

MANFRED KRIFKA

PARAMETRIZED SUM INDIVIDUALS FOR PLURAL ANAPHORA*

1. MODELS OF DYNAMIC INTERPRETATION

One of the most important conceptual changes in the study of the semantics of natural languages was a reorientation from a static notion of meaning, in which meaning was seen as the truth-conditional CONTENT of expressions, to a dynamic one that takes meanings to be changes in the INFORMATION STATE of the participants in a conversation. The first proposals along these lines, which are due to Stalnaker (1973, 1974, 1978) and Karttunen (1974), concern the treatment of FACTUAL INFORMATION. Information states are modelled as sets of possible worlds, and sentences as mappings from such information states to information states. This opened the door for a new and compelling way of analyzing the distinction between assertional content and presuppositional content of a sentence (Heim 1983b). Later, Heim (1982, 1983a) and others put forward and developed the idea of modelling ANAPHORIC REFERENCE in sentences and texts within a dynamic framework. Information states are modelled by SETS OF VARIABLE ASSIGNMENTS, and sentences are analyzed again as mappings from such information states to information states. In this framework, it is possible to combine the "factual" perspective and the "anaphoric" perspective, as done in Heim (1982). We can also assign meanings to sub-sentential expressions; for example, Rooth (1987) and Groenendijk and Stokhof (1990) give interpretation rules that work with dynamic meanings all the way down.

* A first version of this paper was presented at the conference Semantics and Linguistic Theory III at the University of California at Irvine, March 1993, and a preliminary version was written at the Institute for Logic and Linguistics, IBM Germany, Heidelberg. Increasingly more elaborate talks were given at Rutgers University, February 1995, at the University of Texas at Austin, April 1995, at the Center for the Study of Language and Information, Stanford University, September 1995, and at Brown University, Providence, September 1995. I owe an enormous debt to the audiences at these occasions, in particular to Hans Kamp, Maria Bittner, Godehard Link, Jan-Tore Lønning, Stanley Peters, and to Greg Carlson and two anonymous reviewers of L&P. The final version of this paper was completed during my residency at the Center for Advanced Study in the Behavioral Sciences, Stanford University. I am grateful for the financial support provided by the National Science Foundation, #SES-9022192, and by the University Research Institute, University of Texas at Austin.

DISCOURSE REPRESENTATION THEORY (DRT), as presented by Kamp (1981) and elaborated on in many other contributions, most prominently Kamp and Reyle (1993), can be seen as a part of this general paradigm change from static interpretation to dynamic interpretation. However, DRT introduces an additional level of semantic representation, namely DISCOURSE REPRESENTATION STRUCTURES (DRSs). Sentences are interpreted as functions from DRSs to DRSs, via so-called DRS construction rules. The DRSs, in turn, are interpreted with respect to a model that represents factual information. In this setup DRSs are an essential level of representation. In particular, the DRS construction rules make reference to specific structural features of this representation and are not compositional in the classical sense. It has been shown that classical DRT (the fragment of Kamp 1981) can be reworked into a compositional theory (cf. Zeevat 1989, Asher 1993, Muskens 1994, among others). But the essential use of non-compositional DRS construction rules in newer versions of the theory, such as Kamp and Reyle (1993), seems to preclude a compositional formulation.

The issue of compositionality has played an important role in recent discussion; cf. for example the comparison between DRT and Dynamic Predicate Logic in Groenendijk and Stokhof (1991). Now, compositionality may mean different things to different people. In a certain sense, even the various versions of DRT are compositional on the level of the representations, as the representation of a complex expression can be derived in a well-defined way by the representations of their immediate syntactic parts and the way they are combined. But the operations that are performed by the DRS construction rules are fairly arbitrary and not restricted by general principles, as compared to the few well-defined semantic combination rules like functional application in non-representational accounts. Furthermore, by referring to properties of the semantic REPRESENTATION of expressions, one ascribes a certain reality to the features of the semantic representation. This would call for some independent justification of this level, of which there is little in sight.

Non-representational, fully compositional analyses are clearly to be preferred in the absence of additional evidence for a particular level of representation. However, non-presentational theories have failed to reach the empirical coverage of DRT-type analyses, which have been far more successful in discovering and describing intricate anaphoric phenomena. For example, Partee (1984) has developed a theory of temporal anaphora, Sells (1985) and Roberts (1987) have described modal subordination, Kadmon (1990) has dealt with asymmetric quantification, Asher (1993)

has developed a theory of abstract entity anaphora, and Kamp and Reyle (1993) have treated the highly complex interaction of plural reference with collective and distributive predication. Convincing treatments of these phenomena within non-representational theories are not yet available (but see the remarks on van den Berg 1990 and Elworthy 1995 in the conclusion).

In this paper I will examine the phenomena of plural anaphora that have been discussed and analyzed in Kamp and Reyle (1993) and try to account for them within a non-representational, fully compositional theory. I will show that many of the phenomena that are discussed by Kamp and Reyle can indeed be expressed in a non-representational, compositional way. But this requires substantial changes in our understanding of the central notion of variable assignment.

2. One Example and its Treatment in DRT

Consider the following text:

- (1) Three students wrote an article They sent it to L&P.

One interpretation says that three students wrote an article together, and these three students sent that article to L&P. This COLLECTIVE interpretation is easy to model as soon as we allow for sum individuals in our model structure along the lines of Link (1983). The sentences in text (1), however, also have a DISTRIBUTIVE reading; the first sentence can be read as saying that three students each wrote an article, and the second sentence can be understood as saying that each of these students sent his or her article to L&P. Modelling distributive readings is not a problem as soon as we assume distributive operators that may be covert, as in (1), or overtly expressed by *each*, as in (2).

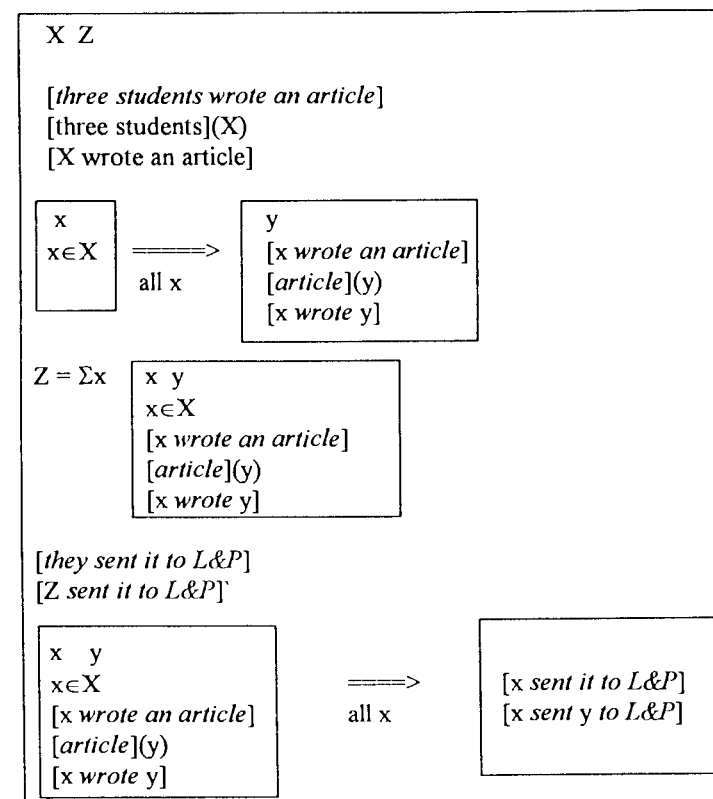
- (2) Three students each wrote an article. They each sent it to L&P.

The problem with (1) on the distributive reading (let me call this (1d)) and with (2) is the pronoun *it*: We are talking about three articles, but it seems that we can use a singular pronoun to refer to these articles. Notice, however, that *it* does not simply refer to the sum individual consisting of the three articles. The second sentence of (1d) and (2) does not say that the students collaborated in sending their three articles; rather, it says that each student sent his or her article separately.

Kamp and Reyle would treat (1d) and (2) in the following way (cf. Sections 4.1.5 and 4.2.6; actually, their book derives various options, and

what I present here is a fairly characteristic analysis). Assume that we start with an empty DRS. The first sentence introduces a plural discourse referent X and the conditions $[three\ students](X)$ and $[X\ wrote\ an\ article]$ (I follow Kamp and Reyle in abbreviating complex syntactic structures by English expressions in brackets). The presence of a distributive operator triggers universal quantification over the elements of X . Quantificational structures are represented by so-called DUPLEX CONDITIONS, which consist of a RESTRICTOR and a MATRIX that are combined by a quantifier. Here, the restrictor introduces a singular discourse entity x that ranges over the elements of X , the quantifier is a universal quantifier, and the content of the verb phrase is spelled out in the matrix – here, a singular discourse entity y is introduced, along with the conditions $article(y)$ and $[x\ wrote\ y]$. If we now would go ahead and assume that *they* in the second sentence picks up X , we would fail, as the pronoun *it* could not be interpreted because y is not accessible from outside its local box. Kamp and Reyle (1993) propose that the duplex condition triggers the introduction of a new plural discourse entity Z that is identified with the sum (Σ) of all the x that satisfy the union of the restrictor and the matrix. Notice that Z and X are anchored to the same sum individual here; they differ only in the way how they are introduced in the DRS, that is, on the representational level. The subject *they* in the second sentence now picks up Z and introduces the condition $[Z\ sent\ it\ to\ L\&P]$. At this point we are entitled to replace Z by a copy of the box that has generated Z . Since the second sentence contains a distributive operator as well, we are entitled to write down another duplex condition, this time with the box associated with Z as restrictor. This makes the discourse referent y available to the matrix; in particular, y can serve as antecedent for the pronoun *it*. We end up with the following DRS:

(3)



Now, the introduction of a second discourse referent Z for the three students looks quite *ad hoc*. For example, why is it that we introduce a second discourse referent Z for the three students? It seems that this is done just so that it may be picked up by a pronoun in a distributive sentence later on. We do not need it for the evaluation of the first sentence itself, nor for other cases of anaphoric pronouns that could just pick up X , e.g. for a sentence like *They were students of linguistics*.

Another problem is: At what point should Z be introduced? There are two options: Either it is introduced upfront, when the representation of the first sentence is construed, since it might be needed later. One can see this forward-chaining strategy as a "generalizing to the worst case", certainly not very attractive for either explaining the linguistic abilities of humans or their implementation on machines. The second option is to introduce Z at the time it is needed; this means in our example, at the point when the pronoun *it* needs to find an antecedent. But then we should observe some evidence of reprocessing previous expressions, that is, a

garden path effect. But this seems to be absent: A text like (2) sounds perfectly smooth. Now, one might object that DRT never was intended to be a realistic model of actual text processing, and this is certainly true. However, it is very tempting to develop formal models of anaphoric relations that one can in principle relate to processing issues, and so any theory that provides the same coverage but a more straightforward analysis of texts like (2) should be preferred.

A more principled objection to the illustrated treatment of plural anaphora concerns the power of the rules involved. The DRS construction rules are virtually unconstrained re-writing rules. For example, the 'Rule for distribution over a set obtained by Abstraction' (Kamp and Reyle p. 389) allows us to rewrite all the conditions in a box created by lambda abstraction over a discourse entity in the restrictor box of a universal quantifier. There is nothing that would restrict the type of copying operations that are possible. For example, it would be possible, in principle, to write quite absurd rules that, say, copy only the first and the last condition of a preceding box into the current restrictor box.

As for the first objection, the *ad hoc* introduction of a discourse referent *Z*, Kamp and Reyle could point out that there seems to be independent motivation for such a move, as we need similar rules to treat cases of so-called MODAL subordination:

- (4) If John sees a new issue of L&P in the library, he checks it out. Usually he xeroxes it the same day.

In (4), the conditional sentence introduces a duplex condition. The second sentence, which contains an adverbial quantifier, introduces a duplex condition as well. The restrictor box of this latter duplex condition is provided by abstracting over the restrictor and matrix of the first duplex condition; the sentence says that whenever John sees a new issue of L&P and buys it, he reads it the same day. This seems quite parallel to Kamp and Reyle's treatment of example (3), in which we also had to construct a complex condition out of the restrictor and the matrix of a preceding duplex condition.

There may indeed be a systematic connection between plural anaphora and modal subordination. But the problem is that the rules that deal with modal subordination that are offered by Kamp and Reyle (1993) seem as unconstrained as the ones for plural anaphora, and hence objectionable for theoretical reasons. Basically, they state that whenever we have constructed a duplex condition $[R](Q)[M]$, where $[R]$ is the restrictor and $[M]$ is the matrix, then either $[R]$ or $[R, M]$, the combination of $[R]$ and $[M]$, can serve as restrictor of subsequent quantifiers. It is quite unclear why

this is so – why, for example, can we not pick up just $[M]$, or why is it not possible to combine just the $[R]$ -boxes of two subsequent duplex conditions?

It is quite obvious that modal subordination involves a special kind of anaphoric dependency. The restrictor of a quantifier that is not spelled out has to pick up the restrictor, or the restrictor and the matrix, of a preceding quantifier. Now, DRT is primarily a theory of anaphoric dependencies, their accessibility restriction, and the way they influence semantic interpretation. The theoretical device introduced for anaphoric dependencies is the notion of discourse referent or discourse marker, as first proposed by Karttunen (1976). DRT can be seen as a highly restrictive theory of how discourse entities are introduced, accessed, and discarded. However, the anaphoric phenomenon of box copying is treated in a quite different and strikingly informal way. This stands in sharp contrast to the narrowly defined and well-motivated constraints for the accessibility of discourse entities that represent standard anaphora.

In the following I will show that a non-representational analysis of examples involving plural quantification like the one presented by Kamp and Reyle (1993) is possible, and I will then address the question whether it leads to a more restrictive overall framework.

