

AD-769 673

NATURAL LANGUAGE INFERENCE

Yorick Wilks

Stanford University

Prepared for:

Advanced Research Projects Agency
Department of Defense

August 1973

DISTRIBUTED BY:

NTIS

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

DISCLAIMER NOTICE

THIS DOCUMENT IS THE BEST
QUALITY AVAILABLE.

COPY FURNISHED CONTAINED
A SIGNIFICANT NUMBER OF
PAGES WHICH DO NOT
REPRODUCE LEGIBLY.

AD 769673

STANFORD ARTIFICIAL INTELLIGENCE LABORATORY
MEMO NO. AI-211

AUGUST 1973

COMPUTER SCIENCE DEPARTMENT
REPORT NO. CS-383

NATURAL LANGUAGE INFERENCE

by

Yorick Wilks

ABSTRACT: The paper describes the way in which a Preference Semantics system for natural language analysis and generation tackles a difficult class of anaphoric inference problems (finding the correct referent for an English pronoun in context): those requiring rather analytic (conceptual) knowledge of a complex sort, or requiring weak inductive knowledge of the course of events in the real world. The method employed converts all available knowledge to a canonical template form and endeavors to create chains of non-deductive inferences from the unknowns to the possible referents. Its method of selecting among possible chains of inferences is consistent with the overall principle of "semantic preference" used to set up the original meaning representation, of which these anaphoric inference procedures are a manipulation.

This paper owes much to discussions with Annette Herskovits, Bill Sagapson, Ken Colby, Bruce Anderson and Horace Enea. The mistakes are all mine of course.

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.

20

Reproduced by
**NATIONAL TECHNICAL
INFORMATION SERVICE**
U.S. Department of Commerce
Springfield VA 22151

INTRODUCTION

The paper describes inferential manipulations of a representation of the meaning structure of natural language. It differs from previous descriptions (1, 2, 3, 4) of this semantics-based system, which have concentrated on the representation itself, and above all on the procedures by which the representation is produced from input sentences and paragraphs in English. In this paper I assume that structure of representation, except for a brief recapitulation, and concentrate on operations upon it for the solution of a class of difficult anaphora problems.

The system described is part of a running system for understanding and translating natural language on the POP6/10 at the Stanford A. I. Laboratory, programmed in MLISP and LISP1.6. I shall not in anyway stress the translation-into-French aspect of the work, but its presence provides a continual empirical check of the adequacy of the inference and "understanding" described here.

The earlier emphasis on the construction of the linguistic base is, I believe, fully justified. The present system is, to my knowledge, the most comprehensive producer of meaning structures for sentences of natural language available at present in terms of implementation: vocabulary, disambiguation of many-sensed words and references, etc. Moreover, as I have argued elsewhere (2), it is not the implementation of a conventional theory from linguistics, but is one with somewhat different principles of content.

In what I call its basic mode, the system already resolves anaphoras depending on superficial conceptual content of the text words. This is done in the course of setting up the initial representation. I shall call these type A anaphoras. For example, in "Give the bananas to the monkey, although they are not ripe. They are very hungry", the system in its basic mode would decide that the first "they" refers to the bananas and the second to the monkeys. It can do that simply from what it knows about monkeys getting hungry because they are animals, and bananas having phases like ripeness because they are plants. All this information is, one might say colloquially, part of the superficial meaning of "banana" and "monkey".

This paper describes an "extended inference mode" of the system that tackles the other kinds of anaphora example that I shall call types B and C. Consider the correct attachment of "it" in "John drank the mineral from the glass, and it felt warm in his stomach". It is clear that the pronoun should be tied to "whisky" rather than "glass", but how it is to be done is not immediately obvious. Analysis of the example (see below) suggests that the solution requires, among other things, an inference equivalent to the sentence "whatever is in a part of X is in the X".

Anaphoras like the last I shall call type B, in that the inferences required to resolve them are analytic but not superficial. By analytic I simply mean that the sentence above, about parts and wholes, is logically true and not in any clear sense a fact about the real world (but rather about the meanings of words). What is meant by "superficial" in the distinction between types A and B will become clear after a discussion of the meaning formalism employed here.

I shall also discuss type C anaphoras, which require inferences that are not analytic at all, but weak generalisations (often falsified in experience) about the course of events in the world. Yet their employment here is not in any sense a probabilistic one. In "The dogs chased the cat, and I heard one of them squeal with pain", we shall, in order to resolve the referent of "one" (which I take to be "cat" and not "dog"), need a weak generalisation equivalent to "animate beings are more often than inanimate beings may be unpleasantly affected". Such expressions are indeed suspiciously vague, and a reader who is puzzled at this point should ask himself how he would explain (say, to someone who did not know English well) the way he knew the referent of "one" in that sentence. It can hardly be in virtue of a probabilistic fact about cats and dogs because the same generalisation would apply to whatever was chasing and being chased. I admit we are puzzled, if he does not come up with something very like the inference suggested, and it may be the nature of natural language itself that is troubling him.

The inferences for type C, then, are general expressions of partial information (in McCarthy's phrase) and are considered to apply only if they are adequately confirmed by the context. What I mean by that will become clear in the course of what follows, but in no case do these expressions yield deductive consequences about the future course of the world. Indeed, they would be foolish if they did because the world's course cannot be captured in that way. In the unhappy example above, it might have been his earlier dinner that made him feel good.

ARTICLE REGARD OF THE SYSTEM'S BASIC MODE OF ANALYSIS

In its basic mode, the system fragments texts (into phrase/clause like items) and attaches a template to each. A template is a canonical form of connectivity of semantic formulas as follows (where a formula is a complex item, to be described, corresponding to a sense of an input word):

```

      F1.....F2.....F3
      /  \  \  /  \  \
F11 F12 F13 F21 F31 F32

```

F1, F2, F3 are the principal formulas of the template and are always agent, action and object (in that order), though any of them may be a dummy in any particular example. (F11, F12, F13) is a list of formulas dependent on main formula F1 etc. It should be said, in view of other current uses of "template", that it is not a surface form at all, but a format underlying meaning representation. Moreover, it does not function within a crude pattern matching technique, such that if some text fragment has no templates matching it, it is thrown away, as it were. Special routines are called in such situations to construct an appropriate template item. All this basic material is set out in earlier papers.

