

THE STRUCTURE OF CATEGORIES: TYPICALITY GRADIENTS, PERCEIVED LINGUISTIC FAMILIARITY AND CROSS-LINGUISTIC COMPARISONS

Norman Segalowitz and Diane Poulin-Dubois

Concordia University, Psychology Department
4171 Sherbrooke Street, West, Montréal, Québec, Canada H4B 1R6

Abstract

This study was concerned with factors that underlie typicality gradients. English and French monolinguals living in the same environment rated items for typicality (perceived representativeness) and linguistic familiarity (perceived frequency of name instantiation). Correlational analyses were conducted to explore the role of rated linguistic familiarity as an underlying factor in determining category typicality gradients. The obtained results have implications for the view that the graded structure of categories may sometimes arise from factors other than the feature similarity relations existing between the items involved; in particular there was evidence that, for some categories, cross-language patterns in typicality gradients reflected cross-language differences in linguistic familiarity.

Key words: Concepts, categories, typicality, cross-linguistic.
Mots clés : Concepts, catégories, typicalité, interlinguistique.

People judge the membership of items in semantic categories (e.g., DOG is a member of the category ANIMAL), not in an all-or-none fashion, but in a graded manner. That is, some members of a semantic category are held to be more typical or more representative of that category than others (Rosch, 1977; Rosch & Mervis, 1975; Smith & Medin, 1981). This results in what has been referred to as the graded structure of categories whereby members of a category are ranked in terms of their goodness of membership.

Numerous studies have demonstrated the psychological relevance of typicality gradients to understanding how people deal with semantic information in various situations. For example, representativeness of items within a category has been shown to affect speed of processing in category verification (Mervis, Catlin, & Rosch, 1976), lexical decision time in priming studies with single word stimuli (Rosch, 1975), and fixation time in studies using words in sentence contexts (Carroll & Slowiczek, 1986). Developmentally, it has been shown that typical category instances are integrated into children's natural language concepts before atypical instances (Anglin, 1977; Carson & Abrahamson, 1976; Nelson, 1974; Saltz, Soller, & Sigel, 1972; White, 1982). Processing differences between typical and atypical category instances have also been found with children in such tasks as sentence verification (Keller, 1982), category priming (Duncan & Kellas, 1978), and recall (Bjorklund, Smith, & Ornstein, 1982; Bjorklund, Thompson, & Ornstein, 1983).

A central theoretical question in this area is why some members of a category come to be judged as more typical than others (see Smith & Medin, 1981, and Medin, Wattenmaker, & Hampson, 1987, for general discussion of this issue). Most theories have focused on the idea that an item's degree of category representativeness depends in some way upon similarity relations among features between that item and others in the category. For example, Rosch and her colleagues (e.g., Rosch & Mervis, 1975) proposed that the representativeness of an item depends on how similar the item is to the category as a whole, that is, on the degree to which the item possesses features in common with other members in that category and the degree to which it is distinct from items in contrast categories. Medin (Medin & Schaffer, 1978) proposed that the strength of an item's category membership is determined by how many other specific members of the category it resembles. Posner (e.g., Posner & Keele, 1975) and Homa and his colleagues (Homa, Sterling, & Trepel, 1981), among others, have proposed that the central tendency of feature patterns across all members of a category determines an abstract prototype for the category and in item's representativeness reflects the degree to which it resembles that prototype.

Recently, however, such approaches have come under critical scrutiny. For example, Barsalou (1985, 1987), Murphy (1982), and Murphy and Medin (1985)

discuss important theoretical problems common to such approaches, such as those concerning definitions of features, the nature of correlational structures, and the insufficiency of theories based on principles of attribute matching and similarity. Beyond such theoretical issues, there is now emerging some empirical evidence to suggest that categories do not necessarily reflect feature similarity relations among category members. For example, it appears that the graded structure of categories does not necessarily remain stable across different linguistic and cultural points of view and contexts (Barsalou & Sewell, 1984; Roth & Shoben, 1983). There is also some evidence that graded structure is relatively unstable within as well as across individuals (Barsalou, 1983, 1987; Barsalou & Medin, 1986). In addition, graded structure has been reported for categories for which membership is clearly well-defined -- as in the case of "even number" -- and where the idea of goodness of membership should not be expected to apply (Armstrong, Gleitman, & Gleitman, 1983). As one alternative to the idea that feature similarity relations underlie the way we represent categories, Murphy and Medin (1985) have suggested that it is people's background knowledge or "naïve theories" about the meanings of words that determine graded structure in categories and that such factors may be more important than feature correlations between members and their categories. According to Murphy and Medin, "... a coherent concept is one that we have a good theory about and that fits with our other knowledge" (1985, p. 313).

This view implies that the different kinds of knowledge we have about the world may differentially influence how we think about various categories. Barsalou's results (1985, see his Table 3) illustrate this point. He found, for example, that for some categories goodness of membership was highly correlated with frequency of instantiation (a subjective measure of how often one encounters a given exemplar) but not with a measure of how much an exemplar possesses ideals, that is, "characteristics that exemplars should have if they are to best serve a goal associated with their category" (Barsalou, 1985, p. 630). For other categories the reverse was true and for still others both ideals and frequency of instantiation seemed to underlie goodness of membership judgments (see also Chaplin, John, & Goldberg [1988] on the role of ideals in determining typicality). He also reported that central tendency was a predictor of typicality for common taxonomic categories whereas it was not for goal derived categories (such as *clothes to wear in the snow*). Such results may reflect the fact that our knowledge about category members derives from a variety of sources. For example, we may come to know about a given category primarily through frequent and direct, perceptual observation of its members. In other cases we may come to know about a category only through verbal reference to its members. In yet other cases we may be formally instructed about category membership through reference to rigorously defined rules. In many cases we can expect

several such sources of knowledge to be involved simultaneously but to varying degrees.

The study of cross-cultural and/or cross-linguistic difference in the graded structure of categories has potentially interesting implications for these issues, particularly with respect to the role played by cultural-linguistic factors as opposed to similarity relations of the physical and functional features of the exemplars themselves. Unfortunately, relatively little research has been directed to the study of factors underlying intergroup similarities and differences in typicality. In one study, Barsalou and Sewell (1984) asked U.S. subjects to rank exemplars of a category as they thought an average U.S., African, Chinese or French person would rank them. They found that when subjects adopted such different points of view they generated different graded structures for the same categories. Hampton and Gardiner (1983) compared British subjects' ratings of typicality to those of Rosch's (1975) U.S. subjects. They found that the correlation between the ratings of these two groups ($r = .85$) was somewhat lower than the split-half reliability measures taken from the British data alone ($r = .92$), indicating possibly that British and U.S. groups differed from each other. Similarly, category structure based on the measure of associative frequency (how often subjects generate an exemplar name when given the category name) showed less similarity between the two cultural groups ($r = .76$) than a split-half reliability measures ($r = .93$) within the British group alone, again suggesting a possible role for cultural factors.