3. A New Proposal Using Parametrized Individuals

The framework that I would like to contemplate here is inspired by theories of dynamic interpretation that use quantification over variable assignments to express anaphoric relations, such as Heim (1982, Chapter 3), Heim (1983a), Rooth (1987) and Groenendijk and Stokhof (1991). But I will have to propose a more complex, recursive notion of variable assignment, which in a way mimicks the recursive DRSs in Kamp and Reyle (1993). In particular, I will elaborate on an idea mentioned by Rooth (1987) as a way of treating partitive quantifiers, such as the following:

- (5) Most of the students who wrote an article₁ talked about it₁.

Rooth observes that when we take *the students who wrote an article* to refer to a standard sum individual (i.e. the sum of x such that x is a student and there is an article y such that x wrote y), and *most of* to be a quantifier over the atomic parts of this sum individual, then there will be no way to interpret the pronoun *it*, which refers to the article of each of the students who wrote an article. Rooth suggests using PARAMETRIZED INDIVIDUALS, a notion that was introduced by Barwise (1985, 1987). A

parametrized individual is an individual that comes with a variable assignment. For (5) we need parametrized individuals x such that x is a student that wrote an article and x is associated with a variable assignment that maps the index 1 to the article that x wrote. Modelling variable assignments by sets of pairs of indices and individuals, and modelling the association of individuals and variable assignments by pair formation, the parametrized individuals for (5) will be of the form $\langle x, \{\langle 1, y \rangle\} \rangle$, where x is a student and y an article that x wrote. We can form the sum of all such parametrized individuals. Representing sum individuals for simplicity by sets, the reference object for *the students who wrote an article*₁ will be

$$\{\langle x, \{\langle 1, y \rangle\} \rangle \mid x \text{ is a student, } y \text{ is an article that } x \text{ wrote}\}.$$

The partitive quantifier *most of* will express a quantification over the atomic parts of this individual, that is, the elements of this set, and the VP *talked about it*₁ will be predicated of these atomic parts. As every such part is of the form $\langle x, \{\langle 1, y \rangle\} \rangle$, for each student x we can access the article y that x wrote: It will be the value of the assignment $\{\langle 1, y \rangle\}$ applied to the index 1.

Assume, for example, that there are three students s , s' and s'' that wrote the articles a , a' and a'' , respectively, and that no other student wrote any article. Then the phrase *the students that wrote an article*₁ refers to the semantic object

$$\{\langle s, \{\langle 1, a \rangle\} \rangle, \langle s', \{\langle 1, a' \rangle\} \rangle, \langle s'', \{\langle 1, a'' \rangle\} \rangle\}$$

which I will write for the sake of clarity as

$$\{s[1 \rightarrow a], s'[1 \rightarrow a'], s''[1 \rightarrow a'']\}.$$

Notice that this representation allows us to identify, for each student $x[f]$, the article that x wrote: it is $f(1)$, the value of f with respect to the index 1.

Let us now define the data structure of parametrized sum individuals. This requires the following auxiliary definitions:

- (6)a. Let $\underline{D}$ be the set of DISCOURSE ENTITIES. This is a countably infinite set, and I will use the set of natural numbers here. I will refer to discourse entities with variables d , d' etc.
- b. Let $\underline{U}$ be the set of URELEMENTS in the model of interpretations. In illustrative examples I will use the letters a , s , a' etc. for elements of $\underline{U}$.
- c. Let $\underline{S}$ be the set of (simple) SUM INDIVIDUALS. For simplicity

of exposition I will model sum individuals by non-empty subsets of $\underline{U}$ that is, $\underline{S} = \text{pow}(\underline{U}) - \{\emptyset\}$. For example, $\{a, a'\}$ is a sum individual. $\underline{S}$ also contains singleton sets like $\{a\}$; I will omit the braces and simply write a .

- d. Let $\underline{P}$ be the set of PARAMETRIZED SUM INDIVIDUALS, or simply P-INDIVIDUALS, for which I will use variables like x , y , x' , etc. P-individuals are defined as sets of pairs of sum individuals (elements of $\underline{S}$) and assignments (elements of $\underline{G}$).
- e. Let $\underline{G}$ be the set of ASSIGNMENTS, for which I will use variables g , h , k , f , i , j , etc. Assignments are functions from discourse entities (elements of $\underline{D}$) to p -individuals (elements of $\underline{P}$).

More specifically, the sets $\underline{P}$ and $\underline{G}$ are constructed recursively in the following way, starting from the basic case of $\underline{P}_0$ and $\underline{G}_0$. I will use the notation $[A \rightarrow B]$ for the set of PARTIAL functions from A to B , that is, the set of all functions from subsets of A to B . Furthermore, let $\text{pow}(A)$ be the powerset of A , and $A \times B$ be the Cartesian product of A and B , as usual. I will introduce various abbreviatory conventions to enhance readability of the resulting structures along the way. Some of these abbreviations are ambiguous, but this should not cause any real problems.

(7) Recursive definition of $\underline{P}$ and $\underline{G}$:

(i) Basic case:

- $\underline{P}_0 := \text{pow}(\underline{S} \times \{\emptyset\}) - \{\emptyset\}$
- $\underline{G}_0 := [\underline{D} \rightarrow \underline{P}_0]$, the set of partial functions from $\underline{D}$ to $\underline{P}_0$.

Examples of $\underline{P}_0$: $\{\langle a, \emptyset \rangle\}$, abbreviated: $a \{\langle a', \emptyset \rangle, \langle a'', \emptyset \rangle\}$, elements of $\underline{P}_0$: abbr. $\{a', a''\} \{\langle \langle a', a'' \rangle, \emptyset \rangle\}$, abbr. $\{a', a''\}$

Examples of $\underline{G}_0$: $\{\langle 1, a \rangle, \langle 2, \{a', a''\} \rangle\}$, abbr. $[1 \rightarrow a, 2 \{a', a''\}]$

(ii) First induction step:

- $\underline{P}_1 := \text{pow}(\underline{S} \times \underline{G}_0) - \{\emptyset\}$
- $\underline{G}_1 := [\underline{D} \rightarrow \underline{P}_1]$

Example of $\underline{P}_1$: $\{\langle s, [1 \rightarrow a, 2 \rightarrow a'] \rangle, \langle s', [1 \rightarrow a', 2 \rightarrow a''] \rangle\}$, element of $\underline{P}_1$: abbr. $\{s[1 \rightarrow a, 2 \rightarrow a'], s'[1 \rightarrow a', 2 \rightarrow a'']\}$

Example of $\langle 3, \{s[1 \rightarrow a, 2 \rightarrow a'], s'[1 \rightarrow a', 2 \rightarrow a'']\} \rangle$,
 element of $\underline{G}_1$: abbr. $[3 \rightarrow \{s[1 \rightarrow a, 2 \rightarrow a'], s'[1 \rightarrow a', 2 \rightarrow a'']\}]$

(iii) General recursive definition for $\underline{G}_n$ and $\underline{P}_n$, for $n > 0$:

- $\underline{P}_n := \text{pow}(\underline{S} \times \cup\{\underline{G}_i \mid 0 \leq i \leq n\}) - \{\emptyset\}$
- $\underline{G}_n := [\underline{D} \rightarrow \underline{P}_n]$

(iv) Definition of $\underline{P}$ and $\underline{G}$:

- $\underline{P} := \cup\{\underline{P}_i \mid 0 \leq i\}$
- $\underline{G} := \cup\{\underline{G}_i \mid 0 \leq i\}$

Individuals with empty assignment functions ($\underline{P}_0$) should stand for simple individuals without any dependent objects. The definition of p -individuals is fully recursive, but we will hardly ever need more than two embeddings, for dealing with natural language examples.

To get some feeling for the way how parametrized sum individuals are put to use, let me give a few examples. I assume that student s wrote article a and sent it to journal j , student s' wrote article a' and sent it to journal j' , student s'' wrote article a'' , professor p talked to students s, s' , and professor p' talked to students s', s'' .

- (8)a. *the students that wrote an article*₂:
 $\{x[2 \rightarrow y] \mid x \text{ is a student, } y \text{ is an article that } x \text{ wrote}\}$
 $= \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}$
- b. *the students that wrote an article*₂ *and sent it*₂ *to a journal*₃:
 $\{x[2 \rightarrow y, 3 \rightarrow z] \mid x \text{ is a student, } y \text{ is an article that } x \text{ wrote,}$
 $z \text{ is a journal to which } x \text{ sent article } y\}$
 $= \{s[2 \rightarrow a, 3 \rightarrow j], s'[2 \rightarrow a', 3 \rightarrow j']\}$
- c. *[the students that wrote an article]*₁:
 $[1 \rightarrow \{x[2 \rightarrow y] \mid x \text{ is a student and } y \text{ is an article that } x \text{ wrote}\}]$
 $= [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}]$
- d. *the professors that (each) talked to [students that (each) wrote an article]*₂₁:
 $\{x[1 \rightarrow y] \mid x \text{ is a professor,}$
 $y \text{ is a set of elements } y'[2 \rightarrow z],$
 $\text{where } y' \text{ is a student such that } x \text{ talked to } y'$
 $\text{and } z \text{ is an article that } y' \text{ wrote}\}$
 $= \{p[1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}],$
 $p'[1 \rightarrow \{s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}]\}$

For an assignment like (8c) I will say that the index 2 is SUBORDINATED to index 1. In the examples so far we did not encounter cases in which assignments are paired with sum individuals. To illustrate this case, assume

that students s and s' wrote article a , and that students s'' and s''' wrote article a' .

- (9) *the groups of students that (each) wrote an article*₂
 $\{x[2 \rightarrow y] \mid x \text{ is a group of students and } y \text{ is an article that } x \text{ wrote}\}$
 $= \{\{s, s'\}[2 \rightarrow a], \{s'', s'''\}[2 \rightarrow a']\}$

These examples illustrate the intended interpretation. In order to show they are built up compositionally from natural language expressions, it is helpful to work with the following additional definitions. Let me illustrate these definitions with the variable assignment:

- (10) $f = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}, 3 \rightarrow s'']$

The definitions we need to get started are the following ones:

- (11) $\in_R$, the RECURSIVE ELEMENT RELATION, is defined as follows (recall that $d \rightarrow x$ stands for the pair $\langle d, x \rangle$ as element of an assignment):
 $d \rightarrow x \in_R g$ iff
 – $d \rightarrow x \in g$,
 – or there are d', x', x'', g' such that
 – $d' \rightarrow x' \in g$,
 – $x''[g'] \in x'$,
 – and $d \rightarrow x \in_R g'$

That is, $d \rightarrow x$ occurs arbitrarily deeply embedded in g .

Example: $2 \rightarrow a \in_R f$, and of course $3 \rightarrow s'' \in_R f$.

- (12) rdom , the RECURSIVE DOMAIN of g , is defined as:
 $\{d \mid \exists x[d \rightarrow x \in_R g]\}$.
 Example: $\text{rdom}(f) = \{1, 2, 3\}$

- (13) g_d , the VALUE of assignment g with respect to d , is defined as $g(d)$.
 Example: $f_1 = \{s[2 \rightarrow a], s'[2 \rightarrow a']\}$; $f_3 = s''$; f_2 : undefined; f_4 : undefined.

- (14) g^d , the CUMULATIVE VALUE of assignment g with respect to d , is defined as:
 $g^d = \cup\{x \mid d \rightarrow x \in_R g\}$, if $d \in \text{rdom}(g)$, else undefined.
 Example: $f^2 = \{a, a'\}$; $f^3 = \{s''\}$, abbr. s'' ; f^4 : undefined.

- (15) $g \cup h$, g INCREMENTED WITH h , is defined as $g \cup h$, provided that $\text{rdom}(g) \cap \text{rdom}(h) = \emptyset$, else undefined.

Examples: $f + [4 \rightarrow a''] =$
 $[1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}, 3 \rightarrow s'', 4 \rightarrow a'']$
 $f + [1 \rightarrow a'']$: undefined
 $f + [2 \rightarrow a'']$: undefined

- (16) $g <_d h$, g EXTENDED to h by d , holds iff there is an x , $x \in \underline{P}_0$, such that $h = g + [d \rightarrow x]$.
 Example: $f <_d [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}, 3 \rightarrow s'', 4 \rightarrow a'']$

Occasionally I will use the notation for extension for more than one discourse entity; e.g. I will write $g <_{d,d'} h$ iff there is a k with $g <_d k$ and $k <_{d'} h$. We will need a few additional definitions, which will be given when need arises.

One important issue at this point is the way how parametrized sum individuals interface with lexical information. Take a simple singular predicate like **student**, which basically applies to an element x of $\underline{S}$ iff x is a singleton $\{u\}$, and u is a student. But then any parametrized individual $x[g]$, where g is an assignment, should fall under **student** as well. That is, we assume a rule for singular lexical predicates α saying that $\alpha(\{u\}[g])$ iff $\alpha(\{u\})$, where the latter information is provided by the model. In general, if we want to find out whether a lexical predicate applies to a p -individual, we first have to “strip off” the assignment function from the p -individual.

In nominal predicates with number words, the number word specifies the cardinality of the elements of its arguments. For example, we have that a p -individual $x[g]$ falls under the predicate *two students* iff $\text{card}(x) = 2$, and $\forall u \in x$: **student**($\{u\}$). A compositional treatment of an English nominal predicate like *two students* is possible when we interpret the plural form *students* as due to grammatical agreement, that is, without semantic significance. There are two pieces of evidence for this treatment: First, there are languages that do not show number agreement in this case, e.g. Turkish. Second, English requires the plural with the number word *one point zero*, as with decimal fractions in general: cf. *one point zero miles*, not **one point zero mile*. As *one point zero* and *one* presumably have the same meaning (perhaps up to a granularity parameter), namely, the number one, the selection of singular vs. plural forms can only be due to syntactic agreement, not to semantic selection restrictions. Hence a singular noun like *student* and a noun like *students* that got its plural by agreement should both involve the same singular predicate, **student**. A plausible analysis for constructions like *two students* then is **2(student)**, with **2** = $\lambda P \lambda x [x[g]] [\#(x) = 2 \wedge \forall u \in x P(\{u\})]$. Of course, the number form of bare plurals is semantically relevant; we may assume, for example, that the bare plural *students* is translated as $\lambda x [x[g]] [\#(x) \geq 1 \wedge \forall u \in x \text{stu-}$

dent($\{u\}$)]. Note that this representation predicts that *students* applies to single students as well. This is well motivated, as, for example, a question like *Do you have students?* cannot be answered by *No, just one*, but can be answered by *yes*, even if the addressee has only one student. The preference of the singular form *a student* over *student* then can be derived by scalar implicature: If it is known that the referent is a single entity, the more informative form *a student* is preferred.

