

Concept Narrowing: The Role of Context-independent Information

PAULA RUBIO-FERNÁNDEZ
Princeton University

Abstract

The present study aims to investigate the extent to which the process of lexical interpretation is context dependent. It has been uncontroversially agreed in psycholinguistics that interpretation is always affected by sentential context. The major debate in lexical processing research has revolved around the question of whether initial semantic activation is context sensitive or rather exhaustive, that is, whether the effect of context occurs before or only after the information associated to a concept has been accessed from the mental lexicon. However, within post-lexical access processes, the question of whether the selection of a word's meaning components is guided exclusively by contextual relevance, or whether certain meaning components might be selected context independently, has not been such an important focus of research. I have investigated this question in the two experiments reported in this paper and, moreover, have analysed the role that context-independent information in concepts might play in word interpretation. This analysis differs from previous studies on lexical processing in that it places experimental work in the context of a theoretical model of lexical pragmatics.

1 CORE FEATURES

Looking at the psycholinguistic literature on lexical processing (see below), it is possible to argue that the properties associated with a given concept all receive an initial boost on word recognition but that only some subset of these remain detectable over the time course of language processing. I am interested in those meaning components which remain active throughout the interpretation process regardless of their contextual relevance (what I call *core features* of a concept). On a strong interpretation of the term, core features would constitute something corresponding to a core semantic interpretation for the word, which would be constant across contexts. However, on a weaker reading, core features would merely be highly accessible during interpretation, not necessarily computed as part of the word's interpretation—whether they are or not depending on the actual context. In order to determine whether there are core features of a concept, a contrast needs to be

established with other associated properties which do not maintain their initial activation context independently (what I call *non-core features*). The lexical priming study reported in this paper will therefore try to establish (i) whether those conceptual properties which are more strongly associated to a concept might behave as core features by maintaining their initial activation across contexts and (ii) whether those properties which are more weakly associated to a concept might behave as non-core features and prolong their automatic activation only in those contexts where they are relevant for interpretation.¹

The distinction between core and non-core features is related to, but not exactly the same as, that drawn by Barsalou between context-independent and context-dependent properties (Barsalou & Bower 1980; Barsalou 1982). According to Barsalou, context-independent properties are activated by the word for a concept on all occasions, regardless of contextual relevance (e.g. BASKETBALL—ROUND). Context-dependent properties, on the other hand, are only activated by particular contexts in which the word appears (e.g. BASKETBALL—FLOATS). In Barsalou's view, context-independent properties constitute the core meaning components of words, whereas the activation of context-dependent properties accounts for semantic flexibility.

Although Barsalou (1982) offers some empirical evidence for context-dependent and context-independent information in concepts (see below), he does not take into account the question of pre- v. post-lexical access processes. In this respect, it is not clear whether the semantic activation he refers to is the result of an early automatic process of spreading activation of associates or rather reflects a later meaning selection. This question is particularly important in the case of context-dependent properties. According to Barsalou, some context-dependent properties are *ad hoc* or computed on the fly (Barsalou 1983), whereas others are part of the information associated to a concept in long-term memory (what Barsalou calls *conceptual frames*; Barsalou 1992). Context-dependent properties therefore include two different types of conceptual properties in this account: those which have to be inferred during processing and those which are accessed automatically by virtue of their association to the concept. The different patterns of activation of these two types of properties make it difficult to make generalizations about the accessibility of context-dependent properties during processing. In contrast, the notion of non-core features refers only to properties

¹ The distinction I make between strong and weak associates refers to the different strength of association between a concept and its properties. In contrast, the distinction between core and non-core features is based on the different time course of activation that the properties of a concept might have during processing by virtue of the strength of their association.

associated to a given concept in long-term memory, although this association would not be strong enough for them to maintain their initial activation context independently, the way core features would.

2 PREVIOUS RESEARCH POINTING TO A CORE/ NON-CORE DISTINCTION

Using a property verification task, Barsalou (1982) showed that context-independent properties (e.g. HAS A SMELL for SKUNK) were primed in neutral sentence contexts ('The skunk was under a large willow') just as much as they were in biasing contexts ('The skunk stunk up the entire neighbourhood'). However, the verification of context-dependent properties (CAN BE WALKED UPON for ROOF) was faster after a biasing sentence ('The roof cracked under the weight of the repairman') than after a neutral sentence ('The roof had been renovated prior to the rainy season'). Using a similarity-judgment task, Barsalou (1982) also showed that the similarity ratings for pairs of words (e.g. SOFA-DESK or RACCOON-SNAKE) were positively affected by the prior presentation of the category name (FURNITURE and CAN BE A PET for the examples above) only when the property shared by the two concepts was context dependent. This would confirm the prediction that the properties shared by instances of a common category (FURNITURE) are usually context independent, whereas the properties shared by instances of an *ad hoc* category (CAN BE A PET) might often be context dependent.

Barsalou acknowledges that the results from these experiments provide only a functional account of property availability, since the procedure did not address the time course of activation. Likewise, the results reported in Conrad (1978), Tabossi & Johnson-Laird (1980) and Tabossi (1982) provide empirical evidence for some distinction between context-dependent and context-independent properties without taking into account the time course of semantic processing (see references for details).

Whitney *et al.* (1985) and Greenspan (1986) carried out different online studies of word interpretation, which converged on the same basic conclusions. Relative to controls, both high- and low-dominant properties of the prime concepts² were activated at 0 ms, regardless of

² In the word processing literature, high-dominant properties of a concept are understood as those properties which are frequently processed together with the word for the concept. This results in a strong association between the concept and the property. Likewise, low-dominant properties are not so often processed together with the word for the concept and so are more weakly associated to the concept in long-term memory.

which property was emphasized by the context. However, low-dominant properties were no longer activated 300 ms after the occurrence of the noun if the context biased interpretation towards a high-dominant property (e.g. 'The fresh meat was protected by the *ice*'—SLIPPERY). In contrast, high-dominant properties remained active at 1000 ms even in contexts emphasizing a low-dominant property (e.g. 'Robert fell on the *ice*'—COLD).

It is therefore reasonable to conclude that previous off- and online studies provide support for the notion of context-independent information in concepts. The present study of core features using the cross-modal lexical priming paradigm was designed as a small-scale follow-up of previous online studies such as those of Whitney *et al.* (1985) and Greenspan (1986). However, unlike previous studies of word processing, I will provide an interpretation of the results in accordance with the theoretical model of lexical pragmatics currently being developed within relevance theory, which is outlined in the next section. The analysis of experimental work from a cognitive-pragmatic perspective takes forward psycholinguistic analyses by adding a necessary communicative dimension to the interpretation of the language processing data.

3 RELEVANCE THEORY: CONCEPT NARROWING

At the sentence level, pragmatics tries to account for the divergence between the encoded meaning of a complex linguistic expression and the meaning that it is used to communicate in context. Likewise, at the word level, lexical pragmatics investigates how the concept communicated by use of a word may differ from the concept encoded by that word. In relevance theory (Sperber & Wilson 1986/1995; Carston 2002), content words encode concepts. What this means is that a content word gives access to a stable concept in the conceptual repertoire of the hearer. As psychological objects, concepts consist of an address in memory, under which various types of information are stored (i.e. logical, encyclopaedic and lexical information). When a conceptual address appears in the mental representation of an utterance being processed, the various types of conceptual information stored at that address would have been activated or made accessible (Sperber & Wilson 1986/1995).

One of the cases where the concept communicated by use of a word differs from the concept encoded by that word is concept narrowing, which results from a general pragmatic process of concept adjustment in online utterance interpretation (Carston 2002; Wilson 2003). In

instances of narrowing, a word is used to convey a more specific concept than the one it encodes.³ Consider the following examples:⁴

- (1) a. John said he would like a *drink*.
- b. This Christmas the *bird* was delicious.
- c. Mary is *happy* here.