The structure of formulas is explained below in some detail. In brief a formula is a nested list (a binary tree in fact) of semantic primitives called elements (expressed here as LISP atoms). Each such formula expresses a sense of at least one English word.

Let me give an example of a template structure at this point by using the following simplifying notation: any English words in square brackets F1 stand for the meaning representation of those words in the Performance Semantics system. This device is important in the presentation of the material in that the content of the coded forms can be seen immediately, whereas the complex coded forms themselves would be awkward to read as such, a sentence read a word at a time. But it is important to state that the rules and formalisms expressed within (1) are really structured primitives, and that their tasks could not be carried out, as some still seem to believe, by massaging the language words themselves to stand for their own meaning representation.

So then, the template for "The black horse passed the winning post easily" could be written (ignoring any ambiguity problems for the moment):

```
[horse]----[passed]-----[post]
      ↑           ↑           ↑
    [the black] [easily] [the winning]
```

If one or all of the agent, action or object formulas are missing, the template code(s) is filled if with a dummy element DTHIS. Thus a template for "The old house collapsed" could be written:

```
[house]----[collapsed]----DTHIS
      ↑
    [the old]
```

In the case of structures like prepositional phrases, we consider the preposition represented as a pseudo action, and the whole template as having a dummy Agent. Thus for "at the Derby", we have:

OTHIS[...](a)[...](Derby)
 ↑
 [the]

The representation of a text (composed of fragments) is then a network of these template networks. The templates are interconnected by case ties. The notion of case is discussed in more detail below, but for the moment a case can be thought of as tying one template to some particular node in another template by a link of a certain type: namely the case type, which specifies the sort of dependence the former template has on the latter. In the sentence "He lost his wallet / in the subway" (fragmented at the stroke) we might say that the second fragment of the sentence depends on "lost" in the first, and that the dependence is the locative case. Thus in the representation, the template for the second fragment would be tied to the central, action, node of the first, by a link labelled LOCA. The node on the first template to which the case tie ties is called the marker of the second template.

Type A anaphoras are dealt with adequately within this framework, and therefore the purpose of this paper is not to describe the basic model or operation of the system) because we get a denser network of links by considering the formula for the appropriate referent substituted for the problem pronoun than with its rivals. A link is correct, or preferred, or strengthened, in the network, when a preference is satisfied. So, if we think of the formula for "ripe" as expressing a preference for application to plants, we see why a denser network arises in the above example for correct solution, rather than for one equivalent to "ripe monkeys". The way in which formulae express preferences of various sorts is described below.

Next, solved, these type A anaphoras also constitute links between templates, from the pronoun variable to its correct referent. Thus the completed first form of the whole representation obtained from the parser looks as:

CASE MARK ANAPHORA F1 F2 F3 (F1 dependents) (F2 dependents) (F3 dependents)

The structural form of representation of a text paragraph is called its IREP (Interlingual Representation), or "semantic block". No emphasis has been placed here on syntax analysis of the input, and a reader who consults (1, 2, 3 or 4) will see that all of conventional grammar analysis has been done in the course of setting up this form of representation. An example of such an IREP, for the monkey-banana example, is given below as computer output.

Having sketched in the basic mode and its representation, we can now get to our muttons and sketch in the extended inference mode that is the heart of this paper.

QUICK SKETCH OF THE "EXTENDED INFERENCE MODE"

The extended inference procedure is called whenever the basic mode cannot resolve an anaphora between two or more candidates by semantic (or density). In the example about John and his stomach, density techniques have nothing to decide whether the glass or the whisky is in the stomach. On a basis of preferred agents and objects of actions (what I have referred to as superficial conceptual information) both are equally good candidates. The extended inference procedure is called, and if it succeeds it returns a solution to the basic mode which then continues with its analysis. If it too should fail to reduce the number of candidates to one, then the top level of the system tries to solve the problem by default, or what a linguist would call focus. Roughly, that means: assume that whatever was being talked about is still being talked about. So, in "He put the bottle on the shelf and when he came back next week it was gone", neither density criteria, nor the extended inferences to be described here, will help at all. So the system may as well assume, in this limited context, that the bottle is still the focus of attention, and thus the reference of "it".

Consider again the following sentence after all the basic mode's routines have been applied:

11: [John drank the whisky / 2 DIRE : DTHIS from a-glass / 3 :
from a stomach / 4: DTHIS in his stomach]

For each of the square brackets, this item is a template representation. The case names DIRE (direction) and IN (containment) indicate the dependencies of templates 2 on 1, and 4 on 3, respectively. The DTHISs are dummies added to fill out the canonical formulae for the in cases of missing agents, objects etc. Further assume that the "it" has been tied to "John" by the basic mode, and present as problem of analysis. And assume too that the basic mode has failed in that of "candidates" for the reference of "it" ("whisky" and "bottle"), because if there had not been such a list of more than one candidates the routine under description would not have been called in the first place.

Extractions are then made from each template in turn, and if and only if a template contains a representation of either an answer word or a pronoun pronoun itself. An extraction is the unpacking of the corresponding case tin: both those in the action (second) formula of the template and those labelling as link to other templates. For the example we obtain the following extractions: which are in the form of $\langle \text{template} \rangle \langle \text{link} \rangle \langle \text{word} \rangle$ (where the first digit refers to the template, the second to the extraction, and "+" links words with a template boundary):

11: [John drank (IN in) John +part]

12: [John drank (DIRE to) John +part]

21: (whisky (DIRE FROM) a+glass)

22: (1 21 IN in) his+stomach)

So, in this internal representation we have acquired new template-like objects that express, in canonical form, new analytic information extracted from the existing templates, and from which new inferences can be made. It is postulated that the generation of this recipient information from the deeper levels of the formulas is essential to the process of understanding. These new forms differ from standard templates only in that their second node, or pseudo-action, has had a case name CONSD onto whatever the node was before. Note here that the form (IN in) is not redundant since the IN becomes the name precisely as containment, while the English preposition can indicate many cases other than containment, as in "in five minutes".