I should mention that I am using sets to model sum individuals purely for expository reasons. Everything in this paper can be expressed within a model that uses a join operation instead (cf. Link 1984 for arguments against using sets for the modelling of sum individuals). Perhaps the relevant operation should rather be non-commutative list formation if we want to treat cases like *John and Mary are twenty and thirty years old respectively*, which would also be compatible with the main points made in this article.

4. COLLECTIVE PREDICATION

Let me start to develop and illustrate the underlying framework of dynamic interpretation that I will be using in this article. I follow the common assumption that NPs come with syntactic indices, that indefinite NPs introduce new indices, and that anaphoric NPs carry the index of their antecedents. Indices are translated into discourse entities in the course of interpretation. (There is a theoretically attractive alternative, namely, that indices come into play only during interpretation itself, with indefinite NPs taking the next available index not used so far; this alternative, however, leads to notational complications that I will try to avoid.) I will use a representational format close to Rooth (1987), a convenient combination of the languages of set theory and predicate logic with quantification over variable assignments. The meaning of a sentence is a relation from input assignments to output assignments, for which I will use the notational format $\{\langle g, h \rangle \mid \dots\}$, where g is the input assignment, h is the output assignment, and “ $\dots$ ” is some description relating g and h . For expressions that do not change the input assignment I will typically use the same letter for input assignments and output assignments and write, for example, $\{\langle g, g \rangle \mid \dots\}$. A one-place predicate has an additional object parameter x ; it will be specified in the format $\{\langle g, x, h \rangle \mid \dots\}$. Two-place predicates then will be specified in the format $\{\langle g, x, y, h \rangle \mid \dots\}$. I will use P and R , also with primes, as variables for 1-place predicates and 2-place predicates, respectively, and p as a variable for sentence meanings. NPs are treated as quantifiers that take a n -place predicate and reduce it to a

($n - 1$)-place predicate. I will use the lambda notation to express NP meanings – an object NP, for example, will be given in the form $\lambda R\{\langle g, x, h \rangle \mid \dots R \dots\}$, which reduces a two-place predicate ρ to a one-place predicate $\{\langle g, x, h \rangle \mid \dots \rho \dots\}$. The scopal order of quantifiers in the fragment I will be developing always corresponds to the syntactic relation of c-command of surface structure (that is, subjects have scope over objects); this is of course a simplifying assumption. The basic semantic composition rule is functional application, except for the combination of sentences to a text, for which it is relational composition. For the sake of brevity and readability I will not specify a complete syntactic fragment with semantic interpretation, but I will discuss various examples and their interpretation in all relevant aspects, and it should be straightforward to construct the underlying fragment from that.

Let me illustrate the framework with the simple, COLLECTIVE interpretation of a sentence like (1) before we deal with more complicated cases. Take the following sentence on its collective reading:

- (17) Two students₁ wrote an article₂. They₁ sent it₁ to L&P.

The meaning of *wrote* is given in (17a) (I will disregard tense throughout this article). As *wrote* does not change the anaphoric potential, input assignments and output assignments are identical. The meaning of *an article*₂ is given in (17b); notice that it changes the input assignment g to an output assignment h that contains a new discourse entity, 2, in its domain. The condition that indefinites bear a novel index is enforced; otherwise the condition $g <_2 k$ cannot hold. (17c) then gives the result of the application of the meaning of *an article*₂ to the meaning of *wrote*:

- (17)a. *wrote*: $\{\langle g, x, y, g \rangle \mid \mathbf{wrote}(x, y)\}$
 b. *an article*₂:
 $\lambda R\{\langle g, x, h \rangle \mid \exists k[g <_2 k \wedge \mathbf{article}(k_2) \wedge \langle k, x, k_2, h \rangle \in R]\}$
 c. *wrote an article*₂:
 $\lambda R\{\langle g, x, h \rangle \mid \exists k[g <_2 k \wedge \mathbf{article}(k_2)$
 $\wedge \langle k, x, k_2, h \rangle \in R]\}(\{\langle g, x, y, g \rangle \mid \mathbf{wrote}(x, y)\})$
 $= \{\langle g, x, h \rangle \mid \exists k[g <_2 k \wedge \mathbf{article}(k_2)$
 $\wedge \langle k, x, k_2, h \rangle \in \{\langle g, x, y, g \rangle \mid \mathbf{wrote}(x, y)\}\}$
 $= \{\langle g, x, h \rangle \mid \exists k[g <_2 k \wedge \mathbf{article}(k_2) \wedge \mathbf{wrote}(x, k_2) \wedge k = h]\}$
 $= \{\langle g, x, h \rangle \mid g <_2 h \wedge \mathbf{article}(h_2) \wedge \mathbf{wrote}(x, h_2)\}$

We end up with a meaning for *wrote an article*₂ to which we can apply the meaning of the subject NP *two students*₁, (17d). This is a plural indefinite NP that introduces a new index. The simplified result of applying (17d) to (17c) is given in (17e):

- (17)d. *two students*₁:
 $\lambda P\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge \langle k, k_1, h \rangle \in P]\}$
 e. *two students*₁ wrote an article₂:
 $\{\langle g, h \rangle \mid g <_{1,2} h \wedge \mathbf{2}(\mathbf{students})(h_1) \wedge \mathbf{article}(h_2)$
 $\wedge \mathbf{wrote}(h_1, h_2)\}$

Let us turn now to the second sentence of (17). Anaphoric pronouns are interpreted as the individual that their index, as a discourse entity, refers to. SINGULAR pronouns impose the requirement that their object has cardinality 1. I will write $\mathbf{sg}(x)$ for $\text{card}(x) = 1$. Clearly, this information has the status of a presupposition, but I will not be concerned with the distinction between presuppositions and assertions here. It may seem that PLURAL pronouns require that their reference object has a cardinality greater than one. However, this is not the case; witness the following example:

- (18) Mary wrote [one or two articles]₂. She sent them₂/*it₂ to L&P.

This example allows for Mary having sent only one article; nevertheless, the discourse entity can and indeed must be picked up by a plural pronoun. One possible move would be to analyze plural pronouns as referring to entities of cardinality greater than or equal to one; but this incorrectly predicts that a text like **Mary wrote an article*₂. *She sent them*₂ to L&P is well-formed. Therefore I suggest that the number in plural pronouns is due to syntactic agreement, and not a semantic requirement that can be expressed as a condition on reference objects: The antecedent of a plural pronoun must be grammatically plural, and an NP like *one or two articles* is grammatically plural, in contrast to an NP like *an article*. To keep things simple, I will refrain from expressing grammatical number agreement here, e.g. by sorting the discourse entities in a “singular” type and in a “plural” type (see Kamp and Reyle 1993 for a proposal along these lines).

The second sentence of (17) then is derived in the following way; I render *sent . . . to L&P* as a simple two-place predicate, *sent*:

- (17)f. *it*₂: $\lambda R\{\langle g, x, h \rangle \mid \langle g, x, g_2, h \rangle \in R \wedge \mathbf{sg}(g_2)\}$
 g. *sent . . . to L&P*: $\{\langle g, x, y, g \rangle \mid \mathbf{sent}(x, y)\}$
 h. *sent it*₂ to L&P: $\{\langle g, x, g \rangle \mid \mathbf{sent}(x, g_2) \wedge \mathbf{sg}(x)\}$
 i. *they*₁: $\lambda P\{\langle g, h \rangle \mid \langle g, g_1, h \rangle \in P\}$
 k. *they*₁ sent *it*₂ to L&P: $\{\langle g, g \rangle \mid \mathbf{sent}(g_1, g_2) \wedge \mathbf{sg}(g_2)\}$.

Notice that the second sentence does not change the input assignment; (17k) is a so called “test” in the terminology of Groenendijk and Stokhof (1991). The semantic combination of two sentences is by relational

composition. I will write $\varphi; \psi$ for a text consisting of a text (or a sentence) φ followed by a sentence ψ . If φ', ψ' are the meanings of φ and ψ , then the meaning of $\varphi; \psi$ is $\{\langle g, h \rangle \mid \exists k[\langle g, k \rangle \in \varphi' \wedge \langle k, h \rangle \in \psi']\}$. For our example we get (17l) as the composition of the meaning of (17c) with (17k):

(17)l. *two students₁ wrote an article₂; they₁ sent it₂ to L&P:*

$$\{\langle g, h \rangle \mid g <_{1,2} h \wedge \mathbf{2(student)}(h_1) \wedge \mathbf{article}(h_2) \wedge \mathbf{wrote}(h_1, h_2) \wedge \mathbf{sent}(h_1, h_2) \wedge \mathbf{sg}(h_1)\}$$

Let me illustrate this with a small model. Assume that the input assignment g is empty, that s, s', s'' and s''' are four students, and that s and s' together wrote an article a , and s'' and s''' together wrote an article a' . Then we have as possible output assignment after the first sentence (17e) the two functions

$$\begin{aligned} - h &= [1 \rightarrow \{s, s'\}, 2 \rightarrow a] \\ - h' &= [1 \rightarrow \{s'', s'''\}, 2 \rightarrow a']. \end{aligned}$$

Assume, furthermore, that s and s' sent their article a to L&P, but s'' and s''' failed to do so. Then we find that h , but not h' , is an output assignment after the whole text, (17l).

5. DISTRIBUTIVE READINGS

We now turn to the more complex case of the distributive reading of examples like (1), which is overtly marked in (2). I will assume that distributive readings arise by a distributive operator **EACH** that stands in an anaphoric relationship to an NP that denotes a sum individual (cf. Link 1987). In the cases at hand I will just consider **EACH** as a VP operator. Before I go into the formal derivation, it is perhaps appropriate to give an example of how this operator is supposed to work. Take the following sentence:

(19) Two students₁ **EACH**₁ [wrote an article₂]. They₁ **EACH**₁ [sent it₂ to L&P].

Assume that we start with an empty input assignment g , and assume furthermore that

- (20) – student s wrote article a and sent it to L&P,
 – student s' wrote article a' and sent it to L&P,
 – student s'' wrote article a'' but did NOT send it to L&P.

The possible output assignments after the first sentence of (19) should be the following, which reflect the possible interpretation of *two students* with respect to the given model:

- (20')a. $[1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}]$
 b. $[1 \rightarrow \{s[2 \rightarrow a], s''[2 \rightarrow a'']\}]$
 c. $[1 \rightarrow \{s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}]$

The important thing to notice here is that the output assignments h after the first sentence provide for a way to identify, for each student $x[f]$ that is an element h_1 , the article that this student x wrote: It will be the value of f , the assignment associated with x , when applied to the index 2, that is, f_2 . The second sentence of (19) then filters out assignments (b) and (c) and accepts assignment (a).

Now let us turn to the interpretation rules that yield this result. For the definition of **EACH** we will need a notation for a particular change of a variable assignment. In the classical interpretation of predicate logic, a notation like $g[v/a]$ is used for a variable assignment that is like g , except that the variable v is mapped to the individual a . For our purposes we need a slightly more complex way of changing variable assignments, which is given in the following definition:

- (21) $g[d/\pi]h$ holds iff h is a variable assignment like g , except that every $x[f] \in g_d$ is replaced by $x[f+i]$ such that $\langle g+f, x[f], g+f+i \rangle \in \pi$.

Here, π stands for some dynamic one-place predicate. Let me illustrate this definition with an example. Let g be the following assignment:

$$g = [1 \rightarrow \{s[3 \rightarrow b], s'[3 \rightarrow b']\}]$$

Furthermore, let π be the following predicate (we derived that in (17c) as the meaning of *wrote an article₂*):

$$\pi = \{\langle g, x, h \rangle \mid g <_2 h \wedge \mathbf{article}(h_2) \wedge \mathbf{wrote}(x, h_2)\}$$

Let us assume the model given in (20). Then the only assignment h that stands in the $g[1/\pi]h$ -relation is the one given in (20'a):

$$\text{if } g[1/\pi]h, \text{ then } h = [1 \rightarrow \{s[2 \rightarrow a, 3 \rightarrow b], s'[2 \rightarrow a', 3 \rightarrow b']\}]$$

Notice that, according to the definition (21), each element $x[f]$ in g_1 got

replaced by $x[f + i]$ such that $\langle g + f, x[f], g + f + i \rangle \in \pi$. In particular, $s[3 \rightarrow b]$ got replaced by $s[2 \rightarrow a, 3 \rightarrow b]$, and $s'[3 \rightarrow b']$ got replaced by $s'[2 \rightarrow a', 3 \rightarrow b']$.

Had we started with the following input assignment:

$$g = [1 \rightarrow \{s, s'\}], \text{ which is shorthand for } [1 \rightarrow \{s[\emptyset], s'[\emptyset]\}]$$

with π as above, then the only h for which $g[1/\pi]h$ holds with respect to the model given in (20) would have been the following:

$$\text{if } g[1/\pi]h \text{ then } h = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}]$$

Notice that in contrast to the ordinary notion of assignment variants, this notion does not affect the domain of an assignment g itself, but rather the domain of assignments that are paired with certain entities that are values of g . Also, the type of change is not arbitrary, but is specified using the descriptive apparatus of the language itself by making reference to the predicate π .