In the examples under (1), a content word is used to convey a more specific concept than the lexical one. In (1a), it is understood that John would like an alcoholic drink, rather than any type of drink, just as in (1b) *bird* would refer to some edible bird, in particular one of the typical birds served at Christmas dinner (i.e. goose or turkey). Similarly, interpreting (1c) would involve understanding the term *happy*, which covers a wide range of positive emotional states, in a specific way that applies to Mary's. Understanding the examples in (1a), (1b) and (1c) would involve constructing an *ad hoc* concept DRINK*, BIRD* and HAPPY*, respectively, with a more restricted denotation than the concepts linguistically specified.

Relevance Theory seems to share with psycholinguists such as Barsalou (1987) the assumption that an *ad hoc* concept maybe formed from a stable concept by using the information associated with the stable concept to restrict or extend its extension (Carston 2002; Wilson 2003). Some selection of the logical and encyclopaedic information stored under the address of the lexical concept would therefore take place in concept narrowing. According to Carston (1996, 2002), when an *ad hoc* concept results from a process of narrowing, some encyclopaedic property of the lexical concept is promoted to the status of logical or content constitutive. For example, in processing (1a), the *ad hoc* concept DRINK* would result from strengthening the lexical concept DRINK by making the encyclopaedic property ALCOHOLIC content constitutive of the concept communicated.

Unfortunately, relevance theory and other theoretical work in lexical pragmatics do not offer a more detailed explanation of the mechanisms involved in concept narrowing, especially in terms of the time course of

³ The opposite process would be concept loosening, where a word is used to convey a more general concept than the one encoded. Most instances of metaphor interpretation, for example, involve both narrowing and loosening the concept for the metaphor vehicle (see Carston 2002). Metonymy, on the other hand, would involve semantic transfer: a third type of process where the outcome is not a narrower or broader concept but a different concept altogether (Recanati 2004).

⁴ Different versions of these examples are usually cited in the relevance theory literature as instances of narrow and loose use (e.g. Carston 2002). However, it is possible that these uses might be so common as to have given rise to secondary meanings of the corresponding words—with the corresponding concepts not being *ad hoc* anymore. If that were the case, I would be interested in how online interpretation would have worked *before* these uses were lexicalized.

activation of the relevant properties (e.g. Wilson 2003; Recanati 2004). This is particularly problematic when trying to derive testable predictions from the theory. However, if certain conceptual properties were so strongly associated with a concept as to behave as core features (i.e. become and remain activated in processing the corresponding word regardless of their contextual relevance), then their role in word interpretation would seem particularly interesting in cases of concept narrowing. It seems a reasonable assumption that those encyclopaedic properties of the lexical concept which are promoted to the status of content constitutive in narrowing the encoded concept would remain active during the interpretation process (Rubio-Fernández 2005).⁵ In this view, it is possible that core features might function as some sort of default meaning component of the prime words, with the encoded concepts being always narrowed down to their stereotypical members, provided the result is consistent with the context (Levinson 2000). For example, if YELLOW is a core feature of the lexical concept BANANA, in processing the sentence 'The tourist gave the monkey a banana', the concept BANANA* might be interpreted as referring to yellow bananas. However, it is also possible that core features are only upgraded to the status of definitional of the concept expressed in particular contexts where they are relevant for interpretation and can reasonably be taken to have been intended by the speaker. In all other contexts, core features would still be highly accessible during interpretation, although not actually computed as part of the information communicated by use of the corresponding word. These two alternatives will be further discussed in the light of the experimental results.

4 THE SCOPE OF THE MECHANISM OF SUPPRESSION

Another question that is worth discussing around the notion of core features is the scope of the mechanism of suppression in word processing, that is, the scope of the cognitive mechanism which reduces the activation of lexical information that is irrelevant for word interpretation. According to some views (e.g. Simpson & Kang 1994; Simpson & Adamopoulos 2002), suppression is a rather specific mechanism only operating when candidates are mutually exclusive (e.g. the various meanings of an ambiguous word or the literal and the

⁵ It also seems reasonable that conceptual properties that might be relevant in a certain context remain active during processing without necessarily being promoted to the status of content constitutive. For example, in 'Mary travels so often that she cannot even keep a cactus at home', the non-core feature RESILIENT would probably be highly accessible in the context even though the concept CACTUS would not need to be narrowed down to include only resilient cacti.

figurative interpretations of a metaphor vehicle). In these cases, keeping the inappropriate meaning active or simply leaving it to passively decay could interfere with sentence comprehension, hence the need for an executive mechanism that actively reduces the activation of information that is inconsistent with the interpretation of the sentence. However, according to other views (e.g. the Structure Building Framework; Gernsbacher 1990), suppression is a general cognitive mechanism involved not only in disambiguation and metaphor interpretation but also in reducing the activation of less relevant information of unambiguous words like the ones I tested in the experiments reported here.

In another study (see Rubio-Fernández 2006), I found evidence that core features get actively suppressed in contexts where they are inconsistent with the mental representation of the prime concept. For example, FAST, a core feature of the concept CHEETAH would drop in activation between 0 and 400 ms when processing a conflicting sentence like 'In their final exam the biology students had to dissect a cheetah'. In other type of contexts where this property was not inconsistent with interpretation, it was still active at 400 ms. This empirical evidence is consistent with both of the above views of suppression. However, what remains to be seen is whether in a context where high-dominant properties or strong associates of a concept are simply irrelevant for interpretation (and not necessarily inconsistent), they would maintain their initial activation at 400 ms. If this is the case, the very notion of core features would support the narrow scope view of suppression, where competition between candidates is necessary to trigger the mechanism. Core features would therefore become and remain activated regardless of their contextual relevance, unless actively suppressed in a conflicting context.

5 A PRIMING STUDY OF CORE FEATURES

In order to investigate the notion of core features, I carried out two experiments using a cross-modal lexical priming procedure adapted from the literature on disambiguation (Swinney 1979; Tanenhaus *et al.* 1979). This procedure requires that participants make a lexical decision on a series of visual targets while listening to different sentential contexts. Word targets are sometimes related to the sentence (e.g. 'John found a *bug* in his room'—ANT), so facilitation of word recognition relative to an unrelated control is taken as a measure of activation.

Unlike the experiments discussed in the previous section, I was interested not only in investigating the role of context in semantic activation but also in finding some empirical basis for distinguishing core and non-core features among the associates of a given concept.

The activation of strong and weak associates of the prime concepts⁶ was tested across three time delays (i.e. 0, 400 and 1000 ms) in two different types of contexts (i.e. neutral and weak-associate biasing).

If strong associates behave like core features of the prime concepts, they should become and remain activated both in neutral and weak-associate biasing contexts where they are irrelevant for interpretation. In contrast, if weak associates behave like non-core features of the prime concepts, then they should maintain their initial activation only in weak-associate biasing contexts where they are relevant for interpretation.

In terms of concept narrowing, if strong associates behave like core features and remain active across contexts regardless of their contextual relevance, the important issue to discuss is whether they play the same role in lexical interpretation in both types of contexts or rather have a different function in concept comprehension depending on their contextual relevance.

5.1 Experiment 1

In the first experiment, I tested the time course of activation of strong and weak associates in neutral contexts, which do not make salient any particular aspect of the primes (e.g. ‘The tourist gave the monkey a *banana*’—YELLOW [SA]/BENT [WA]). I worked on the assumption that the interpretation of the prime words in this type of context would be rather stereotypical. According to even the most expansive views on the role of suppression (e.g. Gernsbacher 1990), the mechanism of suppression would not have sufficient grounds to operate in neutral contexts as it would not have enough of a basis on which to select the conceptual information that maybe irrelevant for interpretation. In this respect, if associates decreased in activation during the processing of neutral contexts, it would be due to passive decay for lack of contextual stimulation rather than to active suppression.