We now describe how these particular extractions were made, even in the absence of any detailed knowledge of the structure of formulas:

11 is inferred from the template for "John drank the whisky" because from the structure of the formula for "drink" it follows that the liquid drunk is subsequently inside the drinker. This is because, when making up the formula for the action "drink", we express in it what the action consists in causing a liquid to be inside the agent of the action.

Analog, 12, is inferred because the same formula specifies that the liquid enters the drinker through a specific part of the drinker (his mouth, if you wish).

13 is inferred from the direction case of the second template, whence we know that it was the whisky that was moved from the glass.

Lastly, 21 is inferred from the direction case of the fourth template, because whatever the referent of it is, it is also in John's stomach.

Let us see where we are: we have obtained new template items that grade relative information, but did not appear in the original text. (As we shall see in the detailed treatment below, some of the above are obtained from extraction, more strictly defined, and some from what I shall call "repacking the semantic block".)

In the pool for inference procedures we now have the original templates that mention either the variable pronoun or the possible "answer" referents, plus the extractions. We also have access to an inventory of Common Sense Inference Rules (CSIRs) which are of the form (T1 T2), where T1 and T2 are T-forms, that is, templates or extractions.

We now try the strategies in turn: first we try a zero-point strategy, which is to try to identify an answer template (or extraction) and a variable template (or extraction) without the use of any rules.

The general assumption here, and it is a strong psychological assumption, is that in order to resolve these painful ambiguities the understanding system is going to use the shortest possible chain of inferences, if any. And a zero-point strategy will, as it were, have no inferences at all (in terms of a chain of CSI inferences) and so if it works, it will always provide the shortest chain.

This strategy is adequate for the example under discussion, because in the first step (which is the initiation of matching) identify extractions 1 and 21, and identify 21 and the whisky, and we are home. This is the resolution of a B type anaphora, requiring only analytic, not inductive, conceptual information. It should be noticed that this type anaphora (defined earlier by the need of weak inductive information for their solution) can also be solved by the zero point strategy, because some extractions, and in particular those from the good wine (free output), are inductive and non-analytic.

If the zero-point strategy fails, we bring down all the CSI rules that contain in action subformula occurring in an answer or problem formula in the proof, and attempt to find the shortest chain that leads from some answer to some variable.

Thus, in the sentence "The soldiers/ fired+at the women/and I saw several fall", we extract a form equivalent to [soldiers strike women], since we can tell from the formula for "fired+at" that the action is intended to strike the object of the action. We are seeking for partial confirmation of the assertion [?several fall DTHIS], and such a chain is completed by the rule [X strike Y] $\rightarrow$ [Y fall DTHIS], which is not form equivalent to, say, [X strike Y] $\rightarrow$ [Y die DTHIS], 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 natural language 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 fire: [soldiers fire DTHIS] $\wedge$ [X fire DTHIS] $\rightarrow$ [Y fire+at X] $\rightarrow$ [X strike Y] $\rightarrow$ [X fall DTHIS] etc., based on the assumption that though that fire guns get fired at ("...he who lives by the sword shall perish and...."). But such a chain would be longer than the one already constructed and would not be preferred.

MORE ON THE BASIC MODE

Formulas

Formulas are structures corresponding to senses of words, expressing their meanings. Much of the body of this paper is concerned with the manipulation of such structures, and the extraction of information from them, so it is important to have some general idea of their construction and interpretation.

Formulas are binary trees, expressed as lists, of semantic elements, denoted by right and left brackets. The elements are either case-envelopes or functions such as CAUSE, STRIK, CHANGE, or items such as HUMAN, MAN, BEAST, etc. (standing, as examples here, element names that are mostly Anglo-Saxon monosyllables, but there are about 1000 in all, and some need informal explanation, such as GRAIN, used to denote grains of wheat). There are also elements like KIND indicating categories, and elements indicated by and (initial *) that stand for "not" or "is not" (e.g. *ANIMATE) to cover MAN, BEAST and other human properties. In addition, most elements have a negated form, marked with a tilde (~) on the element name. I assume here that the use of logical primitives of this general sort, that are not logical themselves, need no special defense at this point.

The head of an element in a formula is its rightmost, called its head. This is the most general sort of item, or action, or type of action, expressed corresponds to: for example, any word corresponding to a human being will have MAN as its head.

Formulas are binary trees of unlimited depth, they can be decomposed into pairs of elements and subformulas, down to the most primitive elements. This process is called building up the formula, or decomposing it into its parts. At each stage there is a dependence of the left part on the right: this dependence is

1. the left part is a function or action

2. the left part is an item, or act of

3. the left part is a function or action or item.

In interpretation, the interpretation is unambiguous, once we know the frame and functions of the elements in play. So, the word HUMAN always means "a human" (envelops a human being), WRAP is always an action when in the right hand position (and always a qualifier when in the left dependent position). It is never an item and MAN is a agent not an object in the frame. It should be an object if in the representation of "a human being" is enveloped" because agents of actions may be unmarked, and objects are never unmarked. Conversely (WRAP THING) is a

container, since WRAP is always a qualifier when in the left hand, dependent, position in a pair.

An important notion is that of the semantic preferences that formulas can express. Consider the formula for "grasp" of objects:

```
"grasp"(action) = ((*AN) SUBJ)((*PHYSOB OBJE)((THIS (MAN  
WRAP)) INST(ODD SENSE))))
```

The case of agents SUBJ and OBJE occur at the top levels on the left of the formula, and at that level in an action formula they express the preferred agent and object of the action concerned.