Let me now derive example (19) in a way that illustrates the use of the distributivity operator *EACH*, which appears as a VP operator here. The following shows the derivation of the first sentence:

- (19)a. *wrote an article*₂: $\{\langle g, x, h \rangle \mid g <_2 h \wedge \mathbf{article}(h_2) \wedge \mathbf{wrote}(x, h_2)\}$
 b. *EACH*₁: $\lambda P\{\langle g, x, h \rangle \mid x = g_1 \wedge g[1/P]h\}$,
 or $\lambda P\{\langle g, g_1, h \rangle \mid g[1/P]h\}$, for short.
 c. *EACH*₁ *wrote an article*₂:
 $\{\langle g, g_1, h \rangle \mid g[1/\{\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)\}]h\}$
 d. *two students*₁:
 $\lambda P\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge \langle k, k_1, h \rangle \in P]\}$.
 e. *two students*₁ *EACH*₁ *wrote an article*₂:
 $\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge k[1/\{\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)\}]h\}\}$.

Notice that the conjunct $x = g$ in (19b) guarantees that the distribution is over the subject argument of the predicate. With respect to an empty input assignment g and the model given in (20), this relation will give us the three assignments specified in (20') as output. For example, $g (= \emptyset)$ can be changed to $k (= [1 \rightarrow \{s, s'\}])$ by the condition $g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1)$. Then the condition $k[1/\dots]h$ requires us to change every element $x[f]$ of k_1 in a way such that f is extended by i such that $\langle k + f, x[f], k + f + i \rangle$ satisfies the given description, that is,

$$k + f <_2 k + f + i \text{ and } \mathbf{article}([k + f + i]_2) \wedge \mathbf{wrote}(x, [k + f + i]_2).$$

As k_1 is $\{s, s'\}$, which is short for $\{s[\emptyset], s'[\emptyset]\}$, we have $f = \emptyset$. For $x = s$ the model specified in (20) gives us $i = [2 \rightarrow a]$ as the only value for i : We have $k + \emptyset + i = k + i = [1 \rightarrow \{s, s'\}, 2 \rightarrow a]$, hence $[k + i]_2 = a$, and the condition $\mathbf{article}(a) \wedge \mathbf{wrote}(s, a)$ is satisfied. Replacing the empty assignment of $s[\emptyset]$ by i will give us $s[2 \rightarrow a]$. The derivation for $x = s'$ is quite similar; it will give us $i = [2 \rightarrow a']$, and change $s'[\emptyset]$ to $s'[2 \rightarrow a']$. We end up with the output assignment function $h = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}]$. The other two assignments specified in (20') can be construed in similar ways.

The second sentence of (19) is interpreted in the following way:

- (19)f. *they*₁ *EACH*₁ *sent* *it*₂ *to* *L&P*:
 $\{\langle g, h \rangle \mid g[1/\{\langle j, x, j \rangle \mid \mathbf{sent}(x, j_2) \wedge \mathbf{sg}(j_2)\}]h\}$

What is $g[1/\{\langle j, x, j \rangle \mid \dots\}]h$ in this description? According to the definition of assignment variants, h stands in that relation to g iff it is like g , except that every $x[f]$ in g_1 is replaced by $x[f + i]$ such that $\langle g + f, x[f], g + f + i \rangle \in \{\langle j, x, j \rangle \mid \mathbf{sent}(x, j_2) \wedge \mathbf{sg}(j_2)\}$. For this to be the case the input assignment $g + f$ and the output assignment $g + f + i$ must be identical, that is, i must be the empty assignment $\emptyset$. This means that $x[f]$ is not changed at all (or rather, replaced by itself), provided that the condition $\mathbf{sent}(x[f], [g + f]_2) \wedge \mathbf{sg}([g + f]_2)$ is met. This means that we could have written, instead of (19f),

$$\{\langle g, g \rangle \mid g[1/\{\langle j, x, j \rangle \mid \mathbf{sent}(x, j_2) \wedge \mathbf{sg}(j_2)\}]g\}$$

The input assignment $g = (20'a)$ indeed meets this condition with respect to the given model. In particular, g_1 is $\{s[2 \rightarrow a], s'[2 \rightarrow a']\}$, and $x[f] \in g_1$ ranges over $s[2 \rightarrow a]$, $s'[2 \rightarrow a']$. For take $x[f] = s[2 \rightarrow a]$; then $g + f$ is $[1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}, 2 \rightarrow a]$, and $[g + f]_2$ is a . The conditions are met, in particular $\mathbf{sent}(x, [g + f]_2)$. The case of $x[f] = s'[2 \rightarrow a']$ is parallel. The other assignments, (20'b) and (20'c), do not meet this requirement with respect to the model in (20).

The text (19) can be derived by relational composition of (19e) and (19f) according to the same rules that were used before.

A few comments about the crucial definition of assignment variants (21) are in order. First, it is important that the notion of the increment of one assignment by another, e.g. $g + h$, which figures in the definition of assignment variants, be defined with respect to the recursive domain (cf. 15). If we would define the increment with respect to the standard domain,

then indexations like the following ones, in which *an article* and *a picture* share an index, would be well-formed:

- (22) Two students₁ EACH₁ wrote an article₂. Two children₃ EACH₃ drew a picture₂.

Second, the definition of assignment variants makes the input assignment to the whole sentence (*g* in 21) available to the interpretation of *P*, the expression in the scope of EACH. This is necessary, as we may have reference to previously introduced discourse entities:

- (23) A professor₁ came in. Two students₂ EACH₂ [drew a picture₃ and showed it₃ to her₁].

Furthermore, notice that the definition of assignment variants is set up in such a way that the discourse entities that are introduced within the scope of EACH are not accessible at the uppermost level of the resulting assignment function. This predicts that the following text is ill-formed, as the pronoun *it*₂ cannot find an antecedent that meets its number requirement if the two students wrote different articles. That is, the second sentence is forced to have a distributive interpretation.

- (24) Two students₁ EACH₁ wrote an article₂. They₁ sent it₂ to L&P.

The analysis given so far is able to handle the following case, in which the first sentence is interpreted collectively (there is one article), and the second sentence is interpreted distributively (each of the students sent (a copy of) that article to L&P):

- (25) Two students₁ wrote an article₂. They₁ EACH₁ sent it₂ to L&P.

Assume that *s* and *s'* together wrote *a*, and *s* sent *a* to L&P, and *s'* sent *a* to L&P. Given an empty input assignment, the first sentence will output the assignment

$$h = [1 \rightarrow \{s, s'\}, 2 \rightarrow a],$$

and the second sentence will accept this assignment. To see this, notice that the second sentence of (25) will check whether *h* satisfies the following property:

$$h[1/\{\langle j, x, j \rangle \mid \text{sent}(x, j_2) \wedge \text{sg}(j_2)\}]h$$

Now, when we replace in *h* every element $x[f]$ in h_1 (namely, $s[\emptyset]$ and $s'[\emptyset]$, i.e. $f = \emptyset$) by an element $x[f + i]$ such that $\langle h + f, x[f], h + f + i \rangle \in \{\langle j, x, j \rangle \mid \text{sent}(x, j_2) \wedge \text{sg}(j_2)\}$, we get $i = \emptyset$ in both cases;

hence we get back *h* as a result. Notice that *it*₂ can be interpreted, as the index 2 is present in the domain of *h*.

I would like to stress that the analysis of EACH, as presented, certainly needs refinement. For one thing, it needs to be specified that its antecedent denotes a plural object to exclude sentences like **A student EACH₁ wrote an article*. Also, overt *each* can occur in other positions as well, e.g. as an NP modifier, as in *Two students wrote three articles each*. In this paper I will not try to account for these aspects of *each*.

6. CUMULATIVE AND CORRESPONDENCE INTERPRETATIONS

In Section 4 we treated cases of collective readings of sentences like *Two students wrote an article*. We can have collective readings involving two plural NPs, such as *two students wrote three articles*, with the interpretation that two students collaborated in writing three articles. The framework developed in (4) can handle such sentences without any problem. But there is another interpretation of such sentences that appears more clearly with such examples as (25), the CUMULATIVE interpretation (cf. Scha 1981). This sentence may be true in the indicated situation:

- (25) Three students wrote five articles.
True if; e.g., student *s* wrote article *a*,
students *s'* and *s''* wrote article *a'*,
students *s'* and *s''* wrote article *a''*,
student *s''* wrote article *a'''*,
students *s* and *s''* wrote article *a''''*.

In Scha's original treatment of cumulative interpretations, (25) is true only under the additional restriction that no other students besides *s*, *s'* and *s''* wrote articles, and no other articles besides the ones mentioned were written by students. However, it seems that this meaning component only has the status of a (scalar) implicature, and I will disregard it here (but see the end of Section 7.3).

One way to arrive at cumulative interpretations in a general manner is to assume the meaning postulate (26) for lexical transitive predicates *p* (cf. Krifka 1992):

- (26) General cumulativity for two-place predicates *p*:
If $p(x, y)$ and $p(x', y')$, then $p(x \cup x', y \cup y')$.

Interpreting plural NPs like *two students* or *they* as before, the rule for general cumulativity will give us suitable interpretations. For example:

- (25)a. *three students*₁ *wrote five articles*₂:
 $\{\langle g, h \rangle \mid g <_{1,2} h \wedge \mathbf{3}(\mathbf{student})(h_1) \wedge \mathbf{5}(\mathbf{article})(h_2) \wedge \mathbf{wrote}(h_1, h)\}$

Assume an empty input assignment g , and the model given in (25); one possible output assignment after (25a) is

$$h = [1 \rightarrow \{s, s', s''\}, 2 \rightarrow \{a, a', a'', a''', a''''\}],$$

as we have **3(student)**(s, s', s''), **5(article)**(a, a', a'', a''', a''''), and, under the assumption that **wrote** is cumulative, **wrote**($\{s, s', s''\}, \{a, a', a'', a''', a''''\}$). Notice that cumulative interpretations are fairly unspecific under this analysis: From the interpretation given, we cannot infer which student(s) wrote which article(s). This is perhaps as it should be, as cumulative interpretations allow for a wide range of scenarios in which they can be true (cf. Gillon 1987, Lasersohn 1989, Verkuyl 1993, Kamp and Reyle 1993: 414ff. for discussion).

General cumulativity may appear to be the right way to model the meaning of cumulative sentences when considering them in isolation. However, we may need a more fine-grained representation when anaphoric dependencies come into play, as in the following case:

- (27) Three students wrote five articles. They sent them to L&P.

The second sentence of (27) can be interpreted in a number of ways. First, it can be interpreted collectively (i.e., all the students collaborated in sending the articles); in this case, *them* can be treated as regular pronoun. Assuming that *they* and *them* have the indices 1 and 2, respectively, we get the following interpretation:

- (27)a. *they*₁ *sent them*₂ *to L&P*: $\{\langle h, h \rangle \mid \mathbf{sent}(h_1, h_2)\}$

Second, it can be interpreted in a cumulative fashion that leaves open which of the three students sent which of the five articles. The representation format (27a) can be employed for this reading as well if **sent** is cumulative. Assume that the three students sent the five articles in some way or other – for example, s sent a and a' , s' sent a'' , and s'' sent a''' and a'''' . If **sent** is cumulative, then h , the output function of (25a), will be accepted by the interpretation of the second sentence, (27a).

But the second sentence of (27) can also be interpreted as saying that each student (or group of students) sent exactly the articles they have written. In our example, it may say that s sent a , s' and s'' sent a' and a'' , s'' sent a''' , and s and s' sent a'''' . Let me call this the CORRESPONDENCE INTERPRETATION; this interpretation is also recognized by Elworthy (1995).

It is perhaps not entirely clear that this is a distinct interpretation, rather than a preferred model of the general cumulative interpretation. However, note that the correspondence interpretation can be marked by *each* (which shows that it is nothing other than the distributive interpretation.) This is particularly evident in simpler cases for which intuitions are clearer:

- (28) Three students wrote three articles. They each sent them to L&P.

The first sentence, under its cumulative interpretation, is most likely interpreted as saying that each of the three students wrote one of the three articles. The second sentence then says that each of the three students sent the article he or she wrote to L&P; it would be considered false if, say, s wrote article a and s' wrote article a' , but s sent article a' and s' sent article a . We should therefore have a representation of the meaning of the first sentence of (27) that is articulate enough to retrieve this kind of information. The output function h of (25a) does not provide for that kind of information. What we rather need is the output assignment

$$h^* = [1 \rightarrow \{s[2 \rightarrow a], \{s', s''\}[2 \rightarrow \{a', a''\}], s''[2 \rightarrow a'''], \{s, s''\}[2 \rightarrow a''']\}],$$

which does record the specific way in which the students wrote the articles. How can this output assignment be derived from the first sentence of (27), and how can the second sentence of (27) make use of this information?

I propose to give up general cumulativity, (26), and rather assume the following principle of RESTRICTED CUMULATIVITY as a general meaning postulate for lexical two-place predicates:

- (29) Restricted cumulativity for two-place predicates ρ :
- (i) If $\rho(x, y)$ and $\rho(x', y)$, then $\rho(x \cup x', y)$.
 - (ii) If $\rho(x, y)$ and $\rho(x, y')$, then $\rho(x, y \cup y')$.

Given the situation specified in (25), we could not derive **wrote**($\{s, s', s'', \{a, a', a'', a''', a''''\}\}$) under restricted cumulativity; all we can derive is **wrote**(s, a), **wrote**($\{s', s'', \{a'', a'''\}\}$), **wrote**($\{s''', a'''\}$), and **wrote**($\{s, s''\}, a''''$).

