

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Large-Scale Empirical Analyses of the Abstract/Concrete Distinction

Permalink

<https://escholarship.org/uc/item/21s9j46t>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 35(35)

ISSN

1069-7977

Authors

Hill, Felix
Korhonen, Anna
Bentz, Christian

Publication Date

2013

Peer reviewed

Large-Scale Empirical Analyses of the Abstract/Concrete Distinction

Felix Hill (fh295@cam.ac.uk)¹, Anna Korhonen (alk23@cam.ac.uk)¹, Christian Bentz (cb696@cam.ac.uk)²

¹ Computer Laboratory, University of Cambridge

² Department of Theoretical and Applied Linguistics, University of Cambridge

Abstract

We present original evidence that abstract and concrete concepts are organized and represented differently, based on statistical analyses of thousands of concepts in publicly available datasets. First, we show that abstract and concrete concepts have differing patterns of association with other concepts. Second, we test recent hypotheses that abstract concepts are organized according to association, whereas concrete concepts are organized according to (semantic) similarity. Third, we present evidence suggesting that concrete representations are more strongly feature-based than abstract representations. We argue that degree of feature-based structure may fundamentally determine concreteness, and discuss implications for cognitive and computational models of meaning.

Keywords: Concreteness; concepts; similarity; association.

Introduction

A large body of empirical evidence indicates important cognitive differences between abstract concepts, such as *guilt* or *obesity*, and concrete concepts, such as *chocolate* or *cheeseburger*. It has been shown that concrete concepts are more easily learned and remembered than abstract concepts, and that language referring to concrete concepts is more easily processed (Schwanenflugel, 1991). Moreover, there are cases of brain damage in which either abstract or concrete concepts appear to be specifically impaired (Warrington, 1975). In addition, functional magnetic resonance imaging (fMRI) studies implicate overlapping but partly distinct neural systems in the processing of the two concept types (Binder et al., 2005). Despite these widely known findings, however, there is little consensus on the cognitive basis of the observed differences (Schwanenflugel, 1991). Indeed, while many studies of conceptual representation and organization focus on concrete domains, comparatively little has been established empirically about abstract concepts.¹

In this paper we test various theoretical claims concerning the abstract/concrete distinction by exploiting large publicly-available experimental datasets and computational resources. By analyzing thousands of abstract and concrete concepts, our approach marginalizes potential confounds more robustly than in smaller-scale behavioral studies. In Analysis 1 we show that abstract concepts are associated in the mind to a wider range of other concepts, although the degree of this association is typically weaker than for concrete concepts. In Analysis 2 we explore the basis of these associations by testing the hypothesis that similarity

predicts association for concrete concepts to a greater extent than for abstract concepts. In Analysis 3, we show that free-association is a more symmetric relation for abstract concepts than for concrete concepts. The findings together suggest contrasts in both the organization and representation of abstract and concrete concepts. We conclude by discussing the implications of the findings for existing theories and models of conceptual representation.

Data

Our analyses exploit three publicly available resources compiled to assist psychological modeling and analysis.

USF Norms All three experimental analyses use the University of South Florida (USF) Free-association Norms (Nelson & McEvoy, 2012). The USF data consists of over 5,000 words and their associates. In compiling the data, more than 6,000 participants were presented with cue words and asked to “*write the first word that comes to mind that is meaningfully related or strongly associated to the presented word*”. For a cue word *c* and an associate *a*, the Forward Association Probability (FAP) from *c* to *a* is the proportion of participants who produced *a* when presented with *c*. FAP is thus a measure of the strength of an associate *relative to other associates of that cue*.

Many of the cues and associates in the USF data have a concreteness score, derived from either the norms of Paivio, Yuille and Madigan (1968) or Toglia and Battig (1978). In both cases contributors were asked to rate words based on a scale of 1 (very abstract) to 7 (very concrete).²

WordNet WordNet is a tree-based lexical ontology containing over 155,000 words produced manually by researchers at Princeton University (Felbaum, 1998). The present work used WordNet version 3.0.

Brown Corpus Word frequencies were extracted from the one million-word Brown Corpus (Kucera & Francis, 1967), chosen because it is an American corpus compiled at a similar time to the USF data. Word tokens in the Brown Corpus are tagged for their part of speech (POS). For a word type it is then possible to extract the *majority POS* (the POS with which the type is most frequently tagged).