Thus, grasping, in this sense is an action preferably done by animate beings (*AN) to physical objects (*PHYSOB), and consists in a act of reaching, holding, and done with an instrument (INST is the case agent) which is a part of the body. When I say "prefers" here, I mean that, if the preferred agent or object cannot be found, a template is constructed with whatever is available. Thus, "The robot grasped the block" would never be rejected; it would only be less preferred than any possible competing interpretation that had an animate agent. I have argued in (2) that this approach to rules of formation has unexpected consequences for linguistic theory.

We can now also expect another formula to be available for "grasp", one such as:

```
"grasp"(action) = ((*HUM SUBJ)((STGN OBJE)(TRUE THINK)))
```

In this case, we have an action, preferably done by human beings to grasp (understand) ideas, principles etc), namely of thinking them to be true, or adequate, signs.

The procedure of the basic mode always fit this last formula into a template structure for "He grasped the principles", and the other formula for "grasp" into the template for "the boy grasped the toy", by means of the preference and semantic density techniques described in earlier papers. These preferences for agent and object are part of the "superficial conceptual information" referred to earlier.

A few other rules will help to clarify the notion of "knowing our way round a formula" when interpreting it:

Agents (case implicit need not be specified by SUBJ case) must always occur at the top level in an action formula as described above, or they attach to the head of a formula, as in:

```
"grasp"(action) = (NOTPLEASE FEEL)(SUBJ MAN)
```

Here, the normal priority of agents being to the left of (= dependent of) the corresponding action, is violated, since MAN is the agent for FEEL, while at the same time being the head of the whole formula.

This violation of order in search is indicated by also violating the order restriction that normally makes the SUBJ case element the dominant (initial hand member) of the pair in which it occurs. The corresponding rule of analysis is: On encountering SUBJ as dependent, expect action for the agent to follow to the left".

There is, however, one never implicit. Moreover, an object is considered a part of all actions to its right. This enables us to express the important notion of real and apparent agents of actions. For example in:

"fire+at" (xAN) =

((XAN) (XAN) ((XAN) (OBJE) ((STRIK GOAL) ((THING MOVE) CAUSE)))

This action (fire, preferably by human beings to animate beings) is one of firing and firing to move (the bullet) with the aim (GOAL case) of striking something. Since xAN is the object of all actions to its right, it is the object not only of CAUSE, but also of STRIK. Hence, the striking is also of the same animate being. Moreover, THING (the bullet) is internally the agent of MOVE, not the object of CAUSE, which is correct as far as the meaning of "fire+at" is concerned.

Notes:

At present we operate with a distinction system of ten cases, which are listed below, together with (in capital letters) the semantic concept that represents them, the questions that define them, and examples of subformulas expressing them. Defining a case is a tricky matter, but the distinction method is reasonably adequate. Note that the subformulas examples are of those parts of a formula that would express that notion AS PART OF THE MEANING OF A WORD. The subformulas are not, of course, how the system would express the quoted words if encountered in a text, when they would be represented by a template.

recipients: FOR "for a woman" → ((FEM MAN)FOR)
what/why to? what/who for?

instruments: INST "with a stick" → ((LINE THING)INST)
what with? by what means?

directions: xDIR (see below), TO, FROM, UP
"from the top" → ((UP POINT)FROM)
where to where from? at what? out of where? by what?

possessives: POSS "owned by a man" → ((MAL MAN)POSS)
who owns the thing mentioned?

location: LOCA "at that time" → ((THIS(WHEN POINT))LOCA)
when? where? where at? by what? in what time? near what? at what
time? where? where? before when?

in what? "in a glass" → (((((FLOW STUFF)OBJE)WRAP)THING)IN)

out of what? "out of wood" → ((PLANT STUFF)SOUR)

do it? "to strike a woman" →
(((FEM MAN)OBJE)STRIK)GOAL)
to what end? for what purpose?

with what? "with a glass" → (((((FLOW STUFF)OBJE)WRAP)THING)NOTW/TH)
accompanied by what/who? with what/whom? without what/whom?

subject? "Still who did this?"

object? "What was this done to?"

Certain cases above have negative forms leading to additional
elements: NOTIN, NOTPOSS, NOTIN, NOTWITH.

These elements have two functions, and occur in two sorts of
constructions: formulas and IREPs. In formulas they express part
of the meaning of a word sense. Thus in

"The glass is full of liquid" → ((FLOW STUFF)SOUR)THING)

the word "full" has a liquid source (FLOW STUFF), and is in a
container (THING). The other function of these elements is, as
already explained, the name of the tie between the template for some
fragment and the part of another template.

Thus, the word "to" is the class of direction case elements (TO and
FROM), and it occurs only as the indicator of the case of a fragment,
never in formulas. Conversely POSS occurs only in formulas, never as
the indicator of a fragment case.

The word "to" is only included in a formula when it is specific:
the aspect of the case is involved. In the formula
for "pour", for example, we include a direction specification for
movement (TO). However, in the formula for "move" we do
not include the element TO or FROM, even though movement must in fact
be in some direction, since we have no reasonable expectation about it
in the sentence "John moved". Sentences containing "move" may very well go
on to specify the direction involved, but its association with "move"
is completely arbitrary, and we cannot expect any confirmation of
expectations that would, say, resolve ambiguities. In this respect

the system differs from other systems that do create case expectations for wide classes of actions, which are essentially unhelpful, as in this example, and so we would claim unhelpful semantically.

The IREP, or semantic block representation

Block 100000 is an example of an IREP for a pair of English sentences. The format of the block is the list structure described earlier, as the result from the basic mode of operation. The only difference from that format is the presence in it of the stereotypes from which French is subsequently generated (see Wilks and Herskovits 1977, p. 107). The French, as generated from the block, is written under the block itself for diagnostic purposes. The generation of context-sensitive stereotypes are drawn into the block during analysis, along with the formulas. The process of generation is then a recursive unwrapping of the block.