As core features of the prime concepts, I would expect strong associates to maintain their initial activation in neutral contexts past the automatic phase of spreading activation. In contrast, as non-core features, weak associates should lose their activation soon after the spreading activation phase as they would only prolong their initial activation in contexts where they are relevant for interpretation.

⁶ I will use the terms ‘strong’ and ‘weak associates’ in order to emphasize the association between the prime concepts and the target properties in long-term memory (in contrast, for example, with the *ad hoc* properties that Barsalou (1982) would consider as context dependent). However, this distinction is equivalent to that drawn by other terms used in the literature; for example, ‘high-’ and ‘low-dominant properties’ or ‘central’ and ‘peripheral properties’ (Whitney *et al.* 1985; Greenspan 1986).

5.1.1 Method

Participants The participants in this experiment were 60 undergraduate students at Cambridge University who volunteered to take part in the experiment. They all had English as their first language. Each session lasted approximately 12 minutes.

Materials and design A set of 22 common nouns with predictable distinctive properties were selected as primes. Rather than taking these properties from a dictionary definition of the critical nouns, I undertook a direct assessment of property dominance by using two questionnaires based on the literature on prototypes (Rosch 1973; Rosch & Mervis 1975; Barsalou 1983, 1985, 1987). In Questionnaire 1, participants were presented with two tasks. In the first one, they were asked to give a brief definition for each of the 22 words of a list (example given: 'HOUSE: Building where people live'). In the second task, they were asked to list what they thought were the distinctive characteristics of the concepts in the same list of words (example given: 'How could you distinguish a WHALE from other animals? Lives in the sea, large size, mammal'). In Questionnaire 2, participants were presented with a free-association task where they were asked to write down the first characteristic that came to mind when reading the words on a list, which was the same as in Questionnaire 1.

After piloting the questionnaires on 15 participants (5 for Questionnaire 1 and 10 for Questionnaire 2), only minor modifications needed to be made to the instructions. The final versions of the questionnaires were distributed among 65 participants (25 for Questionnaire 1 and 40 for Questionnaire 2). Having chosen a list of words with predictable highly distinctive properties, the results were as expected apart from two terms, *tip* and *spring*, which are ambiguous and did not show a homogeneous result. These latter were discarded. For each of the 20 remaining concepts, the most frequent distinctive property (what I refer to as *strong associate*) was selected from the three different tasks.

A potential problem in selecting the set of weak associates was the strength of the association between the property and the concept. If prime and target were not related closely enough, it could be argued that the property was not associated to the concept in long-term memory (as seems to be the case with Barsalou's context-dependent properties; see Whitney *et al.* 1985 for a critique). Therefore, weak associates were selected not from the least frequent distinctive properties given in the questionnaires but rather from the range between the third and the seventh most frequent properties (see Appendix A for a complete list of primes and targets).

A neutral sentence was then constructed for each one of the 20 primes, which was always the last word in the sentence (see Appendix B for a complete list of the sentential contexts). In order to avoid intra-lexical priming, sentences did not include any associate of the target word other than the critical prime. The 20 critical sentences were divided into two equal groups matched by the word frequency and the length of the corresponding targets (Johansson & Hofland 1989). One group of sentences were paired with related targets (e.g. 'For the dinner Mary brought *champagne*'—BUBBLE [SA]/CELEBRATION [WA]). For the other group, targets were scrambled, and so the sentences were paired with unrelated target words ('My grandmother knew that old *lullaby*'—SMALL [SA]/WINDOW [WA]). The unrelated sentences and targets served as controls. Two lists of materials were constructed by pairing one group of sentences with related targets in List A and with unrelated targets in List B and the other group of sentences with unrelated targets in List A and with related targets in List B. Another set of 20 neutral sentences was constructed and paired with English-like non-words (e.g. Today John was late for his meeting—WASK). Critical and filler sentences were randomized individually for each participant in each of the two lists of materials.

Given that a word is identified at the point in time when information uniquely specifies it, which may actually occur before the physical ending of the word (Marslen-Wilson 1987), for each of the 20 nouns in the experimental materials, a point was selected where the prime would be unequivocally recognized. Targets were presented visually at the end of the acoustic signal 0, 400 or 1000 ms after the word recognition point was selected for each prime. This enabled accurate measuring of initial semantic activation, while controlling for the possibility that an early contextual effect may result from an early word recognition followed by a fast property selection given the length of the primes. Participants were randomly assigned to one Target Type, List and Inter-Stimulus Interval (ISI), so each participant saw each sentence and the corresponding target only once.

Sentences were recorded at a normal rate by a male speaker on an Apple Macintosh computer. The word recognition point of each prime at the end of the context was marked using a sound-editing program. The auditory stimuli and the visual targets were synchronized using a specialized computer program.

The experimental materials were preceded by two sets of practice trials. The first one consisted of a lexical decision task on a list of 10 words and 10 non-words that were presented visually one at a time and randomized separately for each participant. The second set of practice

trials included both sentential contexts in the acoustic modality and visual targets for a lexical decision. This set of trials contained six neutral sentences similar to the critical ones, although the corresponding visual targets were not related to the primes in any of the practice trials.

Apparatus The experiment was conducted on a Toshiba laptop computer. The sentences were presented through a pair of headphones plugged into the laptop. The visual probes were presented in capital letters in the middle of the computer screen on a white background. Responses to the visual targets were made via a response box connected to the laptop. *Word* responses were made with the thumb or the index finger of the right hand on the right-most key of the response box, which was a green key. *Non-word* responses were made with the thumb or the index finger of the left hand on the left-most key of the response box, which was a red key. Target words remained on the screen until the participant made a lexical decision. There was a 1000-ms delay between the offset of the visual target and the onset of the following acoustic context.

Procedure The experiment was presented to the participants as a simple psycholinguistic experiment investigating language processing. Participants were first given standard written instructions, which were then explained by the experimenter. Participants were told that they would be listening to a series of sentences through the headphones and that at the end of each sentence a string of letters would appear on the computer screen. They should try to indicate as fast and accurately as possible whether the string of letters was a word of English or not by pressing the corresponding key on the response box. It was emphasized that both tasks, namely listening carefully to the sentences and making a fast lexical decision, were equally important, although they should be taken as two unrelated tasks. In order to avoid the possibility that participants might have divided their attention and tried to find some underlying coherent structure connecting the sentences, it was stressed that the sentences were unconnected and did not make a story. It was also indicated that non-words would not correspond in principle with words of other languages, but rather be orthographically similar to legitimate English words.

Participants were tested individually. They ran through the two sets of practice trials with the experimenter and got appropriate feedback on their performance. When being tested on the critical materials, participants were left on their own in a closed room or cubicle.

To make sure that participants paid adequate attention to the priming sentences, a short memory test was given at the end of the

experiment. Participants had been told about this memory test in the instructions. Three randomly chosen sentences from the critical set and another three from the filler set were included in the memory test. Another six sentences similar in style but different from any of the sentences used in the experiment completed the memory test. Participants were instructed to tick the sentences that they thought they had listened to in the experiment. It was stressed that no change had been made to the original sentences.

5.1.2 Results The minimum of correct responses required in the memory test was 2.5 standard deviations (SD) below the participants' average of correct responses per ISI. Only one participant was replaced because he failed to meet this criterion.

The mean response time, SD and proportions of missing data for each Relatedness condition together with the facilitation [i.e. the difference between the experimental (related) and the control (unrelated) conditions] and its significance level per Target Type and ISI are presented in Table 1. Across the two experiments, a response time data point was treated as 'missing' if it was either from an erroneous response or over 2.5 SD above the participant's average response time to the word targets in his exercise.