¹ Notwithstanding a body of theoretical work (see e.g. Markman and Stilwell, 2001).

² Although concreteness is well understood intuitively, it lacks a universally accepted definition. It is often described in terms of reference to sensory experience (Paivio et al., 1968), but also connected to specificity; *rose* is often considered more concrete than *flora*. The present work does not address this ambiguity.

Analyses

Each of our analyses is motivated by characteristics of the abstract/concrete distinction proposed in theoretical and behavioral studies.

Analysis 1: Patterns of Association

Motivation Schwanenflugel’s *Context Availability Model* (1991) offers a theoretical basis for the aforementioned empirical abstract/concrete differences. Her exposition of the model relies on the following hypothesis:³

(H1) *Abstract concepts have more (but weaker) connections (to other concepts) than concrete concepts.*

Schwanenflugel presents only small-scale behavioral experiments (64 words, 40 participants) in support of H1. In Analysis 1 we test H1 on a far larger data set.

Method We extracted those 3,255 pairs in the USF data for which the concreteness of the cue-word was known. Since cue words are connected to a finite set of associates by FAP values, we can isolate a probability distribution over associates for each cue. Since our measure of association strength (FAP) is relative, it is not possible to compare these strengths directly across cue words. Nonetheless, we can make inferences about absolute cue associate strength from properties of the FAP distributions. If a cue has many associates with little variance in the FAP distribution, each FAP value must necessarily be low (and absolute association strength intuitively weak). In contrast, for a given number of associates, higher variance implies that some FAP values are notably higher than the mean, and thus likely to be strong absolutely. Therefore, to address H1 we considered both the dimension (number of associates) and the variance of the FAP distribution for each cue word.

In an initial analysis of the data, we noted a moderate but significant negative correlation between concreteness and frequency, $r(3255) = -.16$, $p < .001$. Therefore, a multiple regression analysis was conducted with $\log(\text{Frequency})$, Number of Associates and Variance of FAP as predictors, and Concreteness as dependent variable. Because the Concreteness/Frequency multicollinearity was exacerbated by high frequency abstract prepositions and verbs, a second analysis was conducted solely over cue words with majority POS ‘noun’ ($n = 2,320$).

Results and Discussion In both cases the regression model explained 17% of the variance of Concreteness and was statistically significant. The beta coefficients in Table 1 indicate that concreteness correlates negatively with both #Associates and FAP Variance. Both are highly significant

predictors even when controlling for frequency as an independent predictor.

We have shown that abstract words have more associates than concrete words and lower variance in FAP distributions. This is consistent with the idea that the strength of their associates is on average weaker than for concrete words. Fig. 1 represents the strength of this effect visually. Whilst this confirmation of H1 is consistent with Schwanenflugel’s Context Availability model, it is also consistent with other theoretical characterizations of the abstract/concrete distinction (Paivio, 1986; Markman and Stilwell, 2001). We thus investigate the distinction in more detail in Analyses 2 and 3.

	<i>All words</i>		<i>Nouns only</i>	
	Coeff. (β)	t	Coeff. (β)	t
# Assocs	-0.04***	-16.70	-0.04***	-15.97
Variance	-18.01***	-5.85	-15.64***	-4.41
$\log(\text{Freq})$	-0.18***	-14.21	-0.12***	-7.87
	$R^2 = .17$,		$R^2 = .17$,	
	$F(3, 3196) = 211.82***$		$F(3, 2319) = 157.51***$	
* $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$				