THE BANANAS TO THE MONKEYS ALTHOUGH THEY ARE NOT RIPE! THEY ARE VERY HUNGRY.

DONNER DES BANANES AUX SINGES BIEN QU' ELLES NE SOIENT PAS MURES !
ALORS ILS SONT FAIMS.

```
((EX (THE BANANAS) ((EX NIL NIL ((IMPL)))) 4 (((THIS DTHIS)
HUNGRY)) ((GIVE ((ENT OBJE) GIVE)) GIVE (DONNER)) ((MUCH ((*ANT
OBJE) ((GIVE ((ENT OBJE) WANT)) (OBJE PLANT))) BANANAS (FEMI BANANE)) NIL
THIS DTHIS) ((TO THE MONKEYS ((PTO (GIVE) RECI ((PREOB 4)))) 6
((THIS DTHIS) HUNGRY)) ((THIS PTO) PTO NIL) ((THE (MUCH ((MAN LIKE)
BEAN) ((MONKEYS (MASC SINGE)) NIL NIL NIL)) ((ALTHOUGH THEY ARE
NOT RIPE) ((ALTHOUGH (GIVE) COND (BIEN QUE (SUBC~ L)))) 1 (((MUCH
((GIVE ((ENT OBJE) WANT)) (OBJE PLANT)) (THEY BANANAS)
((GIVE ((ENT OBJE) WANT)) (OBJE PLANT)) ((PRES ~ (BE BE)) ARE ((IS_OBJECT HUNGRY)
AVOIR ((DIROB Q FAIM)) ((IS_OBJECT THIRSTY) AVOIR ((DIROB Q SOIF))
((IS_OBJECT AFRAID) AVOIR ((DIROB Q PEUR)) (ETRE)) ((PLANT POSS)
((CAN USE)) KIND)) RIPE (MUR)) NIL NIL NIL)) ((THEY ARE~ VERY
HUNGRY) 2) ((NIL NIL NIL ((INDCL)))) 1 (((THE (MUCH ((MAN LIKE)
BEAN) ((THEY MONKEYS) ((PRON G MASC PLUR~ )) ((PRES (BE BE)) ARE
((IS_OBJECT HUNGRY) AVOIR ((DIROB Q FAIM)) ((IS_OBJECT THIRSTY) AVOIR
((DIROB Q SOIF)) ((IS_OBJECT AFRAID) AVOIR ((DIROB Q PEUR)) (ETRE))
((PLANT POSS) ((TASTE SENSE) WANT) STATE) KIND)) HUNGRY (AFFAME~ ))
NIL ((MUCH NOW) VERY (TRES))) NIL))
```

THE IMPLEMENTATION OF THE EXTENDED MODEL

There are three parts to the extended inference mode: the REPACK routine that takes the IREP block and repacks it; the EXTRACT routine which produces extractions, new knowledge not explicit in the text analysis; and INFER which tries to link an answer T-form to one which has a problem variable, that is, a text pronoun giving trouble.

The REPACK routine.

It attempts to replace dummy nodes in the IREP wherever possible before handing the whole representation to the extraction procedure. The replacement is itself a complex form of inference, sometimes as complex as the inference routines on which we are concentrating here. However, there is no pretence that these procedures are mutually independent and can be carried out in a single process in this way. The degree of dummy-replacing done by REPACK, in the construction of a new block IREP' from IREP, varies with particular action cases.

If we look back at the informal extractions done from the "John drank the whisky, ..." example, we will see that the new T-forms 21 and 41 are actually obtained by filling in a dummy agent in some template from a node in another template. Thus from [2 DIRE : DTHIS from H1DIRE] we obtained the new T-form, numbered 21, [whisky (DIRE from) alcohol], this was done by filling the dummy agent node of the template for "with a glass" with the formula for "whisky", and shifting the direction case marker into the pseudo-action. This is a repacking, not an extraction proper, since the T-form obtained simply replaces a template "assertion" already in the representation. As we shall see, a true extraction is a new T-form altogether.

Let us now distinguish replaceable and unreplaceable cases roughly as follows. The dummy agent in the second (instrumental) template for "He let his father / with a club" cannot be replaced to yield any form equivalent to [father (INST with) a club]. So we may say that the instrumental case is unreplaceable. But the dummy agent in the second (recipient) template of "He bought the flowers / for his mother" can be replaced to yield a form for [flowers (RECIP for) his mother], and the recipient case is replaceable, and is replaced by the operation of REPACK.

At the top level REPACK can be written in LISP as:

```
REPACK(IREP) -> IREP'
```

The general routine

IREPR takes each T-form, or template-like item, in IREPR in turn and either leaves it modified if necessary, in a new block IREPE, or leaves it as T-forms extracted from it. At present, extractions are only made from templates that contain either one of the possible answers, or one of the variables of the problem. The former are templates containing a formula for a word on the list ANS, the latter are templates one of whose nodes is (QUERYMARK THIS). Any templates not containing either an answer or a variable are simply transferred unchanged from IREPR to IREPE.

Thus, the general form of the extraction routine at the top level is

EXTRACT (IREPR ANS) → IREPE

Taking each template in turn, we first consider those processes that modify it, and then those which produce new T-forms from it. In the first instance, some manipulations to do with negation, and with the split and join cases.

If an agent or object formula is negated, the negative item in its formula is removed and the head of the corresponding action formula is negated, because all the subsequently applied inference rules automatically negate the action. Thus, in $\neg$ notation, we would achieve by this procedure the coded equivalent of

[John drank neg(gin)] → [John not+drink gin]

Each agent and object formula is then scrutinized by the question "does it satisfy the preference expressed by the corresponding action. If it does not, does any "of-phrase" qualifier of it do so in stead". If so, replace the agent or object by that "of-phrase" qualifier as the true agent or object.

