liquid was now inside an animate thing (from the containment case), and that it went in through an aperture of the animate thing (from the directional case). Moreover, since the extractions are partially confirmed, as it were, by the information about actor and object in the surrounding template, we can, by simple tying of variables, extract new quasi-templates equivalent to, in ordinary language, "the water is in John" etc. These and (when in coded form) the extractions to which the CS rules apply as it endeavors to build up a chain of extractions and inferences. The preferred

chain will, unsurprisingly, be the shortest.

So then, in the "women and soldiers" example we extract a coded form, by variable tying in the templates, equivalent to $\{women\}\{hurt\}$, since we can tell from the formula for "fired at" that it is intended to hurt the object of the action. We are seeking for partial confirmation of the assertion $X? \{fall\}$, and such a chain is completed by the rule given, though not by a rule equivalent to, say, $X\{hurt\} \rightarrow X\{die\}$, since there is nothing in the sentence as given to partially confirm that rule in a chain, and cause it to fit here. Since we are in fact dealing with subformulas in the statement of the rules, rather than words, "fitting" means an "adequate match of subformulas".

It is conceivable that there would be an, implausible, chain of rules and extractions giving the other result, namely that the soldiers fall: $\{soldiers\}\{fire\}; X\{fire\} \rightarrow X\{firedat\} \rightarrow X\{hurt\}$ etc. But such a chain would be longer than the one already constructed and would not be preferred.

The most important aspect of this procedure is that it gives a rationale for selecting a preferred interpretation, rather than simply rejecting one in favor of another, as other systems do (see discussion below). It can never be right to reject another interpretation irrevocably in cases of this sort, since it may turn out later to be correct, as if the "women" sentence above had been followed by "And after ten minutes hardly a soldier was left standing". After inputting that sentence the relevant preferences in the example might be expected to change. Nonetheless, the present approach is not in any way probabilistic. In the case of someone who utters the "soldiers and women" example sentence, what he is to be taken as meaning is that the women fell. It is of no importance in that decision if it later turns out that he intended to say that the soldiers fell. What was meant by that sentence is a clear, and not merely a likelihood matter.

It must be emphasized that, in the course of this application, the CS rules are not being interpreted at any point as rules of inference making truth claims about the physical world. It is for that reason that I am not contradicting myself in this paper by describing the CS approach while arguing against deductive and IP approaches. The clearest way to mark the difference is to see that there is no inconsistency involved in retaining the rule expressed informally as $X\{fall\} \rightarrow X\{hurt\}$ while, at the same time, retaining a description of some situation in which something animate fell but was not hurt in the least. There is a clear difference here from any kind of deductive system which, by definition, could not retain such an inconsistent pair of assertions.

Such rules are intended to cover not only "world knowledge" examples like the last example, but also such cases as:

[9] "In order to construct an object, it usually takes a series of drawings to describe it"

where, to fix the second "it" as "object" and not "series" (though both yield equivalent semantic densities on expansion) we need a CS inference rule in the same format that can be informally expressed as "an instrument of an action is not also an object of it". The point of such rules is that they do not apply at a lexical level like simple facts (and so become an unmanageable totality), but to higher level items like semantic formulas and cases. Moreover, their "fitting" in any particular case is always a "fitting better than" other applicable rules, and so is a further extension of the uniform principle of inference by density.

Discussion

Two points about the general procedures I have described are of some topical theoretical importance. Firstly, the notion of preferring a semantic network with the greatest possible semantic density is a natural way of dealing not only with normal semantic disambiguation, like the "policeman" example [1] above, but with metaphor. For example, if we know from the lexical tree for "drink" that, as an action, it prefers human actors, then, in any given context in which a human actor is available, it will be preferred to any non-human actor, since its presence creates a dependency link and increases the semantic density of that context. So, in

[10] "The crook drank a glass of water"

it would correctly opt for the human and discard the shepherd's staff sense of "crook". Yet in the case of

[11] "My car drinks gasoline"

it would accept the automobile sense, since no animate actor is available to be preferred. This all seems obvious and natural, but is in fact very hard to accommodate within the derivational paradigm of TG and GS, where there must either be a stipulational rule requiring, say, animate actors for drinking (in which case [10] is rejected, although perfectly correct), or there is a rule which permits both [10] and [11] to be "derived", in which case it is hard to see how a structure involving a shepherd's staff is to be excluded (as it properly should be). PS cannot formally be accommodated within the conventional derivational paradigm because it is equivalent to running another derivation with a different set of rules (after dropping a stipulation about actors for "drink" in this case). Yet it makes no real sense in the conventional paradigm to talk of re-running a derivation after an unsatisfactory result.