The statistical analysis of Experiment 1 examined the effects of Target Type (strong/weak associate), Target Relatedness (related/unrelated), ISI (0/400/1000 ms) and List (A/B). Mean reaction times were entered into four-way analyses of variance (ANOVAs), with participants (F_1) as random variables.⁷ There was a significant overall

Target	Relatedness	ISI		
		0	400	1000
Strong associates	Related	601 (119, 0.07)	536 (53, 0.06)	518 (50, 0.03)
	Unrelated	649 (121, 0.00)	585 (58, 0.05)	522 (56, 0.11)
	Facilitation	−48***	−49**	−4
Weak associates	Related	634 (120, 0.03)	645 (85, 0.03)	740 (188, 0.06)
	Unrelated	680 (114, 0.04)	639 (74, 0.06)	763 (177, 0.07)
	Facilitation	−46***	6	−23

Table 1 Mean reaction times (in milliseconds), SD, proportions of missing data and facilitation in each condition in Experiment 1
*** $P < 0.05$, ** $P < 0.01$.

⁷ Because this study was a small-scale follow-up of previous experiments, the design was not powerful enough to carry out reliable analyses per item. Even though this is clearly a limitation of the study, List was included as an independent variable to see whether the distribution of the materials had had any significant effect on the ANOVAs.

priming effect.⁸ Reaction times to targets in the related condition were 27 ms faster than those in the unrelated condition, $F_1(1, 48) = 20.77$, $MSE = 1077.4$, $P < 0.001$. There was also a significant main effect of Target Type, $F_1(1, 48) = 17.26$, $MSE = 22\ 864$, $P < 0.001$. The $ISI \times$ Target Type interaction was also significant, $F_1(2, 48) = 4.750$, $MSE = 22\ 864$, $P < 0.02$. This interaction can be explained as an effect of the different speed of reaction of the different groups of participants tested in each condition, the greatest difference between strong and weak associates being observed at the 1000-ms delay (232 ms slower in the weak associate condition). The Relatedness $\times$ ISI $\times$ Target interaction, which was critical for the investigation, was significant, $F_1(2, 48) = 3.390$, $MSE = 1077.4$, $P < 0.05$. List did not show any significant main effect or interaction with any other variable.

A $2 \times 2 \times 3 \times 2$ (Target $\times$ Relatedness $\times$ ISI $\times$ List) ANOVA was carried out on the arcsine transformation of the missing data using participants as the random factor. The missing data were arcsine transformed to stabilize variances (Winer 1971). A main effect of ISI was observed, $F_1(2, 48) = 5.001$, $MSE = 0.015$, $P < 0.02$, the highest missing rate being observed at the 1000-ms delay (0.193). A significant Target $\times$ ISI $\times$ List interaction was observed, $F_1(2, 48) = 4.008$, $MSE = 0.015$, $P < 0.03$. Since the missing rates for each Target Type did not vary consistently across the List and ISI conditions, this interaction can be understood as an effect of the different accuracy of the different groups of participants tested in each condition. The critical Relatedness $\times$ ISI $\times$ Target interaction was marginally significant, $F_1(2, 48) = 2.488$, $MSE = 0.044$, $P < 0.1$. The average arcsine-transformed missing data were higher in the unrelated than in the related condition across ISIs (0.155 v. 0.140).

5.1.3 Discussion Both strong and weak associates were significantly primed at 0 ms in neutral contexts. This facilitation would have been the result of an automatic process of spreading activation of associates. Therefore, the choice of targets included associates of the prime concepts in both the strong and the weak associate conditions. However, the level of priming for the two types of associates diverged at 400 ms, where only strong associates remained active. Since neutral contexts did not prime any particular aspect of the prime concepts, strong associates must have maintained their initial activation because of

⁸ In both experiments, only significant results will be reported, with the exception of the critical interactions, which will be reported for every ANOVA.

their strong association to the prime concepts and not because they were relevant for interpretation. In this respect, strong associates behaved like core features of the prime concepts in neutral contexts. The loss of activation of weak associates at 400 ms and strong associates at 1000 ms would have been the result of passive decay rather than active suppression since neither of these target types was clearly irrelevant for interpretation in neutral contexts. Most importantly, the overall patterns of activation of strong and weak associates in neutral contexts were significantly different.

It is possible that strong associates would have maintained their initial activation past the 400-ms delay if participants had had a discourse background with which to integrate the incoming information. Because it was emphasized in the instructions that the sentences did not make a story, it is possible that participants may have processed the sentences rather shallowly, without constructing a very detailed mental representation of each sentence. In a different study, I have observed that strong associates remain active up to 1000 ms in contexts where they are highly relevant for interpretation (see Rubio-Fernández 2007).

I take the results of the first experiment as supportive of the core/non-core distinction. Strong associates behaved like core features of the prime concepts since they became and remained active in neutral contexts where they were not particularly relevant for interpretation. In contrast, weak associates lost their initial activation by 400 ms, as predicted for non-core features when they are not relevant for interpretation.

5.2 Experiment 2

In the second experiment, I tested a twofold hypothesis. First, I investigated whether strong associates behave like core features in weak-associate biasing contexts where they are irrelevant for interpretation (e.g. ‘The child said the moon was shaped like a *banana*’—YELLOW). If this was the case, there should be no difference between their pattern of activation in Experiments 1 and 2, with strong associates getting and remaining active in both neutral and weak-associate biasing contexts. Second, I investigated whether weak associates maintain their initial activation in weak-associate biasing contexts where they are relevant for interpretation (e.g. BENT would be the weak associate target for the biasing context above). If they behave like non-core features, weak associates should show a different pattern of activation in the two experiments, since their initial activation decayed by 400 ms in neutral contexts. Overall, the patterns of activation of strong and weak associates should be similar in weak-associate biasing contexts since core features would maintain

their initial activation context independently and non-core features would do so in contexts where they are relevant for interpretation.

5.2.1 Method

Participants The participants in this experiment were 60 undergraduate students at Cambridge University who volunteered to take part in the experiment. They all had English as their first language. Each session lasted approximately 12 minutes.

Materials and design Sentences biasing weak associates were used as irrelevant contexts for strong associates. Therefore, biasing contexts had to (i) make weak associates salient while (ii) making strong associates irrelevant for interpretation. Bearing these criteria in mind, 20 biasing sentences were constructed, one for each of the 20 critical nouns used in Experiment 1. The prime was always the last word in the sentence. In order to avoid intra-lexical priming, sentences did not include any associate of the target word other than the critical prime word. The 20 weak-associate biasing sentences were evaluated in a questionnaire that was distributed among 34 students, who had to rate using a 1–5 scale how close each target was to the point of the sentence (example given: ‘In the sentence “Despite his poor reading skills, the boy was very good at *scrabble*,” WORD would be much more closely related to the point of the sentence than TILE’). The strong and weak associates were randomized and divided into two questionnaires so that each participant got only one type of target per sentence. Not being biased by the sentences, I expected the strong associates to be rated between 1 and 2.5, whereas the ratings for the weak associates should have ranged from 3.5 to 5 given their contextual relevance. In view of the results, minor changes were made to five of the sentences (see Appendix B for a complete list of the sentential contexts).

The same set of primes and targets that were used in Experiment 1, as well as the same experimental design, were used in Experiment 2 to measure the activation of strong and weak associates in weak-associate biasing contexts. As in Experiment 1, the 20 weak-associate biasing sentences were divided into two equal groups matched by the word frequency of the corresponding targets (Johansson & Hofland 1989). One of these groups was paired with related targets (‘When they got the results, John went to buy a bottle of *champagne*’—BUBBLE [SA]/CELEBRATION [WA]) and the other with scrambled and therefore unrelated targets (‘They could still remember when their grandmother taught them that *lullaby*’—TALL [SA]/THIN [WA]). As in the

previous experiment, the unrelated sentences and targets served as controls. The two sentence groups and prime-target pairings were combined into two lists of materials. A set of 20 similar filler sentences was paired with English-like non-words. Critical and filler sentences were randomized individually for each participant in both lists of materials. Participants were randomly assigned to one Target Type, List and ISI.