Thus:

[John drank at+class(of+wine)] → [John drank wine]

[At+class(of+women drank wine)] → [Women drank wine]

The sub-prime of EXTRACT takes the action formula of a template and moves leftwards through it seeking case heads (other than SUBJ and OBJ). If it finds one, it asks is it replaceable, and, if it is, EXTRACT looks at subsequent fragments to see if REPACK has already replaced it. If it has been replaced it is forgotten, thus avoiding the same case information being extracted twice. It detects that REPACK has made such a replacement by finding the case name itself in the preceding action of a succeeding template, and a replaced dummy as the corresponding agent.

With the next case, for example, the dependent of the case element becomes the action of the new T-form. In this case, as with every other, an attempt is made, on finding a potential agent or object for the new T-form at the top level of the action formula of the template under inspection, to identify it with the main agent or object formula of the template. If this can be done, the agent or object formula of the original template is used, as being more specific. For example, in extracting from the action formula in the template for "John hit the deer", we find the goal case in [fired-at], with dependent STRIK, which is the attempted action. The object of that action, found to its left, is *AN, which can be instantiated by the formula for "deer" in the main template, namely (THIS BEAST). So the latter is used as the object of the new, extracted, T form [John strikes deer], since "deer" is more specific than "animate being".

For most other cases (agent, direction, location, containment and support), the case element provides the new pseudo-action, and the main subject of it is specified as the dependent of the case element. The pseudo-action is defined as follows: it is the highest level object actually available of the action that dominates the case (to its immediate right in the formula).

So in the formula for "pour" in "I pour the wine"

```
((FLOW STUFF) ((WRAP THING) TO) MOVE) CAUSE)))
```

we encounter (moving leftwards through the the formula) the direction case in the sub-formula ((WRAP THING) TO), implying that the (FLOW STUFF), which is the highest level object in the formula, is moved in the direction of a container, or (WRAP THING). The case element TO is dominated by MOVE, whose highest level object is (FLOW STUFF), which would become the pseudo-agent of the new extracted T-form, but since it can be identified with the object of the old template, namely "wine", it is, and that becomes the pseudo-agent of the new T-form, since it is more specific than "liquid". It is, (WRAP THING), the container, becomes the object of the new T-form and the direction case element becomes the action so we get an extracted formula

```
((CONTAIN (CONTAINER) (FLOW STUFF)) (TO PDO) (WRAP THING) )
```

where (CONTAINER) is (some+container)

The Hierarchy

This routine has access to the representation IREPE produced by the rules AN and CSIR, the inventory of common sense inference rules. Its form at the top level is

```
INTER(IREPE ANS CSIR) - ANS'
```

where ΔL is either ΔN or some sublist of it, preferably containing only a single node, the solution.

When ΔL is the zero-point strategy: trying to match some of the nodes in some variable T-form directly, with no use of CSIRs. This kind here means that the two T-forms as arguments of a function ΔL produce a non-nil result, which will be a list of the corresponding, but non-identical, nodes in the two matched T-forms. For the solution of the example "John drank the whisky. . . ." is obtained by the zero-point strategy, and rests on the matching of the two T-forms:

[whisky (IN in) John+part]

[stomach (IN in) stomach]

where MATCH returns the list ((2) whisky (John+part stomach)) comparing the answer. No such match can be made for the alternative solution "glass".

The main principle of inference at work is to select the shortest possible chain of inferences, on the assumption that an ambiguity of understanding of this sort should be solved in the most shallow way possible; that the situation becomes intolerable for the understanding. Thus a zero-point solution, if available, will always be the shortest possible chain of inference.

If the zero-point strategy fails, the CSIRs are called, stored as a list, indexed by their action subformulae, and, moreover accessible from both "antecedent" and the "consequent" action subformula. At present we experiment with inferences of length one: those which require only a single CSIR for their solution. However, it should be possible to extend the present strategy to at least length two; and apparently they will almost never be any longer.

Let us look once more at the example "The soldiers fired at the women and several fell". We have to resolve "several", which cannot be done by the basic mode since both soldiers and women can equally be killed. Let us set out the fragment representations and the explanations obtained as follows:

[1 (THINGS (THAT WOMEN))

[2 (SOLDIERS (THAT THING))

[3 (THINGS (THAT WOMEN))

[4 (THINGS (THAT THING))

[5 (THINGS (THAT THING))

[6 (THINGS (THAT THING)) (THINGS (THAT THING))

The inventory of rules is searched for those containing any action simultaneously occurring in a T-form in the pool that also contains either an "answer" or a "problem variable". In this case we pull in a rule informally expressed:

[I STRIKE (*ANI 1)] ~ [(*ANI 2) fail DTHIS]

Here variables are indicated by numbers ; *ANI expresses a restriction on the variable that any value of it must be animate, and the double ~ indicates that this rule can be considered as running in either direction.

This CSIR form is of course a more perspicuous form of:

[I (THIS STRIKE) (*ANI 1)] ~ [(*ANI 2) (NOTUP BE) DTHIS]

which, of course covers a wider class of activities than simply the English verb "strike". It would cover at least "hit" "batter" etc.

That a chain of length one is established by the rule from T-forms 12 to 14, since the "animate condition" is satisfied and the variable *x* is identified by the rule with the formula for "women". It should be noted here that the inference rules are very weak in that the application of a rule like the present one is perfectly consistent with the description of a situation where an animate being *do* strike in some way but does not fail. And this weakness is wholly intentional.

One important question in CSIRs is whether or not negation is significant or not. The negation of the action in a T-form is not at all significant. Consider "John drank no gin / in his martini / but it left warm / in his stomach nonetheless". In the template for the first fragment, shifting the negation to the action, and extracting for the containment case from the formula for "drink", we shall obtain a T-form

[gin (MATCH BE) container]

and another:

[gin (MATCH BE) John]

Conversely we shall obtain, by the same method, from the second template, [martini (IN BE) container] and [martini (IN BE) John]. In trying to tie only one of these drinks by matching to the extraction [gin (IN BE) John], we shall, without the use of CSIRs be able to cut the ANS list down to a single member, namely [martini], since MATCH will show us that [gin] cannot stay on ANS.