A second lacuna in the derivational paradigm is the lack of a natural way of dealing with what one might call knowledge of the real world, of the sort required for the analysis of [8] above. Lakoff(1971) seems to think that such cases are to be dealt with, within the derivational paradigm, by calling such assumptions as are required "presuppositions" and using a conventional first order deductive apparatus on them. The question I want to discuss briefly is whether such apparatus can be fitted into the derivational paradigm.

The new development in linguistic theory that GS brought, it will be remembered, can be expressed in Lakoff's(1972) claim that "the rules relating logical form to surface form are exactly the rules of grammar". In order to make my most general point below, let me digress briefly upon the last quotation, and summarise the results of detailed argument established elsewhere (Wilks 1972b). The difficulty in discussing the quoted claim hinges upon what exactly "relate" in that sentence is to be taken to mean.

With GS, as with all such theses, there are two ways of looking at them: one is to take the words as meaning what they appear to mean; the other is to assume that they mean something quite different. The first approach gives us what I shall call the TRANSLATION view or the CONSEQUENCE view depending on how we take the word "relate" in that last quotation. The second approach would give what I could call the RENAMING view. By that I mean that when Lakoff speaks of logical form he doesn't mean that in any standard sense, but as some linguistic structure, either familiar or of his own devising. In either case, on the renaming view, GS would not really be ABOUT logic at all, and disputes about it would be wholly an internal matter for linguistics. When Chomsky (1971) and Katz(1971) write of GS as "notational variant" of Chomsky's work they are taking the renaming view.

The consequence view is the most obvious possibility, namely that the "relates" is by inference, valid or otherwise, and that the well formedness of sentences is settled by whether or not they can be inferred from logical forms. Much of the evidence for the assumption that this is Lakoff's view is circumstantial, but it is reinforced by his introduction of rules of inference with "It is clear that there is more to representing meanings than simply providing logical forms of sentences" (1972, p. 686). That quotation seems to rule out the translation view: that logical forms are the meaning, or "backbone", of sentences and can be related to them by mere rules of translation. The translation view also becomes less plausible when one remembers how much of Lakoff's discussion of these matters is about inference: if GS were really about translation into logical form(which may be equivalent to the "transformations preserve meaning" view, (see Partee 1971) then inference would have no place at all in a discussion of natural logic. So then, the consequence view must be Lakoff's view if he has a firm view. Three clear and simple considerations tell against it:

(1) There is just no clear notion available of inference going from logical forms to sentences. Rules that cross the logical form-sentence boundary are rules of translation.

(2) There is the problem of "reverse direction" : how could one, even in principle, analyse sentences with reverse inference rules to produce logical forms. Reversing inference rules is to produce falsehood, as in "if this is not colored then it is not red" What possible interpretation could we attach to such a procedure in the context of GS? It is true that the relation of a sentence to its presuppositions has the required "inferential direction", but no one has ever seriously suggested that the premisses required for the solution of examples like sentence [8] will in general be presuppositions, in any sense of that over-worked word. In the case of [8], it is clear that the information required for its solution is NOT presuppositional.

(3) Any "consequence interpretation" of GS will find itself committed to the view that logical falsehoods are ill-formed in some sense, and so should not be generated by a proper linguistic system. This will lead to difficulties with apparently well formed sentences that might well be held to express implicit logical falsehoods, such as:

[12] I have just proved arithmetic complete

An immediate result at this point in the argument is that, given the consequence interpretation of GS, a GS system could never be used as an analysis system, and so could surely never function so as to take account of "social context" in the way Lakoff would like. At the very least it requires some more explanation as to how that can be done with a system that is, in principle, non-analytic.