In order to arrive at an output assignment like h^* , we will have to assume some kind of distributive interpretation. Only by applying a distributive operator can we construct an output assignment in which the assignments for the object index are subordinated under the assignment of the subject index. Furthermore, the predicates of the nominal arguments will have to apply to the cumulative value of the relevant indices, which we have defined in Section 3, Definition 14. Also, it seems that the indefinite NPs in the scope of a cumulative sentence do not introduce new discourse

entities; rather, they are predicates that measure the size of the set of entities involved in the cumulative interpretation: The first sentence can be paraphrased as 'three students wrote something, and what they wrote is 5 articles all together'. I would like to propose the following derivation

- (27)b. CUM_2^3 :
 $\lambda R \lambda P \{ \langle g, x, i \rangle \mid \exists k [\text{cov}(g, x, 2, k)$
 $\wedge k[2/\{ \langle g, x, h \rangle \mid \exists k [g <_3 k \wedge \langle k, x, k_3, h \rangle \in R \} h$
 $\wedge \langle h, h^3, i \rangle \in P] \}$
 c. *wrote*: $\{ \langle g, x, y, g \rangle \mid \text{wrote}(x, y) \}$
 d. CUM_2^3 *wrote*:
 $\lambda P \{ \langle g, x, i \rangle \mid \exists k [\text{cov}(g, x, 2, k)$
 $\wedge k[2/\{ \langle g, x, h \rangle \mid g <_3 h \wedge \text{wrote}(x, h_3) \} h \wedge \langle h, h^3, i \rangle \in P] \}$
 e. *five articles*: $\{ \langle g, x, g \rangle \mid \mathbf{5}(\text{article})(x) \}$
 f. CUM_2^3 *wrote five articles*:
 $\{ \langle g, x, h \rangle \mid \exists k [\text{cov}(g, x, 2, k) \wedge k[2/\{ \langle g, x, h \rangle \mid g <_3 h$
 $\wedge \text{wrote}(x, h_3) \} h \wedge \mathbf{5}(\text{article})(h^3) \} \}$
 g. *three students*₁:
 $\lambda P \{ \langle g, h \rangle \mid \exists i [g <_1 i \wedge \mathbf{3}(\text{student})(i_1) \wedge \langle i, i_1, h \rangle \in P] \}$
 h. *three students*₁ CUM_2^3 *wrote five articles*:
 $\{ \langle g, h \rangle \mid \exists i, k [g <_1 i \wedge \mathbf{3}(\text{student})(i_1) \wedge \text{cov}(i, i_1, 2, k)$
 $\wedge k[2/\{ \langle g, x, h \rangle \mid g <_3 h \wedge \text{wrote}(x, h_3) \} h \wedge \mathbf{5}(\text{article})(h^3) \} \}$

This derivation makes crucial use of the notion of a *COVER*, that is, a set Y of subsets of X such that $\cup Y = X$ (cf. Gillon 1987 for this notion and its use for cumulative interpretations). It is embodied in the three-place relation *cov* between an input assignment, a discourse index, and an output assignment, defined as follows:

- (30) $\text{cov}(g, x, d, k)$ iff
 – $d \notin \text{DOM}(g)$, $\text{DOM}(g) = \text{DOM}(g) \cup \{d\}$, and $g \subseteq k$.
 – k_d is a set of p -individuals such that
 $\{u \mid \exists (y[f])[y[f] \in x \wedge u \in y]\}$
 $= \{u \mid \exists (y[\emptyset])[y[\emptyset] \in k_d \wedge u \in y]\}$

To see what the meaning (27g) gives us, consider an empty input assignment g , and a model as specified in (25). The first and second conjunct, $g <_1 i$ and $\mathbf{3}(\text{student})(i_1)$, yield the assignment

$$i^* = [1 \rightarrow \{s[\emptyset], s'[\emptyset], s''[\emptyset]\}].$$

Then $\text{cov}(i, i_1, 2, k)$ can change i in a number of ways, for example, to

$$k^1 = i^1 \cup [2 \rightarrow \{s[1], s'[1], s''[1], s'''[1], s''''[1]\}].$$

Notice that the second condition of (30) is satisfied: we have

$$\begin{aligned} & \{u \mid \exists (y[f])[y[f] \in i^*_1 \wedge u \in y]\} \\ &= \{u \mid \exists (y[\emptyset])[y[\emptyset] \in k^*_2 \wedge u \in y]\} = \{s, s', s''\}. \end{aligned}$$

The distributive operator $k[2/\dots]h$ changes k^* further to

$$\begin{aligned} h^* &= i^* \cup [2 \rightarrow \{s[3 \rightarrow a], \{s', s''\}[3 \rightarrow \{a', a''\}], s''[3 \rightarrow a'''], \\ &\{s, s''\}[3 \rightarrow a''']\}] \end{aligned}$$

The assignment h^* satisfies the requirement $k^*[2/\dots]h^*$ in (27h) (cf. (21), Section 5) under the model specified in (25). For example, we have $\text{wrote}(s, a)$, and $\text{wrote}(\{s, s'\}, \{a', a''\})$ (under restricted cumulativity as defined in (29)), and so on. The final condition $\mathbf{5}(\text{article})(h^3)$ is satisfied as well, as $h^{*3} = \cup \{\{a\}, \{a', a''\}, \{a'''\}, \{a''''\}\} = \{a', a', a'', a''', a''''\}$, and all these elements are articles.

The cumulativity operator, as defined in (27b), must have scope over at least one NP (here, the object NP), as this NP cannot be interpreted within the scope of the distribution operator. It will introduce a new index for this NP. I have used a superscript to denote this index (3 in CUM_2^3). As various NPs can be interpreted in a cumulative way (e.g., the subject and the object, or two objects with a ditransitive verb, etc.), we may assume a whole family of CUM-operators.

The subscript 2 of CUM_2^3 indicates a new index that records the specific cover under which the sentence in its cumulative interpretation is true. This information is crucial for the correspondence reading of the second sentences of (27), that is, the reading under which the student(s) that wrote article(s) sent the articles they had written to L&P. This interpretation is actually the distributive interpretation, as derived before. The plural pronoun *them* must be allowed to range over singular entities, but this can be motivated when we assume that its form is due to syntactic agreement with its antecedent, *five articles* (cf. also Example 18).

- (31) *they*₂ *EACH*₂ *sent them*₃ *to L&P*:
 $\{ \langle g, g \rangle \mid g[2/\{ \langle g, x, g \rangle \mid \text{sent}(x, g_3) \}]g \}$

Notice that *they*₂ picks out the index 2 that has as its value the cover as used for the interpretation of the first sentence. Assume the input function h^* , and assume that each student or group of students sent the articles they had written to L&P, then h^* will be accepted by (31). Again, it is necessary that *sent* is interpreted according to restricted cumulativity as defined in (28).

Under a suitable analysis of partitive NPs (see Section 7.5) we can also give an intuitively correct interpretation to the following text. Notice that

the object pronoun *them* does not necessarily refer to all the articles invoked by the first sentence:

- (32) Three students₁ CUM₂³ wrote five articles. [Two of them₂]₄ EACH sent them₃ to L&P.

This means that two of the students sent all their articles to L&P. In the scenario given in (25), the second sentence appears to be true if, for example, *s* sent *a*, *s'* sent *a''*, and *s* and *s'* together sent *a'''*. We get this reading as a distributive interpretation with EACH₄. The partitive noun phrase has to be interpreted in the following way:

- (33) [two of them₂]₄:
 $\lambda P\{\langle g, h \rangle \mid \exists k[g <_4 k \wedge k_4 \subseteq k_2 \wedge$
 $\# \cup\{x \mid x[f] \in k_4\} = 2 \wedge \langle k, k_4, h \rangle \in P\}$

That is, this phrase introduces a new index 4 that is interpreted as a subset of the value of g_2 , where the union of set of the first members of the value of 4 should have cardinality two.

After we have developed a representation for sentences containing anaphoric reference in their correspondence interpretation, let us now turn to such sentences in their collective and cumulative interpretation. Take the second sentence of example (27), and assume that the students sent the articles collectively. The following representation will give us this reading:

- (27)i. *them*³: $\lambda R\{\langle g, x, g \rangle \mid \langle g, x, g^3, g \rangle \in R\}$
 j. *they*₁ sent *them*³ to L&P: $\{\langle g, g \rangle \mid \text{sent}(g_1, g^3)\}$

Notice that the pronoun *them* picks out the cumulative value of the input assignment with respect to the index 3, which is indicated by a superscript in the syntactic representation. In the example under consideration, with $g = h^*$ as input, g^3 is $\{a, a', a'', a''', a''''\}$, g_1 is $\{s, s', s''\}$, and $\text{sent}(g_1, g^3)$ expresses a collective interpretation.

The cumulative interpretation of the second sentence of (27) is the one under which the students who sent their articles need not correspond to the way they wrote them (the way that was contemplated with example (27a)). When we assume general cumulativity for two-place predicates (26), then (27j) would represent the cumulative interpretation as well. However, this has various disadvantages. It would effectively collapse the collective and the cumulative interpretation, and it would force us to assume general cumulativity again (cf. 26) which we abandoned for good reasons. Also, it would make us lose all the information about the way

in which the articles were sent. We may need this information. For example, (27) could be extended in the following way:

- (34) Three students wrote five articles. They sent them to L&P. They got them back within a month.

This first two sentences can be interpreted in the way I have just characterised, that is, the students may have sent the articles in a different way from how they wrote them – say, *s* and *s'* sent *a* and *a'*, *s'* sent *a''* and *a'''*, and *s''* sent *a''''*. The third sentence then can have a correspondence interpretation with respect to the second sentence, that is, it could express that *s* and *s'* got back *a* and *a'*, that *s'* got back *a''* and *a'''*, and that *s''* got back *a''''*. We cannot express this within the representation format illustrated by (27j). Rather, we must assume that the cumulativity operator is involved:

- (27) k. CUM₄⁵ sent:
 $\lambda P\{\langle g, x, i \rangle \mid \exists k, h[\text{cov}(g, x, 4, k) \wedge$
 $k[4/\{\langle g, x, h \rangle \mid g <_5 h \wedge \text{sent}(x, h_5)]h \wedge \langle h, h^5, i \rangle \in P]\}$
 l. *them*³: $\{\langle g, x, g \rangle \mid x = g^3\}$
 m. CUM₄⁵ sent *them*³:
 $\{\langle g, g_1, h \rangle \mid \exists k[\text{cov}(g, x, 4, k) \wedge$
 $k[4/\{\langle g, x, h \rangle \mid g <_5 h \wedge \text{sent}(x, h_5)]h \wedge h^5 = h^3]\}$
 n. *they*₂ CUM₄⁵ sent *them*³:
 $\{\langle g, h \rangle \mid \exists k[\text{cov}(g, g_1, 4, k) \wedge k[4/\{\langle g, x, h \rangle \mid g <_5 h$
 $\wedge \text{sent}(x, h_3)]h \wedge h^5 = h^3]\}$

To see what is going on, assume h^* , the output of (27h), as input to (27n) and assume as before that *s* and *s'* sent *a* and *a'*, that *s'* sent *a''* and *a'''*, and that *s''* sent *a''''*. Now, we have that

$$h_2^* = \{s[3 \rightarrow a], \{s', s''\}[3 \rightarrow \{a', a''\}], s''[3 \rightarrow a'''], \{s, s''\}[3 \rightarrow a'''']\}.$$

The condition $\text{cov}(h^*, h_2^*, 4, k)$ can deliver the following assignment, among other options:

$$k^* = h^* \cup [4 \rightarrow \{\{s, s'\}[\emptyset], s'[\emptyset], s''[\emptyset]\}]$$

The assignment k^* in turn will be changed by the condition $k[4/\dots]h$ of (27n) to the assignment

$$h^* = h^* \cup [4 \rightarrow \{\{s, s'\}[5 \rightarrow \{a, a'\}], s'[5 \rightarrow \{a'', a'''\}], s''[5 \rightarrow a'''']\}]$$

The final condition is satisfied with respect to this assignment: We have $h^{1,3} = h^{1,3} \cup \{a, a', a'', a''', a''''\}$. Now the subject *they* of the third sentence

of (34) could pick up the index 4, leading to a correspondence reading with respect to the second sentence of (34).

I have analyzed cumulative interpretations as involving a distributive interpretation over “covers”, which is structurally quite similar to the interpretation of a distributive sentence. As we have seen, this enables subsequent collective, cumulative, and correspondence interpretations that involve pronouns. We should then expect similar possible continuations for distributive sentences, as in the following example:

- (35) Three students₁ EACH₁ wrote an article₂. They₁ sent them to L&P.
- They₁ sent them² to L&P.*
 - They₁ CUM₃⁴ sent them₂ to L&P.*
 - ?They₁ EACH₁ sent them₂ to L&P.*
 - They₁ EACH₁ sent them² to L&P.*

First, the second sentence can be interpreted collectively. To achieve this, the pronoun *them* must be interpreted as *them²* (cf. (27i)); it then would pick out all three articles (cf. (35a)). Then the second sentence can get a cumulative interpretation, as represented in (35b). Furthermore, a correspondence interpretation is possible, as indicated in (35c); this is somewhat marginal, as it would be synonymous with the more specific and otherwise equivalent *They₁ EACH₁ sent it₂ to L&P*. We find additional combinations of *EACH* and *CUM* and subscripted/superscripted pronouns, for example, (35d), which expresses that each of the three students separately sent (copies of) all the three articles.

7. NOMINAL QUANTIFIERS AND ANAPHORA

In this section I will show how parametrized sum individuals can be used to model certain cases of anaphora involving nominal quantifiers. For example, I will develop a representation of partitive quantifiers like *most students that wrote an article*, for which Rooth (1987) has suggested the use of parametrized individuals in the first place. In order to show how parametrized sum individuals can be put to service for such examples, we have to commit ourselves to specific treatments of the relation of determiners and nouns (Section 7.1), of relative clauses (Section 7.2), and of non-anaphoric definite NPs (Section 7.3). In Section 7.4 I will deal with quantifiers and anaphora. Section 7.4 will then discuss partitive quantifiers, and Section 7.6 will be concerned with issues of distributivity and collectivity in sentences containing quantifiers.

7.1. Determiners

Let me start with some remarks about the compositional semantics of NPs, in particular regarding the role of determiners. The treatment of indefinite NPs like *two students₁* that I have proposed so far abstracts away from the issue of which constituent actually is responsible for the introduction of the index. Clearly, this should be the determiner of a NP. But is *two* the determiner of *two students*? Perhaps so, but notice that we have NPs like *the two students*, in which there is another determiner, *the* and *two* plays just the role of a number word, a kind of adjective (cf. Link 1987 for this analysis). Hence it will be appropriate to either postulate an empty indefinite determiner for the noun phrase *two students* that is responsible for the introduction of an index, or to assume that *two* can either be interpreted as a simple number word or as a full determiner. Let me develop the first view and come back to the second at the end of this section.