There is very little research of a cross-linguistic nature regarding category representation. Recently, Schwanenflugel and Rey (1986) attempted to exploit such cultural-linguistic differences in a study of the role of familiarity in the determination of typicality gradients. These researchers examined the graded structure of English and Spanish categories as determined by English and Spanish speaking monolinguals living in Florida. They found that for certain categories there were indeed cross-language goodness of membership differences (in terms of rated typicality) for exemplars within those categories. They also found cross-linguistic differences in the subjects' reported levels of familiarity with actual, specific exemplars. They concluded that the cross-linguistic differences in typicality gradients could be attributed to the cultural divergences in familiarity with certain exemplars. They further concluded that exemplar familiarity probably plays a role in determining category gradients within a given language.

Other researchers have investigated a role for familiarity in determining English language category gradients. Much of the focus has been on familiarity with the exemplars (referent or item familiarity) rather than with linguistic or label familiarity, and here Barsalou makes the important distinction between familiarity based on experiences with an item across all contexts (thus, *chair*

would usually be considered a more familiar item than *log*) and experiences with an item as instantiations of a specific concept (thus, *log* would be more familiar than *chair* in the context of *firewood*). However, even though the focus of such research has been on familiarity with *items*, often the measures used have been concerned with the ease of processing the names of exemplars. In general, the results of such studies have been inconclusive, partly because different studies have used different measures of familiarity, including number of properties generated given an exemplar label, rated familiarity and frequency of instantiation as a member of a category. For example, Ashcraft (1978) and Malt and Smith (1982) measured familiarity in terms of the number of properties a subject could think of. Because they found that subjects produced fewer properties for atypical than for typical items, they suggested that familiarity, thus defined, plays a role in determining category typicality gradients. However, as Schwanenflugel and Rey (1986) suggest, because highly typical items are likely to share more properties than less typical items (cf. Rosch & Mervis, 1975), property names of highly typical items may serve as prompts for the production of property names for other highly typical items, whereas the property names for low typical items would be less likely to prompt the generation of other property names for low typical items.

McCloskey (1980) found that rated meaning familiarity (degree to which a word's meaning is immediately obvious) may mediate the effect of typicality on response time to categorize. He reported that meaning familiarity correlated significantly ($-.61$) with comprehension time for category names, and when meaning familiarity was statistically held constant, the correlation between category comprehension time and time to decide if an instance belongs to the category dropped substantially, thereby implicating meaning familiarity in category verification. McCloskey further reported evidence that his measure of familiarity is different from word length and from measures related to objective written word frequency (similar to Boster, 1988, and Mervis et al., 1976).

Hampton and Gardiner (1983) collected measures of typicality, associative frequency and meaning familiarity (where a familiar word was defined as one whose meaning was immediately obvious) from a sample of British college students. They obtained significant partial correlations of rated typicality with associative frequency when meaning familiarity was held constant. While they concluded that meaning familiarity does not play a central role in determining the internal category structure, this result does suggest that speed of accessing a word's meaning may affect typicality gradients.

The present study was concerned with cross-linguistic differences in the graded structure of concepts of English and French speakers. In this study we focus on the role of "linguistic familiarity", here defined as perceived frequency of linguistic instantiation or perceived frequency of mention of item names *in*

the context of a particular category. We chose this factor to investigate because it may reflect important differences in the linguistic status accorded to items by a given language community. For example, suppose "ROBIN" is mentioned far more frequently in English than in its French translation equivalent "ROUGE-GORGE". We might expect English linguistic experiences to influence the place the concept robin has in English speakers' category gradient for BIRD differently from the way French linguistic experiences do for French speakers.

Unlike in the Hampton and Gardiner (1983) and the Schwanenflugel and Rey (1986) studies, we chose subjects who grew up and were educated in the same general geographic area, and who therefore shared similar physical environments throughout all or most of their lives (they would have encountered similar cars, buildings, fruits, etc.). In the Hampton and Gardiner and the Schwanenflugel and Rey studies the factors of linguistic groups differences (British versus U.S. English; Spanish versus U.S. English) were possibly confounded with other differences in physical environment (Britain versus USA; South and Central America versus USA origin). By testing subjects from the same general physical environment but who were nevertheless monolingual speakers of different languages, we hoped to explore the possible involvement of linguistic-cultural factors in determining the internal structure of categories.

Now, regardless of one's theory, one would naturally expect some cross-language differences in the representation of certain word meanings, especially in the case of many abstract and culturally dependent terms. For European French speakers, "PRESIDENT" has a somewhat different meaning from what it means for U.S. English speakers, and culturally loaded words such as "LOVE", "HONOUR", or "DUTY" will likely be understood differently too. Cross-linguistic similarities and differences have been observed in a number of different studies, including studies of word association (Kolars, 1963; Lambert & Moore, 1966; Taylor, 1976), semantic differential profiles of word meanings (Osgood, 1964), and cultural differences (Berlin, Breedlove, & Rave, 1973; Cole & Scribner, 1974). In addition, one would expect that different patterns of polysemy in different languages will play a role in determining cross-linguistic differences in word meaning; for example, the different meanings of English "TO STRIKE" do not correspond perfectly with the different meanings of French "FRAPPER".

The situation might be different, however, for words that refer to categories of concrete objects (e.g., fruit, bird, fish) where the referents do not or cannot vary to any great extent in the experience of the different language groups. That is why the present research focuses on such categories. According to a feature similarity approach one would expect robins and sparrows to bear the same relation to the category BIRD for English and French speakers living in similar physical environments since their experience with the physical exemplar would

be similar. On the other hand, according to a "naive theory" approach (Murphy & Medin, 1985), one might expect that cultural-linguistic factors, such as the role robins may play in English versus French literature, could influence how people categorize or select representative examples of categories.

METHOD

Subjects

Subjects were 50 native speakers of English (33 females) and 49 native speakers of French (41 females). The English speakers were students at an English language university in Montreal while the French speakers were students attending a French language university in Trois-Rivières (a French-speaking city near Montréal). The median age for the English speakers was 22 years and for the French speakers it was 28 years. In total 238 participants participated were administered the materials which included a brief biographical questionnaire, but only subjects who reported self-rated skill in the other language as 2 or less on a 4-point scale (4 = fluent bilingual) and who rated their exposure to the other language as being less than 30% of the time were retained for the study. The subjects included in the final sample could therefore be considered as monolinguals in that their exposure and fluency in the other language were minimal.