My principal critical point is that the inductive inferences that analysis of examples like [8] requires cannot be incorporated into the derivational paradigm on the consequence interpretation. The use of inductive premisses is not like the use of entailments, where if something is true then something else must be. If inductive premisses or inferential rules are inserted into a derivational system then that system simply must make mistakes sometimes. And so it must mis-analyse or, within the generative task, it must mis-sort sentences. It only makes sense to use such inferences within a system that is capable, in principle, of finding out it has gone wrong and trying again. Such systems have been developed within what may be called the "artificial intelligence" [AI] approach to language processing. They cannot be TG or GS systems, where the derivation simply runs and that is that.

The weaker way out of this dilemma for GS would be to save the derivational paradigm and give up a consequence interpretation. The latter could be achieved either by accepting the "renaming view", which would be most unpalatable I suspect, or by accepting a weaker interpretation of the inference rules like that adopted by PS. That

is, so to interpret the rules of inference as to remain wholly within what Carnap called the 'formal mode'. As I described them, CS rules equivalent to "hurt things tend to fall" are fitted on by preference, but are never interpreted as making truth claims about the future course of the physical world. They are merely used to make claims about what a sentence asserts, not about the course of events, or about what the speaker meant. So, the successful application of the quoted CS rule to [8] allows us to infer that the speaker asserted that the women fell. However, if the speaker followed [8] immediately with

[13] "And after ten minutes hardly a soldier was left standing as the gas drifted toward them across the marshes"

then we may say that the speaker has merely contradicted himself in some weak sense. And, on this hypothetical approach, one might decline to analyse adequately utterances with self-contradictions. And, if point (3) above is correct, then CS must already decline their analysis, since they are "ill-formed", so this extra proviso should cause no problems. The advantage of this form of interpretation of rules would be that it keeps linguistics self-contained and out of the morass of probability and inductive logic.

However, reconsideration of examples like [10] & [11] shows that the elaborate compromise just described is not possible and that one must adopt the stronger approach or, to change the metaphor, see that the door has been open behind me all the time, and simply give up the derivational paradigm in favour of an "intelligent" meta-system. This is essential for examples like the "car drinks gasoline" case above, whose analysis requires some process equivalent to running the derivation again with different rules---and hence some meta system available to administer such a rerunning. And PS, as I believe, an economical way of describing such a system.

There should be nothing very revolutionary in suggesting that the derivational paradigm be quietly abandoned. Its acceptance has for some time been inconsistent with the real everyday practice of generative linguists, which is to do informal analyses of difficult and interesting example sentences (see Schank & Wilks 1973, for a detailed development of this point), and hardly ever to derive or generate a sentence. In a recent paper Fillmore (1972) too seems to have been questioning, from a very different starting point, the most general description of their activities that modern linguists have inherited without question.

If the stronger approach is to give up the derivational paradigm and adopt the AI one (where I am using that term very loosely to cover any formal approach to language processing that admits of wholly extra-derivational procedures), then the question arises as to whether the deductive procedures that Lakoff now (1971) envisages as part of linguistics can and should be retained. There would clearly

no longer be the barriers to the use of deductive processes that existed within the derivational paradigm. Would there be others?

I think there are several very general difficulties about the use of a deductive system for assigning structure to natural language, and some of these have emerged already within the AI paradigm, and are worthy of attention by generative linguists. One difficulty concerns the theoretical problem of specifying firm procedures that would allow any particular deductive solution to be carried through. Here I refer to the enormous problems of search and strategy within domains of theorems. These are very large problems that cannot be discussed here. A smaller but persistent one can be illustrated again with regard to sentences [8] & [12].

In the AI paradigm, unlike the derivational one, a system analysing [8] would have the opportunity to reconsider its solution (that the women fell) on encountering [13] in some social context. What one might call the "standard AI approach" (for example, Winograd 1972) explains its moves at this point roughly as follows: if we analyse [8] with the aid of the inductive generalization, and later information shows us that the inference was false (i.e. we encounter some form of contradiction), we will simply retrace our steps to some earlier success point in the procedure and try again with the new information.