I will assume a morphologically empty determiner, for which I will use the symbol $\exists_1$ that introduces a new index. The meaning of $\exists_1$ then can be given as follows:

$$(36) \quad \exists_1: \lambda P\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \langle k, k_1, h \rangle \in P]\}$$

But this fails to represent another aspect of determiner meaning (at least as seen in Generalized Quantifier Theory), namely that determiners express how an NP and a verbal predicate should be combined. This meaning component could be captured by operators like those in (37) for subject NPs and object NPs. (I assume that the syntactic component forces these operators to have the same index as their argument NP.)

- (37)a. SUBJ₁: $\lambda p \lambda P\{\langle g, h \rangle \mid \exists k[\langle g, k \rangle \in p \wedge \langle k, k_1, h \rangle \in P]\}$
 b. OBJ₁: $\lambda p \lambda R\{\langle g, x, h \rangle \mid \exists k[\langle g, k \rangle \in p \wedge \langle k, x, k_1, h \rangle \in R]\}$

The meaning of *two students₁* can be derived in the following way; recall that the plural form *students* is due to number agreement (cf. Section 3).

- (38)a. *two students*: $\{\langle g, x, g \rangle \mid \mathbf{2}(\text{student})(x)\}$
 d. $\exists_1$ *two students*: $\{\langle g, h \rangle \mid g <_1 h \wedge \mathbf{2}(\text{student})(h_1)\}$
 e. SUBJ₁ $\exists_1$ *two students*:
 $\lambda P\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\text{student})(k_1) \wedge \langle k, k_1, h \rangle \in P]\}$

We may assume that SUBJ₁, $\exists_1$ and *two* can be composed into one operator, the determiner *two₁SUBJ₁*, which has the following interpretation:

$$(39) \quad \text{two}_{1\text{SUBJ}_1}: \lambda P' \lambda P\{\langle g, h \rangle \mid \exists k, i[g <_1 k \wedge \text{card}(k_1) = 2 \wedge \langle k, k_1, i \rangle \in P' \wedge \langle i, i_1, h \rangle \in P]\}$$

But it is good to keep in mind that the meaning of *two* as a subject determiner embodies three relatively simple meaning constituents, an existential quantifier, a number specification, and an identification mechanism for a verbal argument slot.

7.2. Relative Clauses

Let us turn to relative clauses. I will concentrate here on RESTRICTIVE relative clauses, which are modifiers of N. The relative clause itself consists of a relative pronoun and a sentence that contains an empty element with the same index. This identity of indices is not captured in the current framework for reasons of simplicity, but could be enforced by either syntactic movement or feature percolation. Empty elements are interpreted like pronouns, that is, their index must be already in the domain of the input assignment. Relative pronouns pick up an old index as well, and for grammatical reasons they have to pick up the index of their head noun. The following example illustrates all this with the example *two students that (each) wrote an article*:

- (40)a. *EACH₁ wrote an article₂*:
 $\{\langle g, g_1, h \rangle \mid g[1/\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h\}$
 b. *e₁*: $\lambda P\{\langle g, h \rangle \mid \langle g, g_1, h \rangle \in P\}$
 c. *e₁ EACH₁ wrote an article₂*:
 $\{\langle g, h \rangle \mid g[1/\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h\}$
 d. *that₁*: $\lambda p \lambda P\{\langle g, g_1, h \rangle \mid \exists k[\langle g, g_1, k \rangle \in P \wedge \langle k, h \rangle \in p]\}$
 e. *that₁ [e₁ EACH₁ wrote an article₂]*:
 $\lambda P\{\langle g, g_1, h \rangle \mid \exists k[\langle g, g_1, k \rangle \in P \wedge$
 $k[1/\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h]\}$
 f. *two students*: $\{\langle g, x, g \rangle \mid \mathbf{2}(\mathbf{student})(x)\}$
 g. *two students that₁ [e₁ EACH₁ wrote an article₂]*:
 $\lambda P\{\langle g, g_1, h \rangle \mid \mathbf{2}(\mathbf{student})(g_1) \wedge$
 $g[1/\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h]\}$
 h. $\exists_1$ *two students that₁ [e₁ EACH₁ wrote an article₂]*:
 $\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge$
 $k[1/\langle i, x, j \rangle \mid i <_2 j \wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h]\}$
 i. SUBJ₁ [$\exists_1$ *two students that₁ [e₁ EACH₁ wrote an article₂]]:
 $\lambda P\{\langle g, h \rangle \mid \exists k, f[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge k[1/\langle i, x, j \rangle \mid i <_2 j$
 $\wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]f \wedge \langle f, f_1, h \rangle \in P]\}$*

We can apply this NP to a VP like *EACH₁ sent it₂ to L&P*, as in (19f) which gives us the following result:

- (40)j. [SUBJ₁ $\exists_1$ *two students that₁ e₁ EACH₁ wrote an article₂] *EACH₁ sent it₂ to L&P*:
 $\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \mathbf{2}(\mathbf{student})(k_1) \wedge k[1/\langle i, x, j \rangle \mid i <_2 j$
 $\wedge \mathbf{article}(j_2) \wedge \mathbf{wrote}(x, j_2)]h \wedge h[1/\langle i, x, i \rangle \mid \mathbf{sent}(x, i_2)$
 $\wedge \mathbf{sg}(i_2)]h]\}$*

To understand what is going on here, take the model introduced in (20) above and an empty input assignment g . In a first step g can be extended to $k = [1 \rightarrow \{s, s'\}]$. The assignment k in turn will be extended to an h , $h = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}]$, by the conjunct $k[1/\dots]h$. The conjunct $h[1/\dots]h$ is just a test, and accepts the assignment h .

7.3. Definite NPs

The modelling of partitive quantifiers require a way of treating non-anaphoric definite NPs, such as *the students that wrote an article*. Following Link (1983), the non-anaphoric definite article is interpreted as the supremum of the entities in the denotation of a predicate, if this is in the denotation of the predicate as well. I will call this individual the MAXIMAL individual in the denotation of a predicate. For example, the meaning of *the students* is the supremum of the denotation of *students*, which is always in the denotation of *students* due to the cumulativity of this predicate. The meaning of *the three students* is the supremum of the denotation of *three students* provided that it is the denotation of *three students* as well; this is the case if and only if there are exactly three students. Otherwise, this NP fails to refer. Thus we get the usual uniqueness condition in the special case of non-cumulative, or quantized, predicates.

We can integrate this notion of maximal individuals into the current dynamic framework by defining a notion of a MAXIMAL EXTENSION for dynamic propositions. One way of doing this is the following:

- (41)a. $\langle g, h \rangle \leq \langle g', h' \rangle$, the assignment pair $\langle g, h \rangle$ is SUBORDINATED to the assignment pair $\langle g', h' \rangle$, iff
 – $g = g'$
 – $\text{DOM}(h) = \text{DOM}(h')$
 – and for every d with $d \in \text{DOM}(h)$; $h_d \subseteq h'_d$.
 b. $\text{MAX}(p)$, the MAXIMAL INTERPRETATION of the dynamic proposition p , is the set of assignment pairs $\langle g, h \rangle$ such that:
 $\langle g, h \rangle \in p$
 for every h' such that $\langle g, h' \rangle \in p$ it holds that $\langle g, h \rangle \leq \langle g, h' \rangle$.

The meaning of the non-anaphoric definite article then can be rendered as follows; I give as an example a subject determiner with index 1.

$$(42) \quad the_1: \lambda P[\text{MAX}(\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \langle k, k_1, k \rangle \in P]\})]$$

Let me illustrate how this operator works with two examples, one implying a cumulative interpretation, the other one a distributive interpretation.

$$(43) \quad the_1 [students \text{ that }_1 e_1 \text{ CUM}_2^3 \text{ wrote articles}]: \\ \text{MAX}(\{\langle g, h \rangle \mid \exists k, i[g <_1 k \wedge \text{students}(k_1) \wedge \text{COV}(k, k_1, 2, i) \\ \wedge i[2/\langle i, x, h \rangle \mid i <_3 h \wedge \text{wrote}(h_1, h_2)]h \wedge \text{articles}(h^3)]\})$$

Assume that students s, s' wrote article a together, and s'' wrote articles a' and a'' , and that no other students wrote any articles. Then the following tuples are elements of the meaning of the argument of MAX in (43), under the assumption that the plural predicates **students** and **articles** apply to single students and articles as well.

$$(44) \text{a. } \langle \emptyset, [1 \rightarrow s'', 2 \rightarrow s''[3 \rightarrow \{a', a''\}]] \rangle \\ \text{b. } \langle \emptyset, [1 \rightarrow \{s, s'\}, 2 \rightarrow \{s, s'\}[3 \rightarrow a]] \rangle \\ \text{c. } \langle \emptyset, [1 \rightarrow \{s, s', s''\}, 2 \rightarrow \{\{s, s'\}[3 \rightarrow a], s''[3 \rightarrow \{a', a''\}]\}] \rangle$$

Clearly, all these pairs are subordinated to (44c); hence (44c), but not (44a) nor (44b), is an element of (43). – The following example illustrates MAX with a distributive predicate:

$$(45) \quad the_1 [students \text{ that }_1 e_1 \text{ EACH}_1 \text{ wrote an article}_2]: \\ \text{MAX}(\{\langle g, h \rangle \mid \exists k[g <_1 k \wedge \text{students}(k_1) \wedge k[1/\langle i, x, j \rangle \mid i <_2 j \\ \wedge \text{article}(j_2) \wedge \text{wrote}(x, j_2)]h]\})$$

Assume now that the students s, s' and s'' wrote the articles a, a' and a'' respectively, and that s, s' and s'' are the only students that wrote any articles. Then the following pairs are elements of the meaning of the arguments of MAX in (45). Clearly, (46d) subsumes all the other assignments, and hence is an element of (46).

$$(46) \text{a. } \langle \emptyset, [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\}] \rangle \\ \text{b. } \langle \emptyset, [1 \rightarrow \{s[2 \rightarrow a], s''[2 \rightarrow a'']\}] \rangle \\ \text{c. } \langle \emptyset, [1 \rightarrow \{s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}] \rangle \\ \text{d. } \langle \emptyset, [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}] \rangle$$

Let us study how a phrase like (45) is integrated in a sentence, e.g. in subject position. We can apply the operator SUBJ, defined in (37a), to it, and then apply the result to a verbal predicate:

$$(47) \text{a. SUBJ}_1 the_1 students \text{ that }_1 e_1 \text{ EACH}_1 \text{ wrote an article}_2: \\ \lambda P[\langle g, h \rangle \mid \exists k[\langle g, k \rangle \in \text{MAX}(\{\langle g, k \rangle \mid \exists f[g <_1 f \wedge \text{students}(f_1) \\ \wedge f[1/\langle i, x, j \rangle \mid i <_2 j \wedge \text{article}(j_2) \wedge \text{wrote}(x, j_2)]k]\}) \\ \wedge \langle k, k_1, h \rangle \in P]] \\ \text{b. SUBJ}_1 the_1 students \text{ that }_1 e_1 \text{ EACH}_1 \text{ wrote an article}_2 \text{ EACH}_2 \text{ sent} \\ it_1 \text{ to L\&P:} \\ \{\langle g, h \rangle \mid \langle g, h \rangle \in \text{MAX}(\{\langle g, h \rangle \mid \exists f[g <_1 f \wedge \text{students}(f_1) \\ \wedge f[1/\langle i, x, j \rangle \mid i <_2 j \wedge \text{article}(j_2) \wedge \text{wrote}(x, j_2)]h]\}) \\ \wedge h[1/\langle i, x, i \rangle \mid \text{sent}(x, i_2)]h]\}$$

To see that we get the right result, assume the model as before for (45), and that s, s' and s'' sent their own articles to L&P. Then an empty assignment g is changed to an output assignment $h = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}]$ that satisfies the necessary requirements. In particular, h is the maximal assignment for which it holds that h_1 are students such that they each wrote an article (which is associated with index 2), and it also holds that each element x of h_1 sent x 's article to L&P.

Maximal interpretations are not only needed for definite NPs, but also for the modelling of interpretations that are pragmatically strengthened by pragmatic implicature. For example, a sentence like *Mary wrote three articles* is typically understood as 'the number of articles that Mary wrote is three', although it actually means just 'the number of articles that Mary wrote is at least three'. Similarly, a sentence like *Three students wrote seven articles* is typically interpreted as 'the number of students that wrote articles is three, and the number of articles written by students is seven' (cf. Scha 1981). We can render such interpretations by assuming that a sentence $\{\langle g, h \rangle \mid \dots\}$ is pragmatically strengthened to $\text{MAX}(\{\langle g, h \rangle \mid \dots\})$ under certain circumstances. Also, pragmatic strengthening may explain a certain effect observed by Evans (1980) with texts like *Harry owns some sheep*₁. *Tom vaccinated them*₁. It has been observed that the second sentence is preferably interpreted as 'Tom vaccinated (ALL) THE SHEEP THAT HARRY OWNS', which comes as a surprise under the usual analysis of indefinites like *some sheep* in frameworks of dynamic interpretation. But note that we would arrive at this reading by assuming that the first sentence, *Harry owns some sheep*₁, is pragmatically strengthened by MAX. Then the index 1 would be anchored to the maximal individual x such that x are some sheep that Harry owns. Thus, the major motivation for the analysis of donkey pronouns as definite descriptions (cf. Heim 1990) can be circumvented by assuming pragmatic strengthening of the sentence that contains the antecedent.

7.4. Quantificational NPs

Now we are in a position to deal with cases that involve partitive and other quantifiers. One important phenomenon is that, contrary to the assumptions of classical DRT, quantifiers license anaphoric reference. For one thing, they license plural anaphoric reference to the RESTRICTOR expressed by the nominal predicate. In the following examples, *they* clearly can refer to the students in the domain of discourse:

- (48)a. No student wrote an article. They (all) spent their days on the beach.
 b. Most students wrote an article. They (all) are quite advanced.
 c. Each student wrote an article. They (all) are quite advanced.