However, if we applied the same analyses to a sentence like "John found a car / in the window / and he knew / that he would get it", the basic rule of inference of "it" is "car" and not "window" and he would not be guilty of linking the first and fourth T-forms with some rule like $[(\text{XANI } 1) \text{ want } 2] \rightarrow [(\text{*ANI } 1) \text{ have } 2]$.

But, and this is the point, if the same sentence had concluded "... and he knew he would not get it" we should have required the same rule and this would be wrong. This rule has its "consequent" action which is the condition of its application and it is irrelevant to its application.

Also, the extended mode can be seen to be non-deductive very clearly at this point. It could be said to be of the form $A \rightarrow (B \vee \neg B)$, and so it is compatible with any content whatever in a deductive system, and so it is not a clear and necessary.

GENERAL DISCUSSION AND SOME PARTICULAR COMPARISONS

The system above cannot be considered in any way adequately finished, because no one has any very clear idea of what constitutes "inference" in this area. But even to qualify, the basic mode and the extended mode fall under a considerable vocabulary and range of words and words, and the extended mode must be shown to be able to handle a reasonably sized inventory of CSIRs.

The main question employed in this extended mode will also be those employed in a general discourse ambiguity procedure to back up the extended mode's procedure to resolve ambiguity within a small number of sentence fragments. Ambiguity over a context larger than these fragments and discourse, just as is ambiguity of the sort discussed in the chapter, but he should be prepared for it in an adequate understanding.

The extended mode is a strong, and possibly naive, psychological assumption, namely that chain length is a reasonable metric to compare different inferential interpretations. I think it is reasonable, but that the tension introduced into understanding by extended mode ambiguity has been overlooked. Notice here that the length of a chain of CSIRs employed, not counting extension above, different ways of writing down formulas will not affect chain length.

However, I would justify the principle as being essentially an extension of what is called semantic preference (Wilks 2) used in the extended mode's representation. That preference was justified as selecting the "semantically densest" interpretation which was, I think, the "least meaning" (in the sense in which a statement carries the maximum possible information). Similarly, the shortest chain of inferences also minimises the information in play, and introduces the least extraneous inductive

information into the system. It is clear that such a notion of information based choice is ultimately inadequate. We only have to consider a sentence like "I was named after my father" where it seems clear that we exclude one interpretation simply because it contains virtually no information. This alone shows there must be some qualification to a "minimising information" theory. However, the fact that all available theories are wrong, by no means puts them all in the same position. I think such hypotheses about the overall manner in which an understanding system endeavors to maintain its coherence are well worth making and testing, and that they represent an aspect of human language "competence" almost wholly ignored by current linguistic and artificial intelligence. One could make the point more precisely as follows: virtually all the systems in those areas define "success", that is to say the success of a particular parsed representation with respect to a text. What they do not tell us is what number of success are registered, as is almost always the case in realistic practice. But human understanders do not just accept the first, or opt for the first they find, or pick one at random: they prefer one in particular on some principled basis.

It is for this reason that the subject investigated in this paper cannot be treated in isolation from an adequate linguistic base system, as some seem to think. The inferring of a correct interpretation is intimately related to the systematic exclusion of competing interpretations, and any system that does not allow for ambiguity of sense and structure in at the start can hardly appreciate this point because the difficulty never arises there, but then neither does one essential aspect of natural language either. I have developed elsewhere (6) an abstract view of meaning along these lines: that to have meaning is essentially to have one meaning RATHER THAN ANOTHER. Or, put another way, having meaning essentially involves procedures for the exclusion of alternative interpretations. That, I believe, is the residual truth lurking beneath the "proceduralizing of meaning", a thesis which when taken at face value is patently false.

Let me mention a closely related shortcoming of the micro-world approach to natural language analysis: it concerns what I believe to be an unfortunate mistake in AI about the notion of "inference". Let me begin by stating the obvious, not from dry motives of clarity, but because I believe the muddle has important practical consequences in the area of natural language understanding.

There are two references, in the bare sense of that word, of "transitive" that people might make from one assertion to another.

1) All Englishmen are untrustworthy and Cecil is an Englishman, SO he is untrustworthy.

2) Cecil is an Englishman, SO he is untrustworthy.

3) This is triangular, SO it is three-sided.

I take it that this is a deduction, true in all possible worlds, and hence independent of the meanings of the words "Englishman", "Cecil", and "trustworthy".

And this is an inference, simply and solely, and certainly not valid, whether or not it happens to be true for some English Cecil.

And this is a valid inference, true in all possible worlds, as they say, because of the central meanings of "triangular" and "three sided", a fact that is sometimes expressed by saying that the premise missing, for this to be a deduction, namely "all triangular things are three sided" is in reality false.

What is the point for our purposes of all this dogmatic and self-trustful identification? Simply this: the extra-conceptual content of the so-called dictionary, that is CSIR inferences of the sort we have discussed in this paper problems in text, are of type (iii). These inferences could function as part of a deductive system by the addition of sufficient inductively unreliable premisses to convert them to form (i). Then could then function within established deductive machinery, such as first order PC, PLANNER in one of its modes of operation, etc.

But there may be no need for doing that, at least in the case of natural language analysis, because the conclusions reached can be no more reliable than the dubious generalisations functioning as premisses, whatever the power of the deductive machinery intervening.

In this paper I have described how such weak information can fulfil a problem-solving role in natural language analysis, in terms of a notion of "adequately confirmed" inference in context. But that does not require the deductive machinery at all.

The point will be clarified here by noting two research situations where, by contrast, the deductive machinery may pay its way: (1) in robots, and (2) in simulated micro worlds.