Table 1: Multiple regression analysis of Concreteness

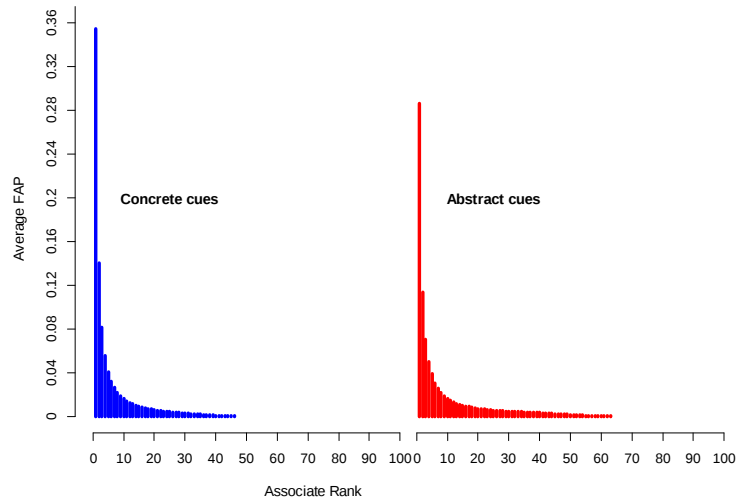


Figure 1: Average FAP mass at each associate rank over the 500 most abstract and concrete cue words in the USF data. Note the stronger initial associates in the concrete case and the longer tail of weak associates in the abstract case.

Analysis 2: Distinct Conceptual Organization?

Motivation Based on recent behavioral studies of healthy and brain-damaged subjects, (see e.g. Crutch et al., 2009), Crutch and colleagues argue that abstract and concrete concepts differ “*qualitatively*” in how they relate to other concepts. More specifically, they propose the following:

³ E.g. she states “What is important to this view is not how abstract words come to have weaker connections [to associated information]...only that they generally do” (1991, p. 243).

(H2) *Concrete concepts are organized in the mind according to similarity whilst abstract concepts are organized according to association.*

The terms *association* and *similarity* refer to the ways the concept pairs [*car*, *bike*] and [*car*, *petrol*] are related: *Car* is said to be semantically similar to *bike*, and associated with (but not similar to) *petrol*. Intuitively, *car* and *bike* may be understood as similar because of their common physical features (wheels), their common function (transport), or because they fall within a clearly definable category (modes of transport). By contrast, *car* and *petrol* may be associated because they often occur together or because of the functional relationship between them. The two relations are neither mutually exclusive nor independent; *bike* and *car* are related to some degree by both association and similarity.

In support of H2, Crutch et al. (2009) asked 20 participants to select the odd-one-out from lists of five words appearing on a screen. The lists comprised either concrete or abstract words (based on ratings of six informants) connected either by similarity (e.g. *dog*, *wolf*, *fox* etc.; *theft*, *robbery*, *stealing* etc.) or association (*dog*, *bone*, *collar* etc.; *theft*, *law*, *victim* etc.), with an unrelated odd-one-out item in each list. Controlling for frequency and position, subjects were both significantly faster and more accurate if the related words were either abstract and associated or concrete and similar. These results support H2 on the basis that decision times are faster when the related items form a more coherent group, rendering the odd-one-out more salient.

Despite the consistency in these findings, each of Crutch et al.'s experiments tested a small sample of subjects (< 20) with a small (< 20) number of concepts. It is therefore possible that the observed differences resulted from semantic factors particular to the subjects and items but independent of concreteness. Analysis 2 exploits the USF data and WordNet to investigate H2 more thoroughly.

Method Because similarity and association are not mutually exclusive, H2 can be interpreted in terms of differing interactions between these two relation types. If concrete concepts are organized in the mind to a greater extent than abstract concepts according to similarity, then the associates of a given concrete concept should be more similar to that concept than the associates of a given abstract concept. In other words, there should be greater correlation between similarity and association over concrete concepts than abstract. We test for this effect with a multiple regression over cue-associate pairs, with FAP as dependent variable and Concreteness, Similarity and their interaction as predictors. Relevant to H2 is the presence or absence of a positive interaction between concreteness and similarity.

Following other studies of conceptual structure (Markman & Wisniewski, 1997), we model similarity as proximity in a conceptual taxonomy, in this case, WordNet. Various measures of similarity have been developed for WordNet (see e.g. Resnik, 1995). *PathSim*, based on the shortest path between two senses, is perhaps the simplest,

and mirrors the manual approach taken by Markman & Wisniewski (1997). For this experiment, *SIM*, a measure of the similarity of two words w_1 and w_2 on the range [0, 1], was defined as the maximum *PathSim* between all senses of w_1 and all senses of w_2 . Since verbs, adjectives and nouns occupy separate taxonomic structures in WordNet, *PathSim* does not effectively measure similarity across these categories. We thus restrict our analysis to those 18,672 pairs in the FAP data for which cue concreteness and FAP are known and the majority POS for both words is 'noun'.