The recording and setting of the materials were the same as in the previous experiment. The two sets of practice trials were also the same as those used in Experiment 1. Practice sentences were recorded again together with the new materials.

Apparatus and procedure The apparatus and procedure in Experiment 2 were the same as in Experiment 1. The only modification that was made to the standard instructions was to ask participants to try to visualize the meaning of the sentences they were listening to, in the same way that they would visualize the story if they were reading fiction. With this modification, I tried to make sure that participants arrived at the intended interpretation of the sentence, for which weak-associate biasing sentences had to be processed more deeply than neutral sentences (e.g. ‘The only thing that I have been able to grow in my garden is a *cactus*’—DRY [WA]).

The secondary task was the same as in Experiment 1, with the memory test including again six sentences from the experimental materials plus another six sentences similar to the original ones.

5.2.2 Results No participants had to be discarded because of their performance on the memory test. The results of Experiment 2 are presented in Table 2.

The statistical analysis of Experiment 2 examined the effects of Target Type (strong/weak associate), Target Relatedness (related/unrelated),

Target	Relatedness	ISI		
		0	400	1000
Strong associates	Related	625 (112, 0.05)	663 (101, 0.03)	631 (98, 0.03)
	Unrelated	652 (115, 0.04)	741 (172, 0.07)	638 (113, 0.07)
	Facilitation	−27*	−78**	−7
Weak associates	Related	743 (263, 0.02)	590 (80, 0.02)	598 (87, 0.04)
	Unrelated	778 (276, 0.08)	642 (81, 0.09)	630 (108, 0.05)
	Facilitation	−35*	−52***	−32*

Table 2 Mean reaction times (in milliseconds), SD, proportions of missing data and facilitation in each condition in Experiment 2
p* < 0.1, *p* < 0.05, ****p* < 0.01.

ISI (0/400/1000 ms) and List (A/B). Mean reaction times were entered into four-way ANOVAs, with participants (F_1) as random variables. A main effect of Relatedness was observed, $F_1(1, 48) = 27.66$, $MSE = 1599.4$, $P < 0.001$, with reaction times to targets in the related condition being 38 ms faster than those in the unrelated condition. The Relatedness $\times$ ISI interaction was also significant, $F_1(2, 48) = 3.555$, $MSE = 1599.4$, $P < 0.04$. Reaction times to related targets were faster than those to unrelated targets across the three time delays, although the difference varied at each ISI (-31 , -65 and -19 ms, respectively). Unlike in Experiment 1, the critical Relatedness $\times$ ISI $\times$ Target interaction was not significant, $F_1(2, 48) = 1.062$, $MSE = 1599.4$, $P > 0.3$. List did not show any significant main effect or interaction with any other variable.

A $2 \times 2 \times 3 \times 2$ (Target $\times$ Relatedness $\times$ ISI $\times$ List) ANOVA was carried out on the arcsine transformation of the missing data using participants as the random factor. Relatedness showed a significant main effect, $F_1(1, 48) = 6.142$, $MSE = 0.040$, $P < 0.02$, the highest missing rate being observed in the unrelated condition (0.192). Like in the analysis of the reaction-time data, the critical Relatedness $\times$ ISI $\times$ Target interaction was not significant, $F_1(2, 48) = 0.954$, $MSE = 0.040$, $P > 0.3$.

Separate analyses for strong and weak associates across Experiments 1 and 2 examined the effect of context type (neutral/biasing), Target Relatedness (related/unrelated), ISI (0/400/1000 ms) and List (A/B). Mean reaction times were entered into four-way ANOVAs, with participants (F_1) as random variables. Regarding the results for strong associates, main effects of Relatedness, $F_1(1, 48) = 24.11$, $MSE = 1573.9$, $P < 0.001$, and Context, $F_1(1, 48) = 11.96$, $MSE = 20191$, $P < 0.002$, were observed. The faster reaction times were observed in the related condition (-35 ms) and neutral contexts (-90 ms). The Relatedness $\times$ ISI interaction was significant, $F_1(2, 48) = 5.466$, $MSE = 1573.9$, $P < 0.008$, with the fastest reaction times being observed in the related condition across the three time delays, although facilitation varied at each ISI (-38 , -64 and -6 ms, respectively). The critical Relatedness $\times$ ISI $\times$ Context was not significant, $F_1(2, 48) = 0.960$, $MSE = 1573.9$, $P > 0.3$. No main effect or significant interaction was observed with the List variable.

A $2 \times 2 \times 3 \times 2$ (Context $\times$ Relatedness $\times$ ISI $\times$ List) ANOVA was carried out on the arcsine transformation of the missing data for strong associates using participants as the random factor. Only the Relatedness $\times$ ISI interaction was significant, $F_1(2, 48) = 3.404$, $MSE = 0.049$, $P > 0.05$. Since the missing rates were consistently lower in the

related condition, this interaction can be explained as an effect of this coefficient varying across the three ISIs. Like in the analysis of the reaction-time data, the critical Relatedness $\times$ ISI $\times$ Context interaction was not significant, $F_1(2, 48) = 1.059$, $MSE = 0.049$, $P > 0.3$.

In the ANOVAs of the reaction-time data for weak associates, only a significant main effect of Relatedness was observed, $F_1(1, 48) = 24.72$, $MSE = 1102.9$, $P < 0.001$, with reaction times to related targets being 30 ms faster than those to unrelated targets. The critical Relatedness $\times$ ISI $\times$ Context interaction was peripherally significant, $F_1(2, 48) = 2.840$, $MSE = 1102.9$, $P < 0.07$. Again, List did not show any significant main effect or interaction with other variables.

A $2 \times 2 \times 3 \times 2$ (Context $\times$ Relatedness $\times$ ISI $\times$ List) ANOVA was carried out on the arcsine transformation of the missing data for weak associates using participants as the random factor. The only significant result was the main effect of Relatedness, $F_1(1, 48) = 6.258$, $MSE = 0.036$, $P < 0.02$, the average missing data being higher in the unrelated than in the related condition (0.190 v. 0.104). The critical Relatedness $\times$ ISI $\times$ Context interaction did not reach significance level, $F_1(2, 48) = 0.207$, $MSE = 0.036$, $P > 0.8$.

Given that the 400-ms delay is the critical one in determining whether weak associates behave like non-core features, a $2 \times 2 \times 2$ (Relatedness $\times$ Context $\times$ List) ANOVA was carried out on the reaction-time data for weak associates in the 400-ms condition of Experiments 1 and 2. The only significant result was the critical Relatedness $\times$ Context interaction, $F_1(1, 16) = 6.714$, $MSE = 1245.1$, $P < 0.03$.

A $2 \times 2 \times 2$ (Context $\times$ Relatedness $\times$ List) ANOVA was carried out on the arcsine transformation of the missing data for weak associates in the intermediate ISI. The only significant result was the main effect of Relatedness, $F_1(1, 16) = 6.860$, $MSE = 0.029$, $P < 0.02$, the average missing data being higher in the unrelated than in the related condition (0.213 v. 0.071). The critical Relatedness $\times$ ISI $\times$ Context interaction did not reach significance level, $F_1(1, 16) = 0.720$, $MSE = 0.029$, $P > 0.4$.