Materials

Six categories were selected from a set of norms for typicality gradients of 30 semantic categories. These gradient norms were based on the frequency of exemplar mention in a production task in which 300 monolingual English- and 300 monolingual French-speaking Canadians from the province of Québec were asked to generate the names of representative exemplars for each of the 30 different categories (Favreau & Segalowitz, 1980). These categories were selected because they had produced high English/French cross-language correlations in terms of the ordering of their exemplars (ANIMAL/ANIMAL, FRUIT/FRUIT, and TREE/ARBRE) and three were selected because they had produced low cross-language correlations (BIRD/OISEAU, FISH/POISSON, and WEAPON/ARME). Twelve exemplars were selected from each category such that they fulfilled the following conditions. The exemplars had to span the high, middle and low range of the associative frequency gradient (i.e., down to items generated by 2% of the 300 subjects in a given language in the Favreau and Segalowitz [1980] study). The selection was made so that all the items in one

language were the translation equivalents of those selected in the other. The new cross-language correlations of rank order of frequency of mention based on the twelve exemplars chosen in this manner for each category were: ANIMAL = .97, FRUIT = .95, TREE = .94, BIRD = .08, FISH = .14, and WEAPON = .16.

Procedure

Subjects were tested during regular university class periods. The subjects were given a booklet containing the Typicality Rating instructions (in English or French, as appropriate), plus six pages each containing one category name and the twelve exemplars with rating scales for that category. The instructions included the following:

"At the top of each page is the name of a category. Under it are the names of some members of this category. You are to evaluate how good an example of the category each item is. For example, if the category is dog and the item is pекinese, you are to evaluate how well a pекinese fits your idea of what makes something belong to the category dog. Use the seven point scale provided.

Notice that this kind of judgment does not necessarily have anything to do with how well you like a thing. You may prefer to own a pекinese without thinking it is the best breed to represent what you think makes something belong to the category dog. Notice, too, that this kind of judgment has nothing to do with how often you have encountered a thing. A church is a better example of a building than is a telephone booth, but you have probably seen telephone booths more often than churches."

The subjects received similar booklets containing the Linguistic Familiarity Rating instructions:

"At the top of each page is the name of a category. Under it are the names of items that belong to that category. You are to evaluate how frequently you encounter each word, in written or spoken form, as referring to a member of the category. You are to use the seven point scale provided."

On each page subjects were reminded to read through the entire list of 12 items before making their ratings in order to ensure that they were aware of the range of the category spanned by the items.

The order of the pages was randomized across subjects as was the order of appearance of the twelve exemplars on each page. The order of presentation of the typicality and linguistic familiarity rating tasks was counterbalanced across subjects within each linguistic group. The first booklet was collected after approximately fifteen minutes. Subjects kept the top sheet containing a code number for purposes of a prize drawing to be held at the end of the class period. The class lecture resumed and approximately thirty minutes later the students were given the second booklet. The students were instructed to write their prize drawing code number on the front page of this second booklet. They

then proceeded to fill the rating scales. The last page contained the language background questionnaire. The subjects were tested in classes of 20-30 students. All the students within one class received the same type of booklet (Linguistic Familiarity or Typicality) at the same time. At the end of the period, five students in the class were awarded \$5.00 each in a drawing of code numbers.

RESULTS

General characteristics of the data

The mean typicality and linguistic familiarity ratings for each exemplar of each category in English and French are shown in the Appendix. Exemplars are listed in rank order according to typicality (from highest to lowest) for English with the rank order for the corresponding French item given in parentheses. Overall, there were no significant patterns of differences as a function of whether the linguistic familiarity task or typicality task was performed first. Correlations between mean typicality ratings from subjects who performed the typicality rating first and those who performed it second ranged from a high of .97 (WEAPON) to a low of .57 (ANIMAL) for English speakers and from .98 (FRUIT) to .56 (ARBRE) for French speakers. Similarly, correlations between mean linguistic familiarity ratings from subjects who performed the linguistic familiarity rating first and those who performed it second ranged from a high of .95 (BIRD) to a low of .83 (FISH) for English speakers and from .99 (FRUIT) to .52 (ARME) for French speakers. Because the inter-order correlations tended to be high and no systematic violations from high correlations were evident, the data are collapsed across task order in all the analyses reported below.

Reliabilities¹ were calculated for the mean item typicality and linguistic familiarity ratings using means from even- and odd-numbered subjects in each language group (similar to the procedure used by Schwanenflugel & Rey, 1986).

1. An alternative way to measure rating reliability is to assess intersubject reliability as opposed to the degree of stability of the mean rating for each item, which can be high even when intersubject reliability is low (Barsalou, 1987; Barsalou & Sewell, 1984). The intraclass correlation (ICC) (Guilford & Fruchter, 1973; Shrout & Fleiss, 1979) is one such method. When the appropriate ICC values were calculated (corresponding to Shrout and Fleiss' index ICC(2,1)) the reliabilities ranged from .02 to .49. The mean ICC was .21 both for typicality and linguistic familiarity ratings; all but one of the ICC values were, however, statistically significant ($p < .001$) according to the test suggested by Shrout and Fleiss (1979). These indices of intersubject reliability are lower than analogous measures based on ranks obtained by Barsalou, typically in the .45 range as reported in Barsalou (1987). One possible source for this relative instability is the narrow range of ratings that were given to items in certain categories. For example, in the case of ANIMAL (ICC = .16), most typicality ratings were given in the high end of the scale, with subjects differing on whether they

Table 1 shows these results. All but three of the reliabilities were .84 or greater, the departures being for French typicality (ARBRE = .58, ANIMAL = .64) and French linguistic familiarity (ARME = .64). Overall reliabilities with data collapsed across category boundaries were .88 or greater.

TABLE 1. Reliabilities for the item ratings.

CATEGORY	English Typicality	English Linguistic Familiarity	French Typicality	French Linguistic Familiarity
Animal/animal	.90	.99	.64	.96
Fruit/fruit	.92	.99	.99	.99
Tree/arbre	.87	.97	.57	.93
Bird/oiseau	.84	.92	.99	.99
Fish/poisson	.99	.92	.99	.99
Weapon/arme	.97	.94	.99	.64
OVERALL	.91	.95	.92	.88

Note: $p < .01$ for all correlations.

TABLEAU 1. Accord inter-sujets dans l'évaluation des items.

