

Introduction

1. Language and meaning

Nothing is as easily overlooked, or as easily forgotten, as the most obvious truths. The tenet that language is a tool for expressing meaning is a case in point. Nobody would deny it — but many influential schools and trends in modern linguistics have ignored it, and have based their work on entirely different and often incompatible assumptions.

In many cases, the conviction that meaning can be more or less ignored in the study of language is clearly linked with a conviction that semantics is an independent field, which can be left to those who happen to be interested in meaning, while other linguists can devote themselves to something else — in particular, to syntax. Grammar in general, and syntax in particular, is seen as more or less autonomous of semantics, and can be pursued independently.

Of course this alleged autonomy of grammar has often been questioned, and a great deal of valuable work has been done on the basis of the opposite assumption. Nonetheless, the idea that grammar is one field and semantics another, is still very widespread. In fact, even those linguists who oppose the idea of ‘autonomous syntax’ as a peculiar aberration in the study of language often oppose to it the slogan of ‘syntax AND semantics’ — implying that syntax and semantics are two separate, if interrelated, domains.

The view on which this book is based is quite different. If semantics is to be defined as a study of meaning encoded in natural language then syntax is simply one part of semantics.

Language is an integrated system, where everything ‘conspires’ to convey meaning — words, grammatical constructions, and illocutionary devices (including intonation). Accordingly, linguistics falls naturally into three parts, which could be called lexical semantics, grammatical semantics, and illocutionary semantics. A Morrisian division of the study of signs into

semantics, syntax and pragmatics may make good sense with respect to some artificial sign systems but it makes no sense with respect to natural language, whose syntactic and morphological devices (as well as illocutionary devices) are themselves carriers of meaning.

In natural language, meaning cannot be defined in terms of a relationship between linguistic units and elements of extra-linguistic reality. Attempts to develop semantic theories based on 'truth conditions', 'denotational conditions' and the like have not been fruitful in terms of actual description of languages. They have usually ended up as manifestos and declarations. (Cf. Wierzbicka 1985c.) In natural language meaning consists in human interpretation of the world. It is subjective, it is anthropocentric, it reflects predominant cultural concerns and culture-specific modes of social interaction as much as any objective features of the world 'as such'. 'Pragmatic meanings' are inextricably intertwined in natural languages with meanings based on 'denotational conditions'. (Cf. e.g. Wierzbicka 1980b and 1987b; see also Padučeva 1985.)

Even concrete concepts such as 'mouse', 'rat' or 'worm' are culture-specific and determined in their content by the speakers' interests and attitudes as much as by any objective 'discontinuities in the world'. "*Men make sorts of things*", as Locke (1690/1959, 2:85, original emphasis) put it a long time ago. The idea that abstract concepts such as 'promise' and 'command' (cf. Searle 1979:ix) or 'shame' and 'disgust' (cf. Blount 1984:130) are language-independent and determined by objective 'discontinuities in the world' is even more of an illusion. (For discussion, see Wierzbicka 1985c, 1985a, and 1986c).

But if no boundary can be drawn between 'denotational meanings' and 'pragmatic meanings' in the area of the lexicon, no boundary can be drawn between them in the area of grammar either. The differences between active and passive sentences, between subjects and objects, between direct objects and oblique objects (complements), and so on, are essentially 'pragmatic' — that is, they are determined, to a considerable extent, by the speakers' interests and attitudes. (For detailed discussion, see Wierzbicka 1980b). The thematic-rhematic structure of sentences, too, can be said to be a matter of their 'pragmatic' organization, but it cannot be separated from their lexical and grammatical structure. (For an illuminating discussion of this problem, see Bogusławski 1977.)

Semantics is one. It encompasses lexicon, grammar and illocutionary structure. It is of paramount importance that we be able to perceive its

essential unity, and that whatever part of the overall task we focus on at a given time, we always keep in mind our main goal: an integrated semantic description of natural languages. (For further discussion, see Wierzbicka 1987a.)

2. Grammatical semantics

The present book is devoted to ‘grammatical semantics’, that is to the study of meaning conveyed by grammatical devices — mainly in English, but also in a number of other languages, such as Russian, French, German, Polish, Italian, Spanish and Japanese. ‘Grammatical semantics’ can be naturally subdivided into the semantics of syntax and the semantics of morphology, and that is how the book is organized. Syntax and morphology of course cannot be rigidly separated from one another, but the distinction offers a useful, if somewhat arbitrary, way of organizing material which ranges over a wide variety of grammatical structures.

The basic idea behind the notion of ‘grammatical semantics’ is this. Every grammatical construction encodes a certain meaning, which can be revealed and rigorously stated, so that the meanings of different constructions can be compared in a precise and illuminating fashion, both within one language and across language boundaries.

Grammar is not semantically arbitrary. On the contrary, grammatical distinctions are motivated (in the synchronic sense) by semantic distinctions; every grammatical construction is a vehicle of a certain semantic structure; and this is its *raison d’être*, and the criterion determining its range of use.

For example, if English has a number of different complement constructions, associated with complementizers such as THAT, ING, TO and FOR TO, the choice between these complement constructions is neither arbitrary nor determined by some formal, non-semantic constraints, but is predictable from the intended meaning. In some situations, the speaker can choose between two or more complement constructions, for example:

- a. Mary started TO work.
- b. Mary started workING.
- a. It is time TO go.
- b. It is time FOR me (us, you) TO go.
- a. I hope THAT I will be able to come.
- b. I hope TO be able to come.