5.2.3 Discussion As in Experiment 1, both strong and weak associates got initially activated in weak-associate biasing contexts by an automatic process of spreading activation. However, their level of activation was only peripherally significant at 0 ms. These weaker results could be an effect of the greater processing load of the task, given that in Experiment 2 participants were asked to form a mental picture of the content of the sentences they were listening to.

Unlike in Experiment 1, both strong and weak associates were significantly primed at 400 ms. These results would have been expected given that, as core features of the prime concepts, strong associates should maintain their initial activation even in contexts where they are irrelevant for interpretation, while, as non-core features of the primes, weak associates should be sensitive to contextual factors and be facilitated in contexts where they are relevant for interpretation. Weak associates still showed a marginal level of activation at the longest delay, whereas strong associates had decayed at that point.

According to an alternative interpretation of the priming levels observed in Experiment 2, the differences observed between Experiments 1 and 2 would not be due to differences in meaning selection but rather due to a shift in the activation patterns of the primes (maybe as a result of the different contexts used or the visualization requirement). Since priming was only marginally significant at 0 ms in Experiment 2, facilitation would have not only risen later but also decayed later, being still significant at 400 ms for both types of associates. The results of the ANOVAs, however, show that context does have an effect in the different patterns of activation observed.

In Experiment 2, using weak-associate biasing contexts, there was no significant difference in the overall patterns of activation of strong and weak associates, as distinct from the results of Experiment 1, which employed neutral contexts. Comparing the time course of activation of strong and weak associates separately, strong associates showed similar results in neutral and weak-associate biasing contexts, whereas the activation of weak associates was significantly different in the two types of context, especially in the 400-ms condition.

If the differences observed in Experiments 1 and 2 only reflected a shift in activation rather than a process of meaning selection, strong and weak associates should show different patterns of activation in both experiments, not only in Experiment 1. Also, if delayed activation resulted in a different priming pattern for weak associates, the same difference should be observed for strong associates between Experiments 1 and 2. The results of the ANOVAs therefore support the twofold hypothesis that strong associates behave as core features of the prime concepts, whereas weak associates behave as non-core features.

Given that strong associates were actively suppressed by 400 ms in contexts where they were inconsistent with interpretation (Rubio-Fernández 2006), I interpret the present results as supporting the narrow scope view of suppression. In this respect, suppression would not have operated on strong associates in weak-associate biasing contexts where they remained active past the 400-ms threshold. On

grounds of parsimony, their loss of activation at the longest delay should be viewed as comparable to that observed in neutral contexts.

Overall, I interpret these results as supportive of the core/non-core distinction, with strong associates behaving like core features and weak associates like non-core features. The advantage of the cross-modal priming paradigm used in these experiments is that it allows presenting the visual target at the critical point during the sentence (e.g. presenting BUBBLE right before the offset of *champagne*), allowing for an accurate measure of activation (Swinney 1979). However, it must be noted that the results of cross-modal priming studies have been controversial in that they might reflect not only purely automatic processing but also post-lexical processing such as participants' strategies to check for relatedness (Shelton & Martin 1992). Eye-tracking methodologies currently used in priming studies (e.g. Myung *et al.* 2006) maybe less susceptible to strategic processing. Also, eye-tracking techniques offer a continuous online measure of activation, rather than relying on an intermittent profile made up of time points corresponding to different delays between the acoustic prime and the visual target.

6 GENERAL DISCUSSION

The present results are consistent with the distinction between core and non-core features of a concept. Strong associates behaved like core features in that they became and remained active regardless of their contextual relevance. In contrast, weak associates showed a different pattern of activation in neutral and weak-associate biasing contexts, maintaining their initial activation only in the latter type of context where they were relevant for interpretation. Most importantly, there was no significant difference in the activation patterns of strong associates in neutral and weak-associate biasing contexts, whereas weak associates showed different results in the two types of context. The present study therefore reinforces the findings of previous studies of lexical processing (e.g. Barsalou 1982; Whitney *et al.* 1985; Greenspan 1986). Further research using eye-tracking techniques should offer even more accurate results bearing on the core/non-core distinction.

It is important to notice that, in the word processing literature, the notion of context-independent properties is different from classical views in semantic theory: in the former, these properties are taken as the most strongly associated conceptual information, whereas in the latter they are understood as those meaning aspects which remain necessarily (and not only typically) constant across contexts. For example, from a purely semantic perspective, the meaning of DALMATIAN would

include the entailment IS A DOG but not the property SPOTTED. I have investigated the time course of activation of superordinate terms (e.g. DALMATIAN—DOG) in other lexical priming studies (see Rubio-Fernández 2005, 2007) and, like the strong associates tested in this study (e.g. SPOTTED), they behave like core features of the prime concepts. In neutral contexts (the same as those used in Experiment 1), superordinates get activated at 0 ms and remain active up until the 1000-ms delay (Rubio-Fernández 2005). Even though their patterns of activation are not significantly different, superordinates are accessible for longer than strong associates. In metaphoric contexts where superordinates are inconsistent with the figurative interpretation of the prime, they remain active up until 400 ms, being suppressed only between the intermediate and the longest delays. These long patterns of activation suggest that the semantic nature of the superordinate relation might make their conceptual association even stronger than that of strong associates. This could be related, among other things, to discourse processing strategies since a given noun can be referred to later on in discourse by a definite description of the corresponding superordinate category (e.g. 'Mary wants to take her Dalmatian to a dog show. She thinks the dog can win').

Regarding the scope of the mechanism of suppression, the results support the view that suppression is a specific mechanism that operates only on conceptual information that is inconsistent with the mental representation of the prime and so could interfere with the interpretation process (Simpson & Kang 1994). Contrary to the predictions of The Structure Building Model (Gernsbacher 1990), suppression does not operate on information that is less relevant for interpretation, since core features maintained their initial activation not only in neutral contexts but also in weak-associate biasing contexts where they were irrelevant for interpretation. Given the empirical evidence found in a previous study (see Rubio-Fernández 2006), core features would only get actively suppressed in conflicting contexts, where they are not simply irrelevant but actually inconsistent with word interpretation.

In view of the distinction between core and non-core features, it is reasonable to conclude that suppression would only need to operate on core features that were contextually inconsistent. Non-core features, being more weakly associated to the corresponding concept, would passively decay for lack of contextual stimulation unless relevant in the context.

Now I would like to address the interesting question for theoretical lexical pragmatics which is raised by the existence of core features of

concepts: do these properties play a part in the pragmatic process of concept narrowing given that they are highly accessible past the point where context has a selective effect on feature activation?⁹ In concept narrowing, one or more encyclopaedic properties of the encoded concept are upgraded to the status of content constitutive of the resulting *ad hoc* concept (e.g. ALCOHOLIC in narrowing down DRINK to DRINK* including only alcoholic drinks in its extension; Carston 2002; Recanati 2004). Given that the set of core features included highly distinctive properties of the prime concepts (e.g. SPOTTED for DALMATIAN), it is possible that the prime concepts were narrowed down to their stereotypical members in neutral contexts, which did not make salient any particular aspect of the primes (e.g. 'This Christmas John wanted a *Dalmatian*').

The pragmatic process of narrowing down to a stereotype maybe a common one in processing content words in broad contexts, since the speaker would be expected to specify when she is not referring to a stereotypical member of a category, rather than vice versa. For example, *Dalmatians* may not necessarily refer to spotted Dalmatians by default. However, as the stereotypical type of Dalmatians, spotted Dalmatians would be the unmarked subset of the category. It would therefore be more relevant to use the word *Dalmatian* to refer to a stereotypical spotted Dalmatian than to a different, marked narrowing of the concept (e.g. albino Dalmatians). This would explain why the following examples might seem intuitively different in comprehensibility:

- (2) I love Dalmatians, but not the spotted ones.
- (3) I love Dalmatians, but not the albino ones.