Because the seven-point rating scale may have been interpreted differently for the Typicality and Linguistic Familiarity rating tasks, we analyzed the variability of responses on the scale. Analysis of variance of the standard deviations of subject responses in the Typicality versus the Linguistic Familiarity conditions give us an indication of the degree to which subjects used similar units when making ratings. Two-way analyses of variance (category $\times$ type of rating) performed separately for each group revealed main effects of category and rating but also a category by rating interaction ($F(5,240) = 12.24, p < .001$, and $F(5,240) = 28.08, p < .001$, for the English- and French-speaking groups respectively). The particular interactions we obtained, however, indicated that it

(Note 1 continued)

assigned a 6 or a 7 to any given item. Subjects agreed that the items were, for the most part, "highly representative" of the category; to the extent that they did assign slightly different ratings to different items, they did so in a relatively inconsistent manner. Thus, there were low intersubject reliabilities for some categories even though subjects gave ratings that were consistently at one end of the scale. By contrast, there were other cases where the ICC was considerably higher (e.g., WEAPON, $ICC = .49$) for which typicality ratings ranged more broadly over the whole seven point scale and did so in a relatively consistent manner.

was not the instructions that led to differential use of the scales; for some categories there were differences in one direction and for other categories there were differences in the other direction. Consequently, any subsequent low correlations between linguistic familiarity and typicality were deemed not likely due simply to subjects using different ranges of the rating scale for the two tasks.

Analyses related to the main hypotheses

The main question posed in this experiment was whether similarities and differences in typicality gradients across English and French reflect the degree to which English and French speakers have the same level of linguistic familiarity with the names for the exemplars concerned. There were two ways we addressed this issue. One was to examine the relation between linguistic familiarity and typicality ratings *within* a given language. Table 2 presents the Pearson product moment correlations between typicality and linguistic familiarity in each language group. As can be seen from Table 2, the obtained correlations vary considerably. The range in correlations across the six categories in English was from $-.27$ to $.70$, and in French from $-.37$ to $.81$. It was difficult to discern a consistent pattern in these data regarding a relationship between typicality and linguistic familiarity. For example, it would appear that for the category TREE there is a significant relationship in both languages (i.e., the perceived frequency of mention of specific tree names correlates with their typicality judgments). On the other hand, for the category ANIMAL there is at best only a weak relation in French and none in English, suggesting that typicality judgments of animals are not strongly related to perceived frequency of mention in the language.

TABLE 2. Intralingual Pearson Product moment correlations.

CATEGORY	English Typicality vs. Linguistic Familiarity	French Typicality vs. Linguistic Familiarity
Animal/animal	.09	.55
Fruit/fruit	.26	.14
Tree/arbre	.70*	.81**
Bird/oiseau	.18	.30
Fish/poisson	-.27	-.37
Weapon/arme	.47	-.30

* $p < .05$ ** $p < .01$

TABLEAU 2. Coefficients de corrélations de Pearson intra-linguistiques.

The second way we considered this issue was to see whether the linguistic familiarity factor accounted for cross-language similarities and differences in the typicality gradients. Our results indicate an interesting but complex relationship between the two variables of linguistic familiarity and rated typicality. The reader is referred to Table 3. Column 1 of the table shows the cross-language correlations for rated typicality. Column 2 shows the cross-language correlations for rated linguistic familiarity. To see the influence of rated linguistic familiarity on rated typicality we followed Schwanenflugel and Rey's (1986) approach of examining the changes in cross-language typicality correlations that result from statistically controlling for each group's rated linguistic familiarity with the items. The appropriate bipartial correlations (Timm & Carlson, 1976) can be seen in the column 3 of Table 3. In three cases the new correlation indicated a substantial change in the shared variance between English and French typicality ratings as a function of removing the influence of linguistic familiarity. Thus the new correlations for WEAPON and TREE accounted, respectively, for 30% and 23% less of the variance than was accounted for before removing the effect of linguistic familiarity. On the other hand, linguistic familiarity acted as a suppressor variable for the BIRD category since the bipartial correlation accounted for 15% more of the variance. These findings suggest that linguistic familiarity may play a role in cross-linguistic overlap in typicality but it does so for some categories only.

TABLE 3. Interlingual Pearson Product moment correlations for typicality (TYP) and linguistic familiarity (LFAM).

CATEGORY	TYP	LFAM	TYP with LFAM removed by bipartial correlation (% change in shared variance in parentheses)
Animal/animal	.91***	.95***	.94*** (+ 6%)
Fruit/fruit	.97***	.90***	.98*** (+ 2%)
Tree/arbre	.82**	.76**	.66** (- 23%)
Bird/oiseau	.34	.31	.52** (+ 15%)
Fish/poisson	.85**	.69*	.86*** (+ 2%)
Weapon/arme	.96***	.51	.79*** (- 30%)

* $p < .05$ ** $p < .01$ *** $p < .001$

TABLEAU 3. Coefficients de corrélations de Pearson Interlinguistiques pour la typicalité (TYP) et la familiarité linguistique (LFAM).

DISCUSSION

The results shown in Table 3 pose several interesting questions. One is to account for the different types of influence linguistic familiarity seems to have. In column 3 of Table 3 it can be seen that in some cases the linguistic familiarity factor had little or no influence on the cross-language pattern of rated typicality (ANIMAL, FRUIT, FISH). In other cases linguistic familiarity appears to have increased cross-linguistic similarity, since removing its influence resulted in a decrease in cross-linguistic similarity of typicality (TREE and WEAPON). Finally, in one case linguistic familiarity masked a cross-linguistic similarity in category gradient (BIRD) since removing its influence resulted in an increase in cross-linguistic similarity of typicality.

These findings suggest that linguistic familiarity may play a role in cross-linguistic overlap in typicality but it does so for some categories only. To explain these results we suggest that the graded structures obtained from typicality ratings reflect several sources of knowledge individuals may have about the items being rated. One is direct experiential knowledge of members of the category, information that would be partially reflected in knowledge about feature similarity relations among the items. Judgments strongly influenced by this knowledge source would be expected to result in category gradients reflecting resemblance patterns among the actual exemplars. A second knowledge source is linguistically based information relating items to each other and to their categories, information that would be largely reflected in knowledge about frequency of mention. Judgments strongly influenced by this knowledge source would be expected to result in category gradients reflecting exemplar status in the specific linguistic culture of the subject. Finally, there are undoubtedly other sources of knowledge based on linguistic experience but not reflected in frequency of mention, for example, knowledge about the uses and occurrences of exemplars in particular situations or knowledge about other people's expectations concerning the items. Judgments based on such knowledge sources could be expected to yield yet again different category gradients from either of the two previous knowledge sources. We propose that when individuals make judgments of typicality, they draw upon one or other of these sources of knowledge, depending on the circumstances. Sometimes this will involve knowledge of feature relations, sometimes it will involve linguistically based knowledge, and sometimes still other sources of knowledge. This approach could be applied to the present results as follows.