As a pre-test, *SIM* was evaluated on Rubinstein and Goodenough's (1965) similarity data for 65 word pairs,⁴ previously used as a benchmark for automatic similarity measures. The correlation between these judgments and *SIM*, $r(63) = .77$, $p < .05$, was comparable to other more complex WordNet metrics such as Resnik's (1995) Information Content, $r(63) = .79$, $p < .05$, and approaching the human replication baseline, $r(63) = .90$ (Resnik, 1995).

Results and Discussion As detailed in Table 2, the regression model was significant, $F(2, 3252) = 194.53$, and, as expected, *SIM* was a significant predictor of FAP. The interaction term *SIM:Concreteness* was positive, as predicted by H2, and a significant predictor of FAP.

Table 2: Multiple regression analysis of FAP over cue (noun) – associate (noun) pairs

	Coeff. (β)	t-value
<i>SIM</i>	0.048	3.66***
Concreteness	0.003	1.64
<i>SIM:Conc</i>	0.005	2.07*
----- $R^2 = .03$, $F(3, 18665) = 194.53$ -----		

* $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$

The positive interaction between similarity and concreteness in our model lends some support to H2. However, the size of this effect is small: the model explains less than .1 of a percentage point more variance in FAP than a model with no interaction term. While statistically significant, this difference is not consistent with a “*qualitative difference*” in conceptual organization between abstract and concrete concepts, as Crutch and Warrington (2005) propose. Rather, our analysis supports a gradual contrast in patterns of organization along a continuum from concrete to abstract. Of course, qualitative or categorical differences may exist that are too subtle to be detected by the current method. We intend to examine this possibility in future work, using the USF data and WordNet to generate appropriate items for larger-scale behavioral experiments.

Analysis 3: Distinct Conceptual Representation?

Motivation Hypothesis H2 (Analysis 2) characterizes the abstract/concrete distinction in terms of conceptual

⁴ Subjects were asked to consider their idea of synonymy and then rate the “*similarity of meaning*” of word pairs (1965, p. 628).

organization. With respect to the differences in representation that cause the H2 effect, Crutch and Warrington offer only speculative hypotheses. For instance, they suggest that that “*abstract concepts are represented in associative neural networks*”, whilst “*concrete concepts have a categorical organization*” (Crutch & Warrington, 2005, p. 624). Weimer-Hastings and Xu (2005) address this question empirically, and find that people tend to generate fewer “*intrinsic*” and proportionally more “*relational*” properties for abstract concepts. Nevertheless, given the untimed, conscious nature of their feature-generation task, and the fact they test only 31 subjects with 36 concepts, the strength of their findings is limited in a similar way to those of Crutch et al. In Analysis 3 we test for evidence of specific representational differences that could explain H2 and the other concreteness effects detailed in the Introduction.

Although the limitations of classical theories of representation with strict binary property specifications are well known, many recent theories characterize representations as *feature-based* in a more dynamic sense (see e.g. Plaut & Shallice, 1993). Indeed, the idea of concepts as complexes of conceptually basic features underlines explanations of various empirical observations, including typicality effects, category learning and category-specific semantic impairments (Tyler et al. 1995).

Feature-based models are not ubiquitous. Competing approaches such as spatial models (See e.g. Shepard, 1957) or associative networks (Steyvers & Tenenbaum, 2005) have also captured various established cognitive phenomena. One criticism of such models, however, is they naturally model relatedness with a symmetric operation: for all concepts x and y , $relatedness(x,y) = relatedness(y,x)$. As often observed, (Griffiths, et al., 2007; Tversky, 1977) empirical measures of conceptual proximity are in general asymmetric. For instance, it is common to find concept pairs X and Y for which subjects judge the statement ‘ X is like Y ’ to be more acceptable than ‘ Y is like X ’. This effect can be particularly evident when one concept is more salient or prototypical than the other (‘Justin Bieber is like Elvis’ vs. ‘Elvis is like Justin Bieber?’). Asymmetries are also observed in priming effects and free-association, for instance with category name/member or whole/part pairs (*Alsatian* primes *dog* more than *dog* primes *Alsatian*).