Furthermore, reference to the INTERSECTION of nominal predicate and verbal predicate is possible if it is not empty, as illustrated in the following example:

- (49) Most students wrote an article. They sent them to L&P.

It has been claimed by Moxley and Sanford (1987) that we may also refer to the complement set of the intersection in cases of downward-entailing quantifiers, as in the following example:

- (50) Few students wrote an article. They rather spent their days on the beach.

However, it seems to me that in such cases *they* rather refers to the restrictor, and the predication is vague. The second sentence of (50) can easily be rendered *The students rather spent their days on the beach*, with the few that did not just being exceptions that are not worth talking about. I will disregard such cases here.

The analysis I would like to propose will be illustrated with the meaning of *most*. As quantifiers introduce two discourse entities, one for the restrictor and another for the intersection, they come with two indices. The following example illustrates *most* as a subject quantifier, with indices 1 for the restrictor and 3 for the intersection.

- (51) $most_{1,3}: \lambda P' \lambda P \{ \langle g, h \rangle \mid$
 $\exists k, f [\langle g, k \rangle \in \text{MAX}(\{ \langle g, k \rangle \mid \exists i [g <_1 i \wedge \langle i, i_1, k \rangle \in P'] \})$
 $\wedge k <_3 f \wedge f_3 \subseteq k_1 \wedge \langle f, h \rangle \in \text{MAX}(\{ \langle f, h \rangle \mid \langle f, f_3, h \rangle \in P \})$
 $\wedge \text{card}(h_3)/\text{card}(h_1) > \frac{1}{2}] \}$

Let me illustrate how things work with an example involving distributive quantification:

- (52)a. *students that*₁ *e*₁ *EACH*₁ *wrote an article*₂:
 $\{ \langle g, x, h \rangle \mid \text{students}(x) \wedge g[1/\{ \langle i, x, j \rangle \mid i <_2 j \wedge \text{article}(j_2) \wedge \text{wrote}(x, j) \}]h \}, \text{abbr. [A]}$
 b. *EACH*₃ *sent it*₂ *to L&P*:
 $\{ \langle g, g_3, g \rangle \mid g[3/\{ \langle i, x, i \rangle \mid \text{sent}(x, i_2) \}]g \}, \text{abbr. [B]}$
 c. *most*_{1,3} *students that*₁ *e*₁ *EACH*₁ *wrote an article*₂ *EACH*₃ *sent it*₂ *to L&P*:
 $\{ \langle g, h \rangle \mid \exists k [\langle g, k \rangle \in \text{MAX}(\{ \langle g, k \rangle \mid \exists i [g <_1 i \wedge \langle i, i_1, k \rangle \in [A]] \})$
 $\wedge k <_3 h \wedge h_3 \subseteq h_1 \wedge \langle h, h \rangle \in \text{MAX}(\{ \langle h, h \rangle \mid \langle h, h_3, h \rangle \in [B] \})$
 $\wedge \text{card}(h_3)/\text{card}(h_1) > \frac{1}{2}] \}$

To check that this is the correct interpretation, assume an empty input assignment g , and assume a model in which the students s, s' and s'' wrote the articles a, a' and a'' , respectively, and in which s sent a to L&P, and s' sent a' to L&P. No other student wrote any article. Then g will first be extended to a k with

$$k = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\},$$

as this is the maximal assignment for which it holds that $\exists i [\emptyset <_1 i \wedge \langle i, i_1, k \rangle \in [A]]$: notice that i will be $[1 \rightarrow \{s, s', s''\}]$ in this case. This assignment k in turn is extended to h , with

$$h = [1 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\},$$

$$3 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a']\},$$

as h is an assignment that satisfies $h_3 \subseteq h_1$ and furthermore is the maximal assignment that satisfies the condition $\langle h, h_3, h \rangle \in [B]$, and it satisfies $\text{card}(h_3)/\text{card}(h_1) > \frac{1}{2}$.

Notice that we have introduced the indices 1 and 3 for good; they stand for the the students that each wrote an article, and the students that in addition sent their article to L&P. Consequently, these indices can be picked up by pronouns, as illustrated in examples (48) and (49). Furthermore, a pronoun with index 2 can pick up the article of each student in a distributive sentence (cf. 53a,b). Also, a cumulative pronoun with the index 2 can pick up the articles that the students have written (cf. 53c).

- (53) *Most*_{1,3} *students that* *e*₁ *EACH*₁ *wrote an article*₂ *EACH*₃ *sent it*₂ *to L&P*.
 a. *They*₃ *EACH*₃ *got it*₂ *published within a year*.
 (i.e. s got a published, s' got a' published)
 b. *They*₁ *EACH*₁ *got it*₂ *published within a year*.
 (i.e. s got a published, s' got a' published, and s'' got a'' published)

- c. They² were written on acid-free paper.
(i.e. a , a' and a'' are written on acid-free paper)

However, the interpretation of (53d) with the intended reading that $\{a, a'\}$ were published cannot be construed in the framework given so far:

- (53)d. They were published within a year.

The only relevant way to interpret *they* is as a cumulative pronoun *they*², which will pick out the sum individual $\{a, a', a''\}$. We may introduce another type of pronouns that are "dependent" on certain indices. For example, we may assume that *they* in (53d) can be interpreted as *they*²⁽³⁾, which, given an input assignment h , is interpreted as referring to $\cup\{x \mid \exists(y[f])[y[f] \in h_3 \wedge x = f_2]\}$. This shows that the data structure of variable assignments with parametrized sum individuals developed here is able to capture cases like (53d) as well.

The formal representation of quantifiers like *most*_{1,3} allows for discourse entities other than 1 or 2 that are introduced in the restrictor or the matrix to remain accessible for future discourse. This seems to be necessary for cases like the following:

- (54)a. Most_{1,3} students that₁ e_1 wrote articles₂ sent them to L&P.
They₂ were written on acid-free paper.
b. Most_{1,3} students wrote articles₂. They₂ were written on acid-free paper.

Although the restrictive clause in (54a) and the matrix in (54b) are not interpreted in a distributive way, reference to the articles in subsequent discourse is possible. The current representation predicts this. Notice that, due to the maximal interpretation of the restrictor and the matrix, *they*₂ in (54a) and (54b) will pick up all the articles that were written by students.

7.5. Partitive Quantifiers

Rooth's original motivation for parametrized sum individuals were cases of partitive quantification, as in *most of the students that wrote an article*. Such cases can be dealt with if we assume that the *of*-phrase creates a new partitive predicate from the meaning of a definite NP. The following interpretation gives us what we want. I assume here that *of* carries an index that is forced to be identical to the index of the definite NP in its argument position.

- (55) $of_i: \lambda p[\lambda x[\lambda y[\lambda g, h \in p \wedge x \in h_i]]]$

This interpretation yields a predicate that applies to subsets of the individual associated with the definite NP. Example:

- (56) $of_4 the_4 students that_4 e_4 each_4 wrote an article_2$:
 $\{\langle g, x, h \rangle \mid \langle g, h \rangle \in \text{MAX}(\{\langle g, h \rangle \mid \exists k[g <_4 k \wedge \text{students}(k_4) \wedge k[4\{\langle i, x, j \rangle \mid i <_2 j \wedge \text{article}(j_2) \wedge \text{wrote}(x, j_2)\}h]\}) \wedge x \subseteq h_4\}$, abbr. [C].

If the students s , s' and s'' wrote the articles a , a' and a'' , respectively, then the following triples are in this set:

- (57) $\langle \emptyset, x, [4 \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}] \rangle$,
where x is one of the following p -individuals:
 $s[2 \rightarrow a]$,
 $s'[2 \rightarrow a']$,
 $s''[2 \rightarrow a'']$,
 $\{s[2 \rightarrow a], s'[2 \rightarrow a']\}$,
 $\{s[2 \rightarrow a], s''[2 \rightarrow a'']\}$,
 $\{s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}$,
 $\{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}$

The predicate derived in (56) in turn can be part of a quantificational NP and a sentence, as illustrated in the following example:

- (58) $most_{1,3} of_4 the_4 students that_4 e_4 EACH_4 wrote an article_2 EACH_3 sent it_2 to L\&P$:
 $\{\langle g, h \rangle \mid \exists k[\langle g, k \rangle \in \text{MAX}(\{\langle g, k \rangle \mid \exists i[g <_1 i \wedge \langle i, i_1, k \rangle \in [C]]\}) \wedge k <_3 h \wedge h_3 \subseteq h_1 \wedge \langle h, h \rangle \in \text{MAX}(\{\langle h, h \rangle \mid \langle h, h_3, h \rangle \in [B]\}) \wedge \text{card}(h_3)/\text{card}(h_1) > \frac{1}{2}]\}$

With respect to the model given above and an empty input assignment g , the various extensions k that satisfy the condition $\exists i[g <_1 i \wedge \langle i, i_1, k \rangle \in [C]]$ are of the following form,

$$[1 \rightarrow x, r \rightarrow \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}],$$

where x ranges over the same p -individuals as in (57). The maximal assignment among these assignments is, of course, the one for which $x = \{s[2 \rightarrow a], s'[2 \rightarrow a'], s''[2 \rightarrow a'']\}$. The net result for the interpretation of (58) is very similar to example (51c), except that the output assignment h contains another discourse entity, 4, which is anchored to the same individual as 1.

7.6. Further Comments on Quantifiers

Several comments are in order about the representation of quantified NPs developed in Sections (7.4) and (7.5). First, the fact that the quantifier does not relate two sets of assignment functions, but rather two sum individuals (that may be associated with assignment functions) has the consequence that only the so-called *ASYMMETRIC* interpretation is generated (cf. Kadmon 1990). This is a welcome result, as nominal quantifiers only have this interpretation.

Another welcome property of the treatment of quantification proposed here is that it can deal with cumulative and collective quantifications in the restrictor and in the matrix that many speakers allow, as in the following cases:

- (59)a. $\text{Most}_{1,3}$ [(of the) students] gathered in the hallway.
 b. $\text{Most}_{1,3}$ [(of the) students that gathered in the hallway] *EACH* carried a backpack.

I have just indicated the indices of the quantificational determiner *most* here. Notice that the predication *gathered in the hallway* in (59a) is “about” what the discourse entity 3 stands for (the intersection). This is a plural individual, and hence the basic requirement for this type of interpretation is satisfied. Similarly, the predication *gathered in the hallway* in (59b) is about what the discourse entity 1 stands for (the noun meaning), which is a plural individual as well. In contrast, the standard analysis of quantifiers in DRT embodies a distributive interpretation in both cases, which would lead to a classification of (59a) and (59b) as ungrammatical.

For singular quantifiers like *every* we seem to have a built-in distributive reading for the restrictor and the matrix. This indicates that these quantifiers have to be analyzed differently from the plural quantifiers. Also, notice that due to the singular form of the head noun the resulting restrictor predicate is not cumulative, and therefore the maximal element of the predicate is undefined in most cases – there is no maximal individual in the extension of, say, *student*, except if there is exactly one student. This is another reason why the interpretation scheme illustrated with *most* does not work for *every*. But anaphoric reference to the restrictor set is still possible (cf. 60a,b), as well as cumulative reference to entities introduced within the restrictor (60b) or the matrix (60c):

- (60) Every guest who brought a present gave it to a child.
 a. They *EACH* had picked it out carefully. (*they*: the guests)
 b. They had picked them out carefully. (*they*: the guests; *them*: the presents)

- c. They did not like them. (*they*: the children; *them*: the presents)

One interpretation of *every* that predicts all these cases is the following. The determiner *every* is associated with two indices, but now one index (here, 4) is associated with the sum of the restrictor individuals, the other (here, 1) serves to identify each element of this sum.

- (61)a. *guest* that₁ *e*₁ brought a present₂:
 $\{\langle g, g_1, h \rangle \mid g \vdash g_1 \wedge \mathbf{guest}(h_1) \wedge \mathbf{present}(h_2) \wedge \mathbf{brought}(h_1, h_2)\}$,
 abbr. [D]
 b. *gave* it₂ to a child₃:
 $\{\langle g, x, h \rangle \mid g \vdash x \wedge \mathbf{child}(h_3) \wedge \mathbf{gave}(x, h_2, h_3)\}$, abbr. [E]
 c. *every*_{1,4} *guest* that₁ *e*₁ brought a present₂ gave it₂ to a child₃
 $\{\langle g, h \rangle \mid \exists k[k = g + [4 \rightarrow \cup\{x[f] \mid \exists i[g <_1 i \wedge x = i_1 \wedge \langle i, i_1, g + f \rangle \in [D]]] \wedge k[4[E]]h]\}$

To see how this works, assume that there are exactly two guests *a*, *a'* that brought presents, namely *b*, *b'*, respectively, and that *a* gave *b* to child *c*, and *a'* gave *b'* to child *c'*. The input assignment *g* should be empty. The assignment *i* then can be either $[1 \rightarrow a]$ or $[1 \rightarrow a']$. Hence *f* can be either $[1 \rightarrow a, 2 \rightarrow b]$ or $[1 \rightarrow a', 2 \rightarrow b']$, and $x[f]$ can be either *a* $[1 \rightarrow a, 2 \rightarrow b]$ or *a'* $[1 \rightarrow a', 2 \rightarrow b']$. Consequently, the only value for *k* is

$$k = [4 \rightarrow \{a[1 \rightarrow a, 2 \rightarrow b], a'[1 \rightarrow a', 2 \rightarrow b']\}],$$

where the indices 1 and 2 are subordinated to 4. The matrix is interpreted distributively with respect to this index 4. In our example, the output assignment *h* will end up having the following value:

$$h = [4 \rightarrow \{a[1 \rightarrow a, 2 \rightarrow b, 3 \rightarrow c], a'[1 \rightarrow a', 2 \rightarrow b', 3 \rightarrow c']\}]$$

Notice that this allows us to refer to the guests that brought a present (by h_4 or h^1), to the presents brought by these guests (by h^2), and to the children that got presents that were brought by these guests (by h^3). Also, we can account for correspondence interpretations of sentences like *They liked them*, as *h* records the information about which child got which present.