In the case of the ANIMAL and FRUIT categories, we may presume that the subjects had relatively frequent experiences with particular exemplars of animals and fruits and so were in a position to make typicality judgments based on knowledge of feature similarity relations. Since both English and French

Table 1 shows these results. All but three of the reliabilities were .84 or greater, the departures being for French typicality (ARBRE = .58, ANIMAL = .64) and French linguistic familiarity (ARME = .64). Overall reliabilities with data collapsed across category boundaries were .88 or greater.

TABLE 1. Reliabilities for the item ratings.

CATEGORY	English Typicality	English Linguistic Familiarity	French Typicality	French Linguistic Familiarity
Animal/animal	.90	.99	.64	.96
Fruit/fruit	.92	.99	.99	.99
Tree/arbre	.87	.97	.57	.93
Bird/oiseau	.84	.92	.99	.99
Fish/poisson	.99	.92	.99	.99
Weapon/arme	.97	.94	.99	.64
OVERALL	.91	.95	.92	.88

Note: $p < .01$ for all correlations.

TABLEAU 1. Accord inter-sujets dans l'évaluation des items.

Because the seven-point rating scale may have been interpreted differently for the Typicality and Linguistic Familiarity rating tasks, we analyzed the variability of responses on the scale. Analysis of variance of the standard deviations of subject responses in the Typicality versus the Linguistic Familiarity conditions gave us an indication of the degree to which subjects used similar units when making ratings. Two-way analyses of variance (category $\times$ type of rating) performed separately for each group revealed main effects of category and rating but also a category by rating interaction ($F(5,240) = 12.24, p < .001$, and $F(5,240) = 28.08, p < .001$, for the English- and French-speaking groups respectively). The particular interactions we obtained, however, indicated that it

(Note 1 continued)

assigned a 6 or a 7 to any given item. Subjects agreed that the items were, for the most part, "highly representative" of the category; to the extent that they did assign slightly different ratings to different items, they did so in a relatively inconsistent manner. Thus, there were low intersubject reliabilities for some categories even though subjects gave ratings that were consistently at one end of the scale. By contrast, there were other cases where the ICC was considerably higher (e.g., WEAPON, $ICC = .49$) for which typicality ratings ranged more broadly over the whole seven point scale and did so in a relatively consistent manner.

was not the instructions that led to differential use of the scales; for some categories there were differences in one direction and for other categories there were differences in the other direction. Consequently, any subsequent low correlations between linguistic familiarity and typicality were deemed not likely due simply to subjects using different ranges of the rating scale for the two tasks.

Analyses related to the main hypotheses

The main question posed in this experiment was whether similarities and differences in typicality gradients across English and French reflect the degree to which English and French speakers have the same level of linguistic familiarity with the names for the exemplars concerned. There were two ways we addressed this issue. One was to examine the relation between linguistic familiarity and typicality ratings *within* a given language. Table 2 presents the Pearson product moment correlations between typicality and linguistic familiarity in each language group. As can be seen from Table 2, the obtained correlations vary considerably. The range in correlations across the six categories in English was from $-.27$ to $.70$, and in French from $-.37$ to $.81$. It was difficult to discern a consistent pattern in these data regarding a relationship between typicality and linguistic familiarity. For example, it would appear that for the category TREE there is a significant relationship in both languages (i.e., the perceived frequency of mention of specific tree names correlates with their typicality judgments). On the other hand, for the category ANIMAL there is at best only a weak relation in French and none in English, suggesting that typicality judgments of animals are not strongly related to perceived frequency of mention in the language.

TABLE 2. Intralingual Pearson Product moment correlations.

CATEGORY	English Typicality vs. Linguistic Familiarity	French Typicality vs. Linguistic Familiarity
Animal/animal	.09	.55
Fruit/fruit	.26	.14
Tree/arbre	.70*	.81**
Bird/oiseau	.18	.30
Fish/poisson	-.27	-.37
Weapon/arme	.47	-.30

* $p < .05$ ** $p < .01$

TABLEAU 2. Coefficients de corrélations de Pearson Intra-linguistiques.

speakers would have had similar direct experiences (similar apples, dogs, etc.) their judgments should be similar, as indeed they were, as reflected in the high cross-language correlations of rated typicality. We can see from Table 3 that there was considerable cross-language similarity on the measure of linguistic familiarity too. However, removal of this source had little effect on the adjusted measure of typicality, a result consistent with the proposal that here typicality judgments were largely based on nonlinguistic sources of knowledge.

Consider now TREE and WEAPON. In the case of TREE we may presume that the subjects had relatively less direct experience with the category. While most city dwellers encounter trees, can probably recognize a few, and know about the leaves, bark and general shapes of trees, they seldom devote nearly as much attention to the featural similarities and differences among trees as they would, say, to animals or fruits. Given the name of a particular tree they probably would have had difficulty retrieving featural information specific to that tree. In the case of WEAPON, it is likely that relatively few subjects would have had extensive, direct experience with revolvers, bombs, swords, and the like. We would expect, then, that typicality judgments for TREE and WEAPON would not be based so much on direct experiential knowledge as in the cases of ANIMAL and FRUIT. The results shown in Table 3 are consistent with this view. For both TREE and WEAPON there were relatively high cross-language correlations for rated typicality and statistical removal of the linguistic familiarity component substantially reduced that cross-language typicality correlation. It would appear, then, that linguistic familiarity was a factor in the typicality ratings for TREE and WEAPON, and that removal of this influence reduced the similarity of the typicality gradients.

In the case of BIRD, we may again expect that speakers of both languages would have relatively less direct experience with exemplars, in so far as one seldom handles a bird or sees one up close for any length of time. Thus, we would expect typicality ratings not to reflect directly acquired experiential knowledge of features. Possibly linguistic sources of knowledge would come into play here as with TREE and WEAPON. This seems to have been the case, since when cross-language patterns of linguistic familiarity were statistically removed from the typicality correlations, the cross-language typicality correlation changed noticeably. In this case, however, the cross-language typicality correlation became more similar, not less similar as in the case of TREE and WEAPON. This suggests that subjects did make use of linguistically based knowledge, but this knowledge was different for English and French speakers. Indeed, Table 3 shows that the cross-linguistic correlation of linguistic familiarity was lowest (.31) for BIRD, indicating that in English and French different birds are given different linguistic status in terms of perceived frequency of mention. It would appear that in this case linguistically based knowledge

contributed to the cross-language difference in the typicality ratings originally observed. That is why the adjusted cross-language typicality correlation indicated greater similarity when this source of difference was removed. (It is important to note that recently Boster [1988] reported a study on the way people classify birds. He found that written word frequency does not correlate well with category gradients. This highlights the importance of distinguishing between objective measures of word frequency such as written word counts, and more subjective measures, such as *rated* word familiarity.)