Thus, just as *Dalmatians* maybe generally interpreted as spotted Dalmatians in the above examples (hence the pragmatic oddity of (2)), it is possible that the prime concepts may have been narrowed down to a stereotype when interpreting the corresponding prime words in neutral contexts.

Unlike in neutral contexts, non-core features maintained their initial activation up to the 1000-ms delay in biasing contexts where they were relevant for interpretation. It is possible that these associates showed a different pattern of activation in neutral and weak-associate biasing contexts because only in the latter were they upgraded to the

⁹ Note that lexical disambiguation, a highly context-sensitive process, takes place 200–300 ms after the offset of the ambiguous prime (Swinney 1979; Tanenhaus *et al.* 1979; Onifer & Swinney 1981).

status of content constitutive of the resulting narrow concept. However, in this case, the process would not have resulted from the context being broad (as would have been the case with core features in neutral contexts), but rather because non-core features would have been part of the information directly communicated by use of the prime word in weak-associate biasing contexts. For example, in processing the sentence 'Mary wished her old dog had the figure of a *Dalmatian*', the hearer would understand that Mary has in mind a stereotypically slender Dalmatian, and not any odd type of Dalmatian. In this respect, the *ad hoc* concept DALMATIAN* constructed in understanding the above example would have been narrowed down around the non-core feature SLENDER to include only slender Dalmatians in its extension.

Core features also prolonged their initial automatic activation in weak-associate biasing contexts. However, unlike non-core features, these properties were irrelevant for interpretation in this type of context. It is therefore less likely that they would have been computed as communicated by use of the word and upgraded to the status of content constitutive of the resulting *ad hoc* concept. In the previous example, the type of Dalmatian that Mary would have referred to did not need to be spotted as long as it was slender. The core and non-core features SPOTTED and SLENDER would have, therefore, played different roles in constructing the *ad hoc* concept DALMATIAN* necessary to understand the above biasing sentence.

I would suggest that, as strong associates of the prime concepts, the automatic activation of core features would last beyond the point where context-sensitive processes take place, without these properties having been necessarily selected for interpretation. In other words, contrary to what happens with non-core features, the sustained activation of core features cannot necessarily be interpreted as the outcome of a selection process of contextually relevant properties. It seems therefore legitimate to wonder what the role of core features would be when they remain accessible during interpretation but are not taken as part of what is communicated by the speaker by use of the corresponding word. One possibility is that core features are part of the mental representation of the concept in working memory across contexts. Given the perceptual character of this set of core features as distinctive properties of the prime concepts, it is possible that these features may have an important role in the visual or imagistic mental representation of the corresponding concepts (see Barsalou 1999). It would be interesting to investigate in future research the relation between the propositional and the perceptual mental representation of concepts.

7 CONCLUSIONS

The present empirical findings and theoretical discussion of the data should have implications for both psycholinguistic accounts of lexical processing and pragmatic models of word interpretation. In previous experimental studies of word processing (e.g. Barsalou 1982; Whitney *et al.* 1985; Greenspan 1986; Tabossi 1988), the high activation of a conceptual property during sentence processing was taken as evidence of it being part of the interpretation—or even the stable meaning, of the corresponding word. Even models of language comprehension that would predict the accessibility of those encoded meaning components that are most frequent or salient (e.g. The Graded Salience Hypothesis; Giora 1997, 2002) fail to differentiate between the roles that highly accessible meanings might play in interpretation depending on whether their long activation is automatic or results from an active process of meaning selection. However, as I have shown, salient meaning components or core features may not necessarily be computed as part of the conceptual information communicated by use of a word since they remain active in contexts where they are irrelevant for interpretation. In other words, there is an important distinction to be made between activation/accessibility, on the one hand, and selection as part of the intended meaning, on the other. It is through bringing together theoretical lexical pragmatics and empirical work on word processing that this insight has emerged.

According to pragmatic models of lexical interpretation (e.g. Sperber & Wilson 1986/1995, 2002; Recanati 2004), the accessibility of a conceptual property would be determined, other things being equal, by its recency of processing and contextual relevance. That is, the more recently a property has been processed and the more relevant it is in a given context, the more accessible it would be for word interpretation. Given that core features remain active regardless of their contextual relevance, it seems that there is another factor involved in accessibility which needs to be acknowledged: the strength of association between a property or feature and the corresponding concept. Nevertheless, even if contextual relevance and recency of processing might not fully determine property accessibility, the selection of a core feature as part of the intended interpretation of a word would ultimately be determined by considerations of relevance, especially the maximization of cognitive effects (Wilson & Sperber 2004; see Rubio-Fernández 2007).

The distinction between core and non-core features should therefore allow us to gain a better understanding of both the automatic

processes involved in lexical interpretation, as well as the pragmatic process of determining the concept communicated by use of a word.

APPENDIX A: PRIMES AND TARGETS

PRIMES	SA	WA
Cactus	Spike	Dry
Lion	Mane	Roar
Slippers	Comfortable	House
Skyscraper	Tall	Window
Lullaby	Sleep	Children
Dalmatian	Spot	Slender
Mercedes	Expensive	Reliable
Chair	Back	Wood
Champagne	Bubble	Celebration
Breakfast	Morning	Toast
Cheetah	Fast	Predatory
Sapling	Young	Thin
Woodpecker	Noise	Beak
Pacific	Large	Deep
Rugby	Tough	Oval
Steel	Strong	Silver
Minnow	Small	River
Banana	Yellow	Bent
Encyclopaedia	Knowledge	Alphabetical
Norway	Cold	North

APPENDIX B: SENTENTIAL CONTEXTS

Neutral Contexts

Mary bought her mother a cactus
 Of all the characters in the story book, John preferred the lion
 When travelling, Mary always packs her slippers
 Mary's office is in a skyscraper
 My grandmother knew that old lullaby

This year for Christmas John wanted a Dalmatian
My friend works for Mercedes
For the new apartment Mary only had to buy a chair
For the dinner Mary brought champagne
John didn't have enough money to buy breakfast
The expedition approached the territory of the cheetah
At school every student planted a sapling
The children went to watch the woodpecker
The hotel room looked over the Pacific
All boys in the school played rugby
In this country, they make a lot of steel
Yesterday John caught a minnow
The tourist gave the monkey a banana
The man at the door was selling an encyclopaedia
They had been planning to visit Norway

WA Biasing Contexts

The child said the moon was shaped like a banana
The zebras didn't notice the hiding cheetah
Mary has lots of light in her office now that she works in a skyscraper
John felt like having marmalade at breakfast
Although the telegraph pole stayed in place, the wind bent the
sapling
Their black & white picture looked better in a frame made of steel
John's job involved lots of driving, so he bought a Mercedes
Mary wished her old dog had the figure of a Dalmatian
The only thing that I've been able to grow in my garden is a cactus
When he got the results, John went to buy a bottle of champagne
It would be impossible to rescue the cargo if a ship sunk in the Pacific
When crossing the woods, John went swimming and saw some
minnows
After dinner, Mary changed into her slippers
The only bird that could have made this hole is a woodpecker
To make the fire John used a broken chair
The plane encountered some turbulence over Germany and headed
towards Norway
They could still remember when their grandmother taught them that
lullaby
Not far away from the camp they could hear a lion

Mary didn't know how to look up words in an encyclopaedia
With a round ball the kids won't be able to play rugby

Acknowledgements

The study reported in this paper is part of the experiments completed by the author towards her Ph.D. degree at Cambridge University. This research was initially supported by an Arts and Humanities Research Board Postgraduate Award, and currently by a British Academy Postdoctoral Fellowship and a Marie Curie Outgoing International Fellowship (Project 022149—Inferential Processes in Language Interpretation). This paper was written at University College London and reviewed at Princeton University, where the author is currently working on a postdoctoral project. I would like to thank Julie Sedivy, Larry Barsalou and an anonymous reviewer for their valuable comments. Thanks also to John Williams for supervising the setting and analysis of the experiments, as well as to Richard Breheny for many interesting discussions of the theory and the data. Thanks as well to Robyn Carston for her insightful comments on this paper.