In the case of a robot, really moving about in the world with perceptually mediated information and plans, the world itself can provide a clear sense of contradiction. If the robot's deductions fall at the door is open, but it bangs into the firmly closed door in fact, then the conclusions are contradicted and the preceding premisses can be discarded, as would be the case with a scientific theory which fails some critical experiment. That is to say, the premisses may be unreliable, not because there can be contradiction of conclusions (the deductive machinery can transfer the "not" back to some premise, as in solving a scientific theory, though the question of which premise it should be transferred back to is very difficult of course).

This problem of attention is quite different from the analysis of natural language where there is little or no expectation of contradiction. In understanding the text, the understander does not know that there is little or no chance of encountering contradiction in the text in the near future. A robot could in principle learn from experiments, and in the case of a dialog in natural language, even learn to step back and ask questions of the speaker, but this method of dealing with contradiction without experimenting on them. The probability of the robot's being able to discover such contradiction is very small and even if it does track back having done so, to the next least probable contradiction. And at the moment no system is in striking evidence of such an ability.

The problem of contradiction in micro-worlds is different. Here there is no payoff to be obtained, but there is no need for it since all promises made are kept; nothing is, and no real information can ever enter the system. The robot, after executing the command "Clear off the top of the stack", it is clear, by definition, apart from the possibility of a contradiction. No lingering and sticky contradiction can be allowed to imperil the stability of the house of cards. It will be clear that such situations have no place in the unreliable inductive information required for the analysis of natural language.

It is clear that there has been that if there is no payoff to be obtained from a deductive approach to natural language understanding, there is no reason for pursuing it. This position is different from the position taken in formal fields, and urges the pursuit of the deductive approach of the former: that is to say, the search for a deductive approach to work in a field rather than a search for an approach to work in a field about whose content we are ignorant. An important point in the deductive position is that their method does also deal with the question of content, or human competence in this area, and not just with the formalization of principles that could be expressed in a deductive way.

After these clarifications, some very brief comparisons follow between the work described here, and three other AI approaches to language understanding: those of Charniak(7), Schank(8) and Harnad(9). In this comparison and criticism of systems is not meant to be a criticism, and I give only brief general remarks, in order to indicate the differences along a number of dimensions: which are (a) the nature of the linguistic base constructed or proposed, in terms of application to everyday texts in English; (b) the degree of implementation, and the definiteness of the task proposed as an application of the elusive notion of "understanding"; (c) the position of the system within the inference-deduction opposition, and (d) the implementation of a preference system that both prefers certain inferences to others on a reasoned principle.

Winograd does not consider the linguistic base essential, and is not particularly interested in the ambiguities of sense of words. He has the most adequate linguistic base of the three, and the present system in general presuppositions. Winograd's system is restricted to unambiguous simple words, and a limited number of their meanings. Even if the words had only one meaning, it is doubtful if the meaning of complex concepts, or relations, could be expressed in that way.

Winograd's system is not intensively implemented but has a very adequate base for it, the resolution of the sorts of anaphora mentioned above in this paper. I think the strategies we advocate are not really different over whether or not the rules form an explicit system, or the facility to chain CSIRs, or what he calls "inference rules". His system is on the verge of implementation through the use of a number of large programs. However, at present it is not really directed to a specific task in the inference field, but to the drawing of inferences per se (as distinct from the drawing of inferences for the solution of some problem or performance of some task). The "inferences to be drawn from x" is not really a task, outside the pages of detective fiction, but a completely implemented within its original domain, a task of inference and assessable task, its merit, as far as I know, from the syntax and the semantics (Minsky and Winograd). In an odd remark, in an Ad context, in that it is not really a task, but a task in semantic analysis by computer (Winograd's work for example) has found unnecessary in view of the same authors' distinction between perception-cognition and inference as a task and a problem.

Winograd's work and deduction I do not feel on sure ground because I am not sure of the authors' work, since, naturally, their work is not designed to answer this question of mine. However, the distinction is complicated by the fact that some of the authors' work is devoted to cover processes that are almost certainly not the same as the feeling is that Winograd's system is not the same as the authors' and Schank's, like the present one, are not the same. The latter call in inference rules whose validity is determined by the possibility of fitting them into the context in hand. Any clarification of the relation of their work to this distinction will be gratefully received.

Winograd's work and preference and choice between inference rules. I think nothing has been done by the authors, and I think there will always be one and only one inference rule in terms of their rules, or that the first inference rule, and that, is, I think, the only worker in the field

also has given attention to this question.

REFERENCES

- Charniak, E. An Artificial Intelligence approach to machine translation. Stanford AI Laboratory Memo #161, and in Schank, R. (Ed.) Computer Models of Thought and Language. San Francisco, 1976 (in press).
- Charniak, E. Reference semantics. Stanford AI Laboratory Memo #162, and in Schank, R. (Ed.) Formal Semantics of Natural Language. San Francisco, 1976 (in press).
- Charniak, E. Understanding without proofs, in Proc. Third. Internat. Joint Conf. on AI, 1973.
- Charniak, E. and Herskovits, A. An intelligent analyzer and parser for natural language. in Proc. Internat. Conf. on Computational Linguistics, Pisa, 1973.
- Herskovits, A. The generation of French from a semantic representation. Stanford AI Laboratory Memo #212.
- Levy, Y. Understandability and natural language. Mind, 1971.
- Levy, Y., Jack and Janet in search of a theory of knowledge.
- Levy, Y. and Rieger, C. Inference and the computer model of natural language. Stanford AI Laboratory Memo #197.
- Levy, Y. Understanding Natural Language. Edinburgh, 1972.
- Levy, Y. and Fagan, S. Artificial Intelligence Progress Report. Stanford AI Laboratory.
- Levy, Y. Semantic memory. in Minsky, M. (Ed.) Semantic Memory. MIT, 1968.