The persistent trouble with this sort of answer (and there is no better one) is that there is no general test of logical consistency available, even in principle, and it is too much to hope that a text would correct our misinferences immediately, by making the interposed sentence between [8] and [13] "But it was not the women who fell".

Paradoxically, it is this sort of deductive approach that Lakoff seems to be embracing in (Lakoff 1971) without seeing that it requires not only the wider AI paradigm but consistency heuristics as well. It may be worth pointing out that, even if this strong deductive AI approach were to have adequate consistency heuristics, it would still be inadequate as a natural language analyser. For example, one of its assumptions is that speakers always use correct logic in their utterances. But consider the following silly children's story:

[14] "I have a nice dog and a slimy snake. My dog has white furry ears. All animals have ears but my snake has no ears, so it is a mammal too. I call it Horace."

Since the story contains a logical error, any deductive analyser for solving anaphora problems must conclude that it is the dog that is called Horace, since that only that conclusion is consistent with what it already knows. Whereas any reader can see that Horace is a snake.

My hope is that PS can at some point be extended, still within the 'formal mode' and not making deductive claims, so as to cover in natural language whatever the human competencies about consistency may turn out to be, and my hunch is that they will require shallow chains of common sense reasonings, drawn from a wide data base, rather than the narrow longer chains of the deductive sciences proper. But even if further research should show this particular approach to be inadequate, the need would still exist for some theory of linguistic inference, one not simply obtained second-hand from logicians, for that will never do. The derivational paradigm has shielded linguists from the pressure to explore this important area, but as the paradigm falls gradually away, the need will become clearer and more acute.

References

- Chomsky, N.: Deep Structure , Surface Structure and Semantic Interpretation. in (eds.Srinberg & Jakobovits) SEMANTICS, Cambridge, 1971.
- Fillmore, C. : "On generativity". in (ed.Peters) GOALS OF LINGUISTIC THEORY, ENGLEWOOD CLIFFS, N.J., 1972
- Joos, M.: Semantic Axiom #1, LANGUAGE, 1971.
- Katz, J. : Generative Semantics is Interpretive Semantics. LINGUISTIC ENQUIRY, 1971.
- Katz, J. & Postal, P.: AN INTEGRATED THEORY OF LINGUISTIC DESCRIPTIONS. Cambridge, Mass., 1964.
- Lakoff, G.: The role of deduction in grammar, in (see 3 below)
- Lakoff, G.: Linguistics and Natural Logic, in (eds. Davidson & Harman) SEMANTICS OF NATURAL LANGUAGE, New York, 1972.
- Lakoff, G.: letter in NEW YORK REVIEW OF BOOKS, 1972.
- Partee, B.: On the requirement that transformations preserve meaning, in (eds. Fillmore & Langendoen) STUDIES IN LINGUISTIC SEMANTICS, New York, 1971.
- Quillian, R. : "Semantic Memory", in (ed. Minsky) SEMANTIC INFORMATION PROCESSING, Cambridge, Mass. 1968.
- Scribner, M.: The concept of comprehension, in (eds. Carroll & Freedle) THE CONCEPT OF COMPREHENSION, Washington, DC, 1972.
- Wilks, Y.: Decidability and Natural Language, MIND, 1971.
- Wilks, Y.: GRAMMAR, MEANING AND THE MACHINE ANALYSIS OF LANGUAGE, London, 1972.
- Wilks, Y. : "Lakoff on Natural Logic", Stanford A.I. Project Memo #161, 1972.
- Wilks, Y.: An Artificial Intelligence Approach to Machine Translation, in (eds. Schank & Colby) COMPUTER MODELS OF THOUGHT AND LANGUAGE, San Francisco, 1973.
- Wilks, Y.: "Understanding without proofs", PROC. THIRD INTERNAT. CONF. ON AI, Stanford, 1973.
- Wilks, Y. (with R. Schank): "The Goals of Linguistic Theory revisited", Stanford A.I. Project Memo, 1973.

Winograd, T: UNDERSTANDING NATURAL LANGUAGE, Edinburgh, 1972.