PAULA RUBIO-FERNÁNDEZ

Department of Psychology,
Princeton University,
Green Hall,
Princeton NJ 08540,
USA

e-mail: paularubio@hotmail.com

REFERENCES

- Barsalou, L. W. (1982), 'Context-independent and context-dependent information in concepts'. *Memory and Cognition* **10**:82–93.
- Barsalou, L. W. (1983), 'Ad hoc categories'. *Memory and Cognition* **11**:211–27.
- Barsalou, L. W. (1985), 'Ideals, central tendency and frequency of instantiation as determinants of graded structure in categories'. *Journal of Experimental Psychology: Learning, Memory, and Cognition* **11**:629–49.
- Barsalou, L. W. (1987), 'The instability of graded structure: Implications for the nature of concepts'. In U. Neisser (ed.), *Concepts and Conceptual Development: Ecological and Intellectual Factors in Categorisation* (pp. 101–140). Cambridge University Press. Cambridge.
- Barsalou, L. W. (1992), 'Frames, concepts and conceptual fields'. In A. Lehrer and E. Kittay (eds.), *Frames, Fields and Contrasts*. Lawrence Erlbaum Associates. Hillsdale, NJ.
- Barsalou, L. W. (1999), 'Perceptual symbol systems'. *Behavioural and Brain Sciences* **22**:577–660.
- Barsalou, L. W. & Bower, G. (September 1980), 'A priori determinants of

- a concept's highly accessible information'. Paper presented at *The Meeting of the American Psychological Association*, Montreal.
- Carston, R. (1996), 'Enrichment and loosening: complementary processes in deriving the proposition expressed?' *UCL Working Papers in Linguistics* 8:61–88. Reprinted in (1997) *Linguistische Berichte* 8:103–27.
- Carston, R. (2002), *Thoughts and Utterances: The Pragmatics of Explicit Communication*. Blackwell. Oxford.
- Conrad, C. (1978), 'Some factors involved in the recognition of words'. In J. W. Cotton and R. L. Klatzky (eds.), *Semantic Factors in Cognition*. Lawrence Erlbaum Associates. Hillsdale, NJ.
- Gernsbacher, M. A. (1990), *Language Comprehension as Structure Building*. Lawrence Erlbaum Associates. Hillsdale, NJ.
- Giora, R. (1997), 'Understanding figurative and literal language: the graded salience hypothesis'. *Cognitive Linguistics* 7:183–206.
- Giora, R. (2002), 'Literal vs. figurative language: different or equal?' *Journal of Pragmatics* 34:487–506.
- Greenspan, S. L. (1986), 'Semantic flexibility and referential specificity of concrete nouns'. *Journal of Memory and Language* 25:539–57.
- Johansson, S. & Hofland, K. (1989), *Frequency Analysis of English Vocabulary and Grammar Based on the LOB Corpus: Vol. 1, Tag Frequencies and Word Frequencies*. Clarendon Press. Oxford.
- Levinson, S. C. (2000), *Presumptive Meanings*. MIT Press. Cambridge, MA.
- Marslen-Wilson, W. D. (1987), 'Functional parallelism in spoken word-recognition'. *Cognition* 25:71–102.
- Myung, J., Blumstein, S. E., & Sedivy, J. (2006), 'Playing on the typewriter, typing on the piano: manipulation knowledge of objects'. *Cognition* 98:223–43.
- Onifer, W. & Swinney, D. A. (1981), 'Accessing lexical ambiguities during sentence comprehension: effects of frequency of meaning and contextual bias'. *Memory and Cognition* 9:225–6.
- Recanati, F. (2004), *Literal Meaning*. Cambridge University Press. Cambridge.
- Rosch, E. (1973), 'On the internal structure of perceptual and semantic categories'. In T. E. Moore (ed.), *Cognitive Development and the Acquisition of Language*. Academic Press. New York.
- Rosch, E. & Mervis, C. B. (1975), 'Family resemblances: studies in the internal structure of categories'. *Cognitive Psychology* 4:573–605.
- Rubio-Fernández, P. (2005), 'Pragmatic processes and cognitive mechanisms in lexical interpretation: the on-line construction of concepts'. Ph.D. thesis, University of Cambridge.
- Rubio-Fernández, P. (2006), 'Delimiting the power of suppression in lexical processing: the question of below-baseline performance'. *UCL Working Papers in Linguistics* 18:119–37.
- Rubio-Fernández, P. (2007), 'Suppression in metaphor interpretation: differences between meaning selection and meaning construction'. *Journal of Semantics, Special Issue on Processing Meaning* 24:345–71.
- Shelton, J. R. & Martin, R. C. (1992), 'How semantic is automatic semantic priming?' *Journal of Experimental Psychology: Learning Memory, and Cognition* 18:1191–210.
- Simpson, G. B. & Adamopoulos, A. C. (2002), 'Repeated homographs in word and sentence contexts: multiple

- processing of multiple meanings'. In D. S. Gorfein (ed.), *On the Consequences of Meaning Selection*. APA Publications. Washington, DC.
- Simpson, G. B. & Kang, H. (1994), 'Inhibitory processes in the recognition of homograph meanings'. In D. Dagenbach and T. H. Carr (eds.), *Inhibitory Processes in Attention, Memory, and Language*. Academic Press. San Diego, CA.
- Sperber, D. & Wilson, D. (1986/1995), *Relevance: Communication and Cognition*. Blackwell. Oxford.
- Sperber, D. & Wilson, D. (2002), 'Pragmatics, modularity and mind-reading'. *Mind & Language* **17**:3–23.
- Swinney, D. A. (1979), 'Lexical access during sentence comprehension: (re)-consideration of context effects'. *Journal of Verbal Learning and Verbal Behavior* **18**:645–59.
- Tabossi, P. (1982), 'Sentential context and the interpretation of unambiguous words'. *Quarterly Journal of Experimental Psychology* **34A**:79–90.
- Tabossi, P. (1988), 'Effects of context on the immediate interpretation of unambiguous nouns'. *Journal of Experimental Psychology: Learning, Memory and Cognition* **14**:153–62.
- Tabossi, P. & Johnson-Laird, P. N. (1980), 'Linguistic context and the priming of semantic information'. *Quarterly Journal of Experimental Psychology* **32**:595–603.
- Tanenhaus, M. K., Leiman, J. M., & Seidenberg, M. S. (1979), 'Evidence for multiple stages in the processing of ambiguous words in syntactic contexts'. *Journal of Verbal Learning and Verbal Behavior* **18**:427–40.
- Whitney, P., McKay, T., Kellas, G., & Emerson, W. A. (1985), 'Semantic activation of noun concepts in context'. *Journal of Experimental Psychology: Learning, Memory and Cognition* **11**:126–35.
- Wilson, D. (2003), 'Relevance theory and lexical pragmatics'. *Italian Journal of Linguistics/Rivista di Linguistica* **15**:273–91. Special issue on pragmatics and the lexicon.
- Wilson, D. & Sperber, D. (2004), 'Relevance theory'. In G. Ward and L. Horn (eds.), *Handbook of Pragmatics*. Blackwell. Oxford. Longer earlier version published in (2002) *UCL Working Papers in Linguistics* **14**:249–87.
- Winer, B. J. (1971), *Statistical Principles in Experimental Design*. McGraw-Hill. New York.

First version received: 03.08.2006

Second version received: 21.01.2008

Accepted: 05.05.2008