A noted strength of feature-based models is that they naturally capture the asymmetry of semantic relations. In the *Contrast Model*, Tversky (1977) proposes that the similarity of conceptual representations is computed as some continuous function of their common and distinctive features. Such operations are generally asymmetric, particularly given a disparity in the number of features. For instance, suppose the concept *jackal* is represented with the features {4LEGS, FUR, HOWLS} and the concept *dog* with the features {4LEGS, FUR, TAIL, COLLAR, LOYAL, DOMESTIC}. Tversky argues that it is more natural to say that *jackals* are like *dogs* than vice versa because two thirds of *jackal* features are shared by *dog*, whereas only one third of *dog* features are shared by *jackal*. As with other theories

of representation mentioned previously, Tversky’s demonstrations are typically confined to concrete words. Nevertheless, his conclusions could be aligned with H2 (Analysis 2) if the following hypothesis held:

(H3) *Concrete representations have a high degree of feature-based structure. Abstract representations do not.*

Indeed, the soundness of H3 could point to a causal explanation of the H2 effect. By H3, computing the similarity of abstract concepts by mapping features would be relatively hard. Alternative types of relation would thus be required to group sets of abstract concepts in the mind.

Proposals similar to H3 have been made by several researchers. Plaut and Shallice (1993) showed that integrating differential degrees of feature-based structure into connectionist simulations of dyslexia leads to more accurate replication of established concrete word advantages. Additionally, Markman and Stilwell’s (2001) analysis of conceptual category subtypes is entirely consistent with H3. On this view, *feature-based categories* include those noun concepts typically considered very concrete, whereas abstract noun, prepositions and verbs are all *relational categories*. Feature-based categories are represented by some configuration of (featural) information ‘*subordinate to*’ or ‘*contained within*’ that representation (p. 330), whereas relational categories are defined by external information, such as the position of the representation in a relational structure. Finally, H3 is also compatible with the feature-generation data of Weimer-Hastings and Xu (2005).

In Analysis 3 we exploit the USF data to test a prediction of H3. If Tversky’s demonstration that asymmetry derives from features is sound, there should be greater asymmetry between concrete concepts than between abstract concepts.

Method Although Tversky’s reasoning pertains to a similarity relation, we use the USF data to explore asymmetries in association. Similarity is an important factor in association in general, as evidenced by the high SIM/FAP correlation (Analysis 2). We thus expect asymmetries deriving from similarity to be reflected in FAP values, noting that asymmetry of free-association has been observed previously (Michels et al., 2011).

For each of the 18,668 ordered cue-associate pairs $[c,a]$ for which the concreteness of c and a is known, we calculate the (additive) asymmetry $|FAP[c,a] - FAP[a,c]|$. We define the total cue asymmetry, $CueAsymm(c)$, as the sum of the additive asymmetries over all associates of that cue. For a given cue item in our analysis, we experiment with three different measures of concreteness. The first is the cue concreteness $Conc(c)$. Since Tversky’s explanation of asymmetry relies on both concepts having a feature-based representation, for each pair $[c,a]$ we also calculate both the sum and the product of concreteness scores. We then define $ConcSum(c)$ as the sum of the sums over all associates, $ConcSum(c) = \sum_a Conc(c) + Conc(a)$, and $ConcProd(c)$ as the sum of products $ConcProd(c) = \sum_a Conc(c)Conc(a)$. To control for the possibility that FAP asymmetries are caused exclusively by a disparity in frequency between cue and

associate, we also define the measure $\text{FreqDisp}(c)$; the sum of the absolute differences between the frequency of a cue word and that of each of its associates, $\text{FreqDisp}(c) = \sum_a |\text{Freq}(a) - \text{Freq}(c)|$. We analyse the relationship between CueAsymm (dependent variable) and the three measures of concreteness (predictors) in separate multiple regression models, with FreqDisp as an independent predictor in each.