Finally, in the case of FISH, we would have once more expected that speakers would not have extensive direct experience with particular exemplars, for reasons similar to the ones given with BIRD above. If this is correct, then we would expect typicality judgments to reflect some other source of knowledge. Now, while it appears that there was a moderate cross-language correlation for linguistic familiarity, statistical removal of this factor did not change the typicality correlation appreciably (by only 2%). Thus, the source of knowledge giving rise to the typicality judgments, in this case relatively similar judgments across English and French, does not seem to be reflected in the measure of linguistic familiarity we chose. We are forced to conclude that they used some other source of knowledge (one possibility, for example, may be knowledge about the specific items as pet fish, as edible fish, or as fish-like [dolphin]). Such other sources of knowledge would presumably be similar for English and French speakers, and hence could explain the cross-language similarity we obtained in rated typicality. These as yet unidentified sources of knowledge may or may not reflect linguistic experience (knowledge gained through listening and reading), but they do not reflect perceived frequency of name instantiation according to our results.

The foregoing account makes a number of assumptions about the relative levels of direct experience people have with the exemplars of the various categories. These assumptions have been supported by the results of a study completed recently in our laboratory (Bauco, 1987). Bauco asked 30 English and 30 French speaking subjects to rate the typicality of *pictures* of the exemplars listed in the Appendix and to rate on a 7-point scale how frequently they encounter the actual item or a picture of the item shown (7 = very frequently). Later they were asked to name the item. Realistic colour photographs were used (an effort was made to use "prototypical" pictures, such as those found in picture dictionaries; in the case of the TREE category, an enlarged picture of a leaf accompanied the picture of the tree itself).

Bauco found that for the categories BIRD and FISH, exemplar encounter was rated relatively low (only 3.4 and 3.5, respectively) and that accuracy of naming was similarly low (53% and 35%). This supports our suggestion that in our study subjects would have had difficulty making typicality judgments based

on direct experience with the exemplars. In contrast, the encounter ratings for ANIMAL and FRUIT were higher (4.1 and 4.8, respectively) as was naming accuracy (96% and 95%) (the difference in encounter ratings for ANIMAL/FRUIT versus BIRD/FISH was significant [$F(1,46) = 7.81, p < .01$]). These figures support our contention that, given the name of an animal or fruit, subjects would have been able to make typicality judgments based on knowledge derived from direct experience with the exemplars. In the case of TREE, while exemplar encounter was rated relatively high (4.3), the rate of correct naming of exemplars was quite low (only 43%). Finally, for WEAPON the mean encounter rating was relatively low (3.6) as expected (naming accuracy was 89%). These findings also support our earlier suggestion that when given the name of a tree or weapon, subjects would have had difficulty making typicality judgments based on direct knowledge about the particular item named.

Thus, our results indicate that perceived linguistic familiarity (rated perceived frequency of name instantiation) can sometimes play a role in determining typicality gradients. The data we obtained are consistent with the idea that, for some categories at least, people do rely on linguistic knowledge for making judgments about representativeness, and that differences in linguistic familiarity may contribute to cross-language typicality patterns.

These results support and extend the view of category gradients proposed by Barsalou (1987) and Murphy and Medin (1985). That is, our representations of categories are not necessarily based primarily on the knowledge we have about the perceptual similarities of the members of the category (family resemblances, central tendencies, or other feature-based approaches). Barsalou (1987) has suggested that category gradients reflect a constructive process, a view consistent with our own conclusion that, depending on the category involved, people will draw on particular sources of information they have about the category when thinking about typical examples. In the present study we have seen that one particular linguistic source -- one related to perceived frequency of name instantiation -- may be relevant when the individual has relatively little direct experience with category members (as in the case of BIRD). The contribution of this factor was highlighted by the cross-language comparisons.

The present results suggest that category structure, as reflected in the existence of membership gradients, is itself a reflection of a number of different sources of information that the individual may draw upon when performing tasks with categories. Such sources may include featural knowledge based on direct experiences with the items being judged, but may also include language-based knowledge concerning the name of the referent being judged. Language-based knowledge may reflect familiarity with the name, but it may of course also reflect other types of knowledge gained via language rather than by direct

experience (i.e., linguistic familiarity is by no means the only appropriate reflection of how rooted in language one's knowledge source is in a given case). The source of knowledge an individual draws upon at any moment will depend on the usefulness of that source for task purposes, and thus will depend on various task factors (e.g., size of the list, nature of the judgment to be made) as well as the person's expectations and past experience with the items. Thus, the results that have emerged from our cross-language comparisons suggest that category gradients are not always to be explained simply in terms of feature relations or familiarity (however defined) or even some other factor, not even for taxonomic categories; rather one or other factor may be involved in particular cases, depending on task demands and possibly on the amount and kind of experience the individual has had with the exemplars of a particular category.

These results raise questions concerning the stability of category representations. At one extreme is the possibility that category representations are not fixed but are always constructed anew to meet the demands of each particular situation. Alternatively, it may be that while category representations are flexible, they are not infinitely so. Perhaps there exists a hierarchy of knowledge sources that influences the category structure observed in any given situation (for example, knowledge based on direct experience, when available, may take precedence over other types of knowledge). In this view the stability of category representations within and across individuals would depend on the degree to which different situations activate the same knowledge sources across individuals and across time and on the degree to which the content of those sources is similar across individuals. Moreover, some sources of exemplar knowledge (e.g., feature knowledge gained through direct experience with exemplars) are likely to be highly similar across individuals and thus result in relatively stable cross individual category representations while other sources (e.g., linguistically derived knowledge) may be more variable across individuals. The stability of category representation could thus be addressed, not just in terms of the structural content (e.g., gradients) observed in particular situations, but in terms of which knowledge sources are available. Thus, it may be important to know what kinds of knowledge an individual possesses about the exemplars -- direct perceptual knowledge, linguistically based knowledge or other types of knowledge -- and which of these various sources are actually accessed in a given situation.

ACKNOWLEDGMENTS

The research reported here was supported in part by grants to the first author from the Natural Sciences and Engineering Research Council of Canada (A2607) and from the Québec Ministry of Education (FCAR-EQ3210) and by a grant to the second author from

the Québec Ministry of Education (FCAR-EQ2963). The authors would like to thank Drs. Tannis Arbuttle-Maag, Rex Kline, and Mel Komoda for their comments on earlier versions of this report, and to Claire Carrière, Johanne Courte, Piero Narducci, Cathy Poulsen, and Beth Sissons for their assistance and comments throughout the various phases of this project.