Let me come back to the case of *most*. It is interesting to notice that the two entities that are always introduced for good by a nominal quantifier, namely, one for the restrictor and one for the intersection of restrictor and matrix, are exactly the ones that we need to define the truth conditions for a quantificational determiner like *most*. In fact, for all *CONSERVATIVE* determiner we just have to refer to the restrictor set and the intersection set, where conservativity is defined as follows:

- (58) A quantificational determiner D is CONSERVATIVE iff $D(X, Y) \Leftrightarrow D(X, X \cap Y)$, where X is the restrictor argument, and Y is the matrix argument.

Now, conservativity is the one general property that holds for all quantificational determiners in natural language (cf. Barwise and Cooper 1981). Hence it is no accident that exactly those two entities are introduced into discourse that are also necessary to establish whether the determiner relation holds (cf. also Peters, Gawron and Nerbonne 1991 for this observation). This link between the truth-conditional semantics and the anaphoric potential that comes with a nominal quantifier leads to an explanatory advantage over the alternative account in DRT that employs modal subordination. In theories of modal subordination, it is generally assumed that two boxes are made available for future reference, namely the restrictor box and the combination of the restrictor box and the matrix box (cf. the discussion in Section 2 and Sells 1985, Roberts 1987, Kamp and Reyle 1993). But the fact that it is exactly these two boxes that can be re-used, rather than, say, the matrix box, is not linked to any other feature of quantifiers.

8. CONCLUSION

Let us come to a conclusion. I have tried to develop in this article an alternative to the DRT treatment of various phenomena involving distributivity, quantification, and plural reference within a non-representational framework. The crucial innovation was a more complex notion of variable assignment which made use of parametrized individuals. In a way, the recursive structure of DRSs is mirrored in the recursive structure of variable assignments with parametrized individuals.

One interesting result is that it is indeed possible to give a non-representational account of these highly complex phenomena, that is, an account that solely relies on the features of the data structure of accessible entities, data structures that are modified in the process of semantic interpretation. An interesting question that was raised in Section 1 is whether the resulting theory is more restrictive, and in particular, whether it is restrictive enough. It is quite obvious that it is more restrictive. For example, a sentence that is to be interpreted can never access the descriptive part of the preceding discourse (in DRT terms, the conditions), but only the discourse entities. However, it may very well turn out that the current framework is not restrictive enough. For example, it would be easy to define an outlandish type of pronoun that can only pick up discourse

entities that are subordinated two steps down (e.g., they could pick up 3 with respect to $[1 \rightarrow a[2 \rightarrow b[3 \rightarrow c[4 \rightarrow d, \dots]]]]$, but not 1, 2 or 4). Pronouns like that obviously do not exist. Hence the data structure presented here allows for more options than what we find in natural language.

There are at least two previous accounts of plurals and plural anaphora within dynamic interpretation that should be mentioned here, although I will not go into a detailed comparison, for reasons of space. Van den Berg (1990) proposes a representation framework that treats sentences as relations between sets of assignment functions. This allows for a treatment of dependency relations between individuals, but it is unclear how to account for cumulative readings. Elworthy (1995) develops an analysis in which the anaphoric potential is captured by "discourse sets", that is, sets of tuples that record anaphoric potential give detailed information about how basic predicates and relations apply to entities. For example, a sentence like (i) *Every student₁ wrote an article₂*, in a situation in which student s wrote article a , and student s' wrote article a' , would generate the discourse set $\{\langle s, a \rangle, \langle s', a' \rangle\}$, where the first member of each tuple is related to index 1, and the second to the index 2. A pronoun like *they₁* can refer to the sum of individuals in the slot indicated by its index, here $\{s, s'\}$. It is also possible to handle correspondence readings, as in the continuation of (i), *They₁ sent them₂ to L&P*, as this type of dependency is encoded in discourse sets. One important difference between this framework and the one developed here is that it lacks any notion of subordination of one discourse referent under another one; every discourse referent (= slot in a tuple) is equally accessible. That (i) cannot be continued by *John read it₂* is then explained by the fact that a singular pronoun cannot refer to a set or sum individual consisting of two entities, $\{a, a'\}$. Elworthy argues that this "referential" analysis, in which number restrictions are essentially determined by what an anaphor refers to, has advantages over a "structural" account, that is, an account that uses discourse referents. However, this seems problematic for texts like *Every student read an article. They liked them. When they talked about them, it turned out that it was the same article*. Here, the plural form *them* is possible although it refers to only one article. Also, it remains unclear how complex texts with independent quantifications should be represented by discourse sets, such as *Every student₁ wrote an article₂, and every professor₃ read a book₄*. Presumably it would be treated like this: Assume that there are two students s, s' and three professors p, p', p'' , that s and s' wrote the articles a and a' , respectively, and p, p' and p'' read the books b, b' and b'' , respectively. The first sentence should generate the discourse set $\{\langle s, a \rangle, \langle s', a' \rangle\}$, the second then should extend this to some set

$\{\langle \dots, p, b \rangle, \langle \dots, p', b' \rangle, \langle \dots, p'', b'' \rangle\}$. It is unclear what the unspecified slots " $\dots$ " should contain. Perhaps we can make use of the element $\perp$ that Elworthy uses to indicate the absence of an individual, and assume the DS $\{\langle s, a, p, b \rangle, \langle s', a', p', b' \rangle, \langle \perp, \perp, p'', b'' \rangle\}$. But this is problematic on several counts: It leads to an inflation of the representational structure (here, the elements $\perp$), and more importantly, it introduces unwarranted dependencies between individuals, e.g. between s and p , as they happen to be instantiated in the same tuple. A framework like the one developed here, with subordination of discourse referents, does not run into such problems.

There are several issues that should be explored further. One is whether a more articulate version of dynamic interpretation can deal with modal subordination phenomena. We have seen that modal subordination can be described quite well in cases where it arises by nominal quantification. It would be interesting to see whether this can be extended to modal subordination in adverbial quantification. Clearly we would have to work with parametrized situation individuals in such cases. Another desideratum is the development of a descriptive language for variable assignments and change of variable assignments that is perspicuous enough for the working linguist, and for which an inference system can be defined.

One crucial difference between DRT and the framework presented here is that anaphoric relations are treated by making use of features of the semantic objects (the variable assignments) instead of making use of features of the representation or DESCRIPTION of these objects. This raises the issue of how far we can go into that direction. We certainly will need descriptions for metalinguistic uses, as in *the man that you have called an idiot*. Perhaps a more essential use of representations is made in the theory of presupposition projection and accommodation of van der Sandt (1992), which assumes that presuppositions are representations that can attach to various constituents within a larger representation. It is an open issue whether the insights of this theory can be rephrased within a framework of direct dynamic representation.

REFERENCES

- Asher, Nicholas: 1993, *Reference to Abstract Objects in Discourse*, Kluwer Academic Publishers, Dordrecht.
- Barwise, Jon and Robin Cooper: 1981, 'Generalized Quantifiers in Natural Language', *Linguistics and Philosophy* 4, 159–219.
- Barwise, Jon: 1987, 'Noun Phrases, Generalized Quantifiers, and Anaphora', in P. Gärdenfors (ed.), *Generalized Quantifiers: Logical and Linguistic Approaches*, D. Reidel, Dordrecht, pp. 233–268.
- Elworthy, David A. H.: 1995, 'A Theory of Anaphoric Information', *Linguistics and Philosophy* 18, 297–332.
- Evans, Gareth: 1980, 'Pronouns', *Linguistic Inquiry* 11, 337–362.
- Gillon, Brendan: 1987, 'The Readings of Plural Noun Phrases in English', *Linguistics and Philosophy* 10, 199–219.
- Groenendijk, Jeroen and Martin Stokhof: 1990, 'Dynamic Montague Grammar', in L. Kálman and I. Polos (eds.), *Proceedings of the Second Symposium on Logic and Linguistics*, Akadémiai Kiadó, Budapest.
- Groenendijk, Jeroen and Martin Stokhof: 1991, 'Dynamic Predicate Logic', *Linguistics and Philosophy* 14, 39–100.
- Heim, Irene: 1982, 'The Semantics of Definite and Indefinite Noun Phrases in English', Ph.D. dissertation, University of Massachusetts at Amherst. Published 1988, New York, Garland Press.
- Heim, Irene: 1983a, 'File Change Semantics and the Familiarity Theory of Definiteness', in R. Bäuerle, Chr. Schwarze and A. von Stechow (eds.), *Meaning, Use and the Interpretation of Language*, Walter de Gruyter, Berlin, New York, pp. 164–190.
- Heim, Irene: 1983b, 'On the Projection Problem for Presuppositions', *Proceedings of the West Coast Conference of Formal Linguistics* 11, 114–126.
- Heim, Irene: 1990, 'E-type Pronouns and Donkey Anaphora', *Linguistics and Philosophy* 13, 137–178.
- Kamp, Hans: 1981, 'A Theory of Truth and Semantic Representation', in J. Groenendijk, T. Janssen and M. Stokhof (eds.), *Formal Methods in the Study of Language*, Mathematical Centre Tract 135, Amsterdam. Also in J. Groenendijk, T. Janssen and M. Stokhof (eds.) (1984), *Truth, Representation and Information*, Foris, Dordrecht, pp. 277–322.
- Kamp, Hans and Uwe Reyle: 1993, *From Discourse to Logic. Introduction to Modeltheoretic Semantics of Natural Language, Formal Logic and Discourse Representation Theory*, Kluwer Academic Publishers, Dordrecht.
- Krifka, Manfred: 1992, 'Definite NPs aren't Quantifiers', *Linguistic Inquiry* 23, 156–163.
- Kadmon, Nirit: 1990, 'Uniqueness', *Linguistics and Philosophy* 13, 273–324.
- Karttunen, Lauri: 1974, 'Presuppositions and Linguistic Context', *Theoretical Linguistics* 1, 181–194.
- Karttunen, Lauri: 1976, 'Discourse Referents', in J. McCawley (ed.), *Syntax and Semantics 7: Notes from the Linguistic Underground*, Academic Press, New York, pp. 363–385.
- Lasnik, Peter: 1989, 'On the Readings of Plural Noun Phrases', *Linguistic Inquiry* 9, 177–197.
- Link, Godehard: 1983, 'The Logical Analysis of Plurals and Mass Terms: A Lattice-Theoretical Approach', in R. Bäuerle, Chr. Schwarze, A. von Stechow (eds.), *Meaning, Use and the Interpretation of Language*, de Gruyter, Berlin, New York.
- Link, Godehard: 1984, 'Plural', Published 1992 in D. Wunderlich and A. von Stechow (eds.), *Semantik. Ein Internationales Handbuch der zeitgenössischen Forschung*, De Gruyter, Berlin, New York.
- Link, Godehard: 1987, 'Generalized Quantifiers and Plurals', in P. Gärdenfors (ed.), *Generalized Quantifiers: Logical and Linguistic Approaches*, Kluwer Academic Publishers, Dordrecht.
- Moxey, L. M. and A. J. Sanford: 1987, 'Quantifiers and Focus', *Journal of Semantics* 5, 189–206.
- Muskens, Reinhard: 1993, 'A Compositional Discourse Representation Theory', *Proceedings of the 9th Amsterdam Colloquium*, pp. 467–486.
- Partee, Barbara: 1984, 'Nominal and Temporal Anaphora', *Linguistics and Philosophy* 7, 243–286.
- Peters, Stanley, Mark Gawron and John Nerbonne: 1991, 'The Absorption Principle and E-type Anaphora', in J. Barwise, J. M. Gawron, G. Plotkin and S. Tutiya (eds.), *Anaphora and Quantification in Situation Semantics*, CSLI Publications, Stanford, pp. 335–362.

- Roberts, Craige: 1987, *Modal Subordination, Anaphora and Distributivity*, Ph.D. Dissertation, University of Massachusetts at Amherst.
- Rooth, Mats: 1987, 'Noun Phrase Interpretation in Montague Grammar, File Change Semantics, and Situation Semantics', in P. Gärdenfors (ed.), *Generalized Quantifiers. Linguistic and Logical Approaches*, Reidel, Dordrecht, pp. 237-268.
- Scha, Remko: 1981, 'Distributive, Collective and Cumulative Quantification', in J. Groenendijk, T. Janssen and M. Stokhof (eds.), *Formal Methods in the Study of Language*, Mathematical Centre Tract 135, Amsterdam, Reprinted in J. Groenendijk, T. Janssen and M. Stokhof (eds.) (1984), *Truth, Representation and Information*, Foris, Dordrecht.
- Sells, Peter: 1985, 'Restrictive and Non-Restrictive Modification', *CSLI Report* 85-28.
- Stalnaker, Robert: 1973, 'Presuppositions', *Journal of Philosophical Logic* 2, 447-457.
- Stalnaker, Robert: 1974, 'Pragmatic Presupposition', in M. K. Munitz and P. K. Unger (eds.), *Semantics and Philosophy*, New York University Press, pp. 197-213.
- Stalnaker, Robert: 1978, 'Assertion', in P. Cole (ed.), *Syntax and Semantics 9: Pragmatics*, Academic Press, New York, pp. 315-332.
- Van den Berg, Martin: 1990, 'A Dynamic Predicate Logic for Plurals', in M. Stokhof and L. Torenvliet (eds.), *Proceedings of the Seventh Amsterdam Colloquium*, ITLI, Amsterdam, pp. 29-51.
- Van der Sandt, Rob: 1992, 'Presupposition Projection as Anaphora Resolution', *Journal of Semantics* 9, 333-377.
- Verkuyl, Henk: 1993, *A Theory of Aspectuality. The Interaction between Temporal and Atemporal Structure*, Cambridge University Press, Cambridge, 1993.
- Zeevat, Henk: 1989, 'A Compositional Approach to Discourse Representation Theory', *Linguistics and Philosophy* 12, 95-131.

Department of Linguistics
University of Texas at Austin
Austin, TX 78712
U.S.A.