Results and Discussion The results in Table 3 show a significant positive correlation between the concreteness measure and CueAsymm in all three models, confirming the prediction of H3. Moreover, the model with ConcProd ($R^2 = .1373$) accounts for more of the CueAsymm variance than with ConcSum ($R^2 = .12$), which in turn accounts for more than with Conc ($R^2 = .08$). These two comparisons show that information about the concreteness of both cue and associate is important for predicting asymmetry, consistent with Tversky’s explanation of the link between features and asymmetry. It is also notable that FreqDisp is a (marginally) significant predictor in only one of the three models. Therefore the predictive relationship between concreteness and asymmetry (illustrated in Fig. 2) does not derive from discrepancies in frequency between words.

Table 3: Multiple regression analyses of CueAsymm

	Coeff. (β)	t
Conc	0.001***	16.28
FreqDisp	-0.000	-1.44

$R^2 = .08, F(2, 3252) = 135.60***$		
ConcSum	0.003***	21.33
FreqDisp	-0.000*	-2.43

$R^2 = .12, F(2, 3252) = 230.92***$		
ConcProd	0.001***	22.60
FreqDisp	-0.000	-0.39

$R^2 = .14, F(2, 3252) = 258.81***$		
* $p < .05$; ** $p < .01$; *** $p < .001$		

Table 4: USF pairs with highest and lowest asymmetry

Cue (conc)	Associate	FAP	Backward AP	Additive asymmetry
<i>Keg</i> (6.87)	<i>Beer</i> (5.83)	0.885	0	0.885
<i>Text</i> (5.80)	<i>Book</i> (6.09)	0.881	0	0.881
<i>Fish</i> (5.84)	<i>Trout</i> (5.93)	0.036	0.913	0.877

<i>How</i> (1.57)	<i>Method</i> (2.2)	0.014	0.014	0
<i>Honor</i> (1.75)	<i>Courage</i> (2.51)	0.014	0.014	0
<i>Immoral</i> (1.81)	<i>Dishonest</i> (2.63)	0.014	0.014	0

In a separate analysis, we observed that the ConcProd model over pairs in which the cue word is a noun ($R^2 = 0.1325$) fits better than the model over pairs in which the cue is a non-noun ($R^2 = 0.0987$) or specifically a verb ($R^2 = 0.114$). Indeed, across all 18,668 pairs, the mean additive asymmetry when both cue and associate are nouns (.071) is

significantly greater than when both are not (.066), $t(9351.3) = 2.78, p < .01$. Together with Tversky’s analysis, these observations are consistent with Markman and Stilwell’s proposal that many noun representations are feature-based whereas representations of verbs and prepositions rely on features to a lesser extent.

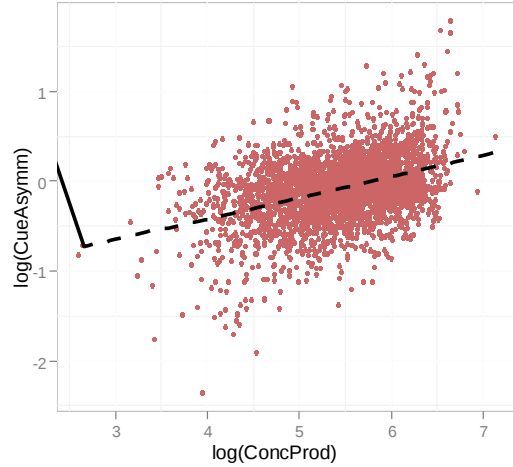


Figure 2: Scatterplot of CueAsymm vs. ConcProd .

Conclusion

In this study we have reported the following effects of increasing conceptual concreteness:

1. Fewer, but stronger associates (Analysis 1).
2. A stronger correlation between the similarity of concepts and the strength of their association (Analysis 2).
3. Greater asymmetry of association (Analysis 3).

These findings derive from analyses of thousands of concepts and data from thousands of subjects, an approach that significantly increases their robustness in comparison with previous behavioral experiments.

Finding 3 is consistent with, and arguably suggestive of, the view that concrete representations are more strongly feature-based than abstract concepts. Instead of a strongly feature-based structure, abstract representations encode patterns of relations with other concepts (both abstract and concrete). We hypothesize that the degree of feature-based structure is the fundamental cognitive correlate of what is intuitively understood as concreteness.