RESUME

La présente étude porte sur les facteurs sous-jacents à l'inégale représentativité des exemplaires d'une catégorie. Des sujets unilingues francophones et anglophones vivant dans le même environnement physique ont évalué la typicalité (représentativité) ainsi que leur familiarité linguistique (fréquence du terme comme exemplaire de la catégorie) avec ces mêmes exemplaires. Des analyses corrélationnelles ont été effectuées afin de déterminer le rôle de la familiarité linguistique comme facteur sous-jacent au degré de typicalité. Les résultats obtenus appuient l'hypothèse selon laquelle la structure graduée des catégories peut, à l'occasion, être attribuable à des facteurs autres que la similarité perceptive entre les exemplaires de la catégorie. Plus spécifiquement, certaines différences interlinguistiques dans le degré de typicalité peuvent correspondre à des différences interlinguistiques de familiarité linguistique.

REFERENCES

- Anglin, J.M. (1977). *Word, object, and conceptual development*. New York: Norton.
- Armstrong, S.L., Gleitman, L.R., & Gleitman, H. (1983). What some concepts might not be. *Cognition*, 13, 263-308.
- Ashcraft, M.H. (1978). Property norms for typical and atypical items from 17 categories: a description and discussion. *Memory and Cognition*, 6, 227-232.
- Barsalou, L.W. (1983). Ad hoc categories. *Memory and Cognition*, 11, 211-227.
- Barsalou, L.W. (1985). Ideals, central tendency and frequency of instantiations determinants of graded structures in categories. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 11, 629-654.
- Barsalou, L.W. (1987). The instability of graded structure: Implications for the nature of concepts. In U. Neisser (Ed.), *Concepts and conceptual development*. New York: Cambridge University Press.
- Barsalou, L.W., & Medin, D.L. (1986). Concepts: Static definitions or context-dependent representations? *Cahiers de Psychologie Cognitive/European Bulletin of Cognitive Psychology*, 6 (2), 187-202.
- Barsalou, L.W., & Sewell, D.R. (1984). Constructing representations of categories from different points of view. Technical Report N° 2. Atlanta, GA: Emory University, Emory Cognition Project.
- Bauco, P. (1987). *The relationship between familiarity and typicality in picture categorization: A cross-linguistic study*. Unpublished Honours Thesis, Concordia University.
- Berlin, B., Breedlove, D.E., & Raven, P.H. (1973). General principles of classification and nomenclature in folk biology. *American Anthropologist*, 75, 214-242.
- Bjorklund, D.E., Thompson, B.E., & Ornstein, P.A. (1983). Developmental trends in children's typicality judgments. *Behavior Research Methods and Instrumentation*, 15, 350-356.
- Bjorklund, D.F., Smith, S.C., & Ornstein, P.A. (1982). Young children's release from proactive interference: The effects of category typicality. *Bulletin of the Psychonomic Society*, 20, 211-213.
- Boster, J. (1988). Natural sources of internal category structures: Typicality, familiarity, and similarity of birds. *Memory & Cognition*, 16, 258-270.
- Carroll, P., & Slowiczek, M.L. (1986). Constraints on semantic priming in reading: a fixation time analysis. *Memory & Cognition*, 14, 509-522.
- Carson, M.I., & Abrahamson, A. (1976). Some members are more equal than others: The effect of semantic typicality on class-inclusion performance. *Child Development*, 47, 1186-1190.
- Chaplin, W.F., John, O.P., & Goldberg, L.R. (1988). Conception of states and traits: Dimensional attributes with ideals as prototypes. *Journal of Personality and Social Psychology*, 54, 541-557.
- Cole, M., & Scribner, S. (1974). *Culture and thought*. New York: Wiley.
- Duncan, E.M., & Kellas, G. (1978). Developmental changes in the internal structure of semantic categories. *Journal of Experimental Child Psychology*, 26, 328-340.
- Favreau, M., & Segalowitz, N. (1980). Semantic category norms for the French and English Quebec populations. Paper presented at the meeting of the American Psychological Association, Montréal.
- Guilford, J.P., & Fruchter, B. (1973). *Fundamental statistics in psychology and education* (Fifth ed.). New York: McGraw Hill.
- Hampton, J.A., & Gardiner, M.M. (1983). Measures of internal category structure: a correlational analysis of normative data. *British Journal of Psychology*, 74, 491-516.
- Homa, D., Sterling, S., & Trepel, L. (1981). Limitations of exemplar-based generalization and the abstraction of categorical information. *Journal of Experimental Psychology: Human Learning and Memory*, 7, 418-439.
- Keller, D. (1982). Developmental effects of typicality and superordinate property dominance in sentence verification. *Journal of Experimental Child Psychology*, 33, 288-297.
- Kolers, P.A. (1963). Interlingual word associations. *Journal of Verbal Learning and Verbal Behavior*, 2, 291-300.
- Lambert, W.E., & Moore, N. (1966). Word-association responses: comparisons of American and French monolinguals with Canadian monolinguals and bilinguals. *Journal of Personality and Social Psychology*, 3, 313-320.
- Mali, B.C., & Smith, E.E. (1982). The role of familiarity in determining typicality. *Memory & Cognition*, 10, 69-75.
- McCloskey, M. (1980). The stimulus familiarity problem in semantic memory research. *Journal of Verbal Learning and Verbal Behavior*, 19, 485-502.
- Medin, D.L., & Schaffer, M.M. (1978). Context theory of classification learning. *Psychological Review*, 85, 207-238.

Medin, D.M., Wattenmaker, W.D., & Hampson, S. (1987). Family resemblance, conceptual cohesiveness, and category construction. *Cognitive Psychology*, 19, 242-279.

Mervis, C.B., Catlin, J., & Rosch, E. (1976). Relationships among goodness-of-example, category norms, and word frequency. *Bulletin of the Psychonomic Society*, 7, 283-284.

Murphy, G. (1982). Cue validity and levels of categorization. *Psychological Bulletin*, 91, 174-177.

Murphy, G., & Medin, D.L. (1985). The role of theories in conceptual coherence. *Psychological Review*, 92, 289-316.

Nelson, K. (1974). Variations in children's concepts by age and category. *Child Development*, 45, 577-584.

Osgood, C.E. (1964). Semantic differential technique in the comparative study of cultures. *American Anthropologist*, 5, 146-169.

Posner, M.I., & Keele, S.W. (1968). On the genesis of abstract ideas. *Journal of Experimental Psychology*, 77, 353-363.

Rosch, E. (1975). Cognitive representations of semantic categories. *Journal of Experimental Psychology: General*, 104, 192-233.

Rosch, E. (1977). Human categorization. In N. Warren (Ed.), *Advances in cross-cultural psychology*, Vol. 1. London: Academic Press.

Rosch, E., & Mervis, C. (1975). Family resemblance: studies in the internal structure of categories. *Cognitive Psychology*, 7, 573-605.

Roth, E.M., & Shoben, E.J. (1983). The effect of context on the structure of categories. *Cognitive Psychology*, 15, 346-376.

Saltz, E., Soller, E., & Sigel, J.E. (1972). The development of natural language concepts. *Child Development*, 43, 1192-1202.

Schwanefflugel, P., & Rey, M. (1986). The relationship between category typicality and concept familiarity: evidence from Spanish- and English-speaking monolinguals. *Memory & Cognition*, 14, 150-163.

Shrout, P.E., & Fleiss, J.L. (1979). Intraclass correlations: Uses in assessing rater reliability. *Psychological Bulletin*, 86, 420-428.

Smith, E., & Medin, D.L. (1981). *Categories and concepts*. Cambridge, MA: Harvard University Press.

Taylor, I. (1976). Similarity between French and English words a factor to be considered in bilingual language behavior? *Journal of Psycholinguistic Research*, 5, 85-94.

Timm, N.H., & Carlson, J.E. (1976). Part and bipartial canonical correlation analysis. *Psychometrika*, 41, 159-176.

White, T.G. (1982). Naming practices, typicality, and underextension in child language. *Journal of Experimental Child Psychology*, 33, 324-346.

Received September 11, 1989
Accepted August 28, 1990

APPENDIX: Mean typicality and familiarity ratings.

CATEGORY	ITEM	Rated typicality		Rated linguistic familiarity	
		English	French	English	French
ANIMAL	horse (3)	6.51	6.70	5.18	5.02
	lion (4)	6.31	6.50	3.32	3.06
	dog (1)	6.22	6.80	6.72	6.25
	tiger (6)	6.22	6.40	3.12	2.98
	cat (2)	6.19	6.71	6.54	6.10
	cow (5)	5.90	6.46	4.94	5.20
	wolf (8)	5.90	6.36	3.14	3.33
	elephant (9)	5.77	6.32	3.32	2.98
	giraffe (10)	5.61	6.14	2.50	2.78
	pig (11)	5.39	6.02	4.08	4.16
	rabbit (7)	5.35	6.36	4.16	4.44
	rat (12)	4.08	5.57	4.96	3.68
FRUIT	apple (1)	6.84	6.88	6.63	6.33
	orange (2)	6.61	6.76	6.31	5.96
	strawberry (4)	6.26	6.58	4.96	6.00
	banana (3)	6.22	6.58	5.92	5.84
	grape (6)	6.16	6.36	5.43	5.38
	cherry (5)	5.98	6.44	4.24	4.37
	raspberry (7)	5.94	6.32	3.90	4.74
	plum (10)	5.84	6.04	4.04	4.00
	pineapple (8)	5.49	6.22	3.75	4.40
	melon (9)	5.37	6.10	3.77	4.37
	lemon (11)	4.55	5.52	4.47	4.33
	tomato (12)	3.22	3.76	5.59	5.69
TREE	oak (3)	6.74	6.53	4.69	4.40
	maple (1)	6.71	6.82	5.29	5.86
	elm (8)	6.28	6.17	3.41	3.57
	pine (6)	6.16	6.26	5.29	5.04
	apple tree (9)	6.08	6.06	4.86	5.06
	birch (2)	6.04	6.71	3.82	5.23
	willow (5)	5.85	6.27	2.87	3.57
	spruce (4)	5.80	6.38	3.75	5.21
	weeping willow (7)	5.77	6.20	3.44	3.73
	walnut (12)	4.88	5.74	2.92	2.62
	palm (11)	4.85	5.83	3.68	2.69
	cherry (10)	4.76	5.88	2.78	3.32

N.B. Items are shown in rank order according to English Typicality rating. French typicality rank is given in parentheses.

APPENDIX (continued)

CATEGORY	ITEM	Rated typicality		Rated linguistic familiarity	
		English	French	English	French
BIRD	robin (5)	6.65	6.25	4.69	3.49
	budgie (8)	6.63	5.98	3.71	4.35
	blue-jay (7)	6.14	6.05	3.42	3.26
	crow (3)	5.88	6.33	4.20	5.04
	hawk (10)	5.86	5.74	3.37	2.71
	swallow (1)	5.84	6.86	2.96	5.10
	seagull (6)	5.82	6.08	4.82	4.04
	hummingbird (11)	5.75	4.98	3.12	2.64
	woodpecker (4)	5.75	6.31	3.06	3.63
	jay (9)	5.66	5.88	2.79	2.93
	finch (2)	4.88	6.58	1.69	3.86
	duck (12)	4.39	4.85	4.80	4.90
FISH	cod (2)	6.14	6.06	4.40	4.14
	goldfish (1)	6.04	6.52	4.49	4.45
	haddock (8)	6.04	5.36	3.84	4.02
	sole (5)	5.94	5.70	4.45	5.04
	halibut (7)	5.92	5.63	3.85	2.75
	herring (3)	5.49	5.75	3.74	2.75
	carp (9)	5.40	5.31	2.29	3.01
	shark (6)	4.96	5.68	4.26	3.08
	sardine (4)	4.92	5.72	4.22	3.65
	dolphin (10)	3.73	5.12	3.45	3.27
	lobster (11)	3.31	4.00	4.96	4.59
	shrimp (12)	3.31	3.84	4.94	5.30
WEAPON	machine gun (2)				
	revolver (1)	6.80	6.67	5.02	3.10
	grenade (6)	6.71	6.74	5.24	4.14
	bomb (5)	6.41	6.04	4.06	2.66
	missile (3)	6.33	6.12	5.73	3.90
	cannon (4)	6.26	6.44	5.37	3.42
	sword (7)	5.43	6.25	2.71	2.90
	bow & arrow (8)	5.33	5.65	2.71	2.74
	axe (9)	4.47	4.96	2.36	3.14
	stick (10)	3.53	4.04	3.18	3.49
	hand (12)	2.92	3.86	3.70	3.60
	rope (11)	2.90	2.43	4.63	3.65
		2.65	2.65	3.45	3.88

N.B. Items are shown in rank order according to English Typicality rating. French typicality rank is given in parentheses.