On this account, computing the similarity of two concrete concepts would involve a (asymmetric) feature comparison of the sort described by Tversky. In contrast, computing the similarity of abstract concepts would require a (more symmetric) comparison of relational predicates such as analogy or *structure-mapping* (Markman & Gentner, 1993). Because of their representational structure, the feature-based operation would be simple and intuitive for concrete concepts, so that similar objects (of close taxonomic categories) come to be associated. On the other hand, for abstract concepts, perhaps because structure mapping is more complex or demanding, the items that come to be associated are instead those that fill neighboring positions in

the relational structure specified by that concept (such as arguments of verbs or prepositions). Intuitively this would result in a larger set of associates than for concrete concepts, as confirmed by Finding 1. Moreover, such associates would not in general be similar, as supported by Finding 2.

If this is correct, it is likely that computational models of meaning could be improved by integrating a dimension of concreteness. For instance, models that connect words via syntagmatic co-occurrence would be particularly appropriate for modeling human association in abstract domains, whereas approaches based on taxonomies, or those measuring paradigmatic co-occurrence, would better reflect similarity and be more apt for concrete domains. In future work we plan to test these hypotheses by analyzing how concreteness is reflected in running text corpora.

Acknowledgments

Thank you to Barry Devereux for helpful comments. The work was supported by St John's College, the Royal Society and the Arts and Humanities Research Council.

References

- Binder, J., Westbury, C., McKiernan, K., Possing, E., & Medler, D. (2005). Distinct brain systems for processing concrete and abstract concepts. *Journal of Cognitive Neuroscience* 17(6), 905-917.
- Crutch, S., Connell, S., & Warrington, E. (2009). The different representational frameworks underpinning abstract and concrete knowledge. *Quarterly Journal of Experimental Psychology*, 62(7), 1377-1388.
- Crutch, S., & Warrington, E. (2005). Abstract and concrete concepts have structurally different representational frameworks. *Brain*, 128(3), 615-627.
- Felbaum, C. (1998). *WordNet: An Electronic Lexical Database*. Cambridge, MA: MIT Press.
- Griffiths, T., Steyvers, M., & Tenenbaum, J. (2007). Topics in semantic representation. *Psychological Review*, 114 (2), 211-244.
- Kucera, H., & Francis, W. (1967). *Computational Analysis of Present-day American English*. Providence, RI: Brown University Press.
- Markman, A., & Gentner, D. (1993). Structural alignment during similarity comparisons. *Cognitive Psychology*, 25, 431-467.
- Markman, A., & Stilwell, C. (2001). Role-governed categories. *Journal of Theoretical and Experimental Artificial Intelligence*, 13, 329-358.
- Markman & Wisniewski, E. (1997). Similar and different: The differentiation of basic level categories. *Journal of Experimental Psychology: Learning, Memory and Cognition*. 23(1), 54-70
- Michelsbaker, L., Evert, S., & Schütze, H. (2011). Asymmetry in corpus-derived and human word associations. *Corpus Linguistics and Linguistic Theory*, 7(2), 245.
- Nelson, D., & McEvoy, C. (2012). *The University of South Florida Word Association Norms*. Retrieved online from: <http://web.usf.edu/FreeAssociation/Intro.html>.
- Paivio, A., Yuille, J., & Madigan, S. (1968). Concreteness, imagery, and meaningfulness values for 925 nouns. *Journal of Experimental Psychology Monograph Supplement*, 76(1, Pt. 2).
- Plaut, D., & Shallice, T. (1993). Deep dyslexia: A case study of connectionist neuropsychology. *Cognitive Neuropsychology*, 10, 377-500.
- Resnik, P. (1995). Using Information Content to Evaluate Semantic Similarity in a Taxonomy. *Proceedings of IJCAI-95*.
- Rubenstein, H., & Goodenough, J. (1965). Contextual correlates of synonymy. *Communications of the ACM* 8(10), 627-633.
- Schwanenflugel, P. (1991). Why are abstract concepts hard to understand? In P. Schwanenflugel. *The psychology of word meanings* (pp. 223-250). Hillsdale, NJ: Erlbaum.
- Shepard, R. (1957). Stimulus and response generalization: a stochastic model relating generalization to distance in psychological space. *Psychometrika*, 22, 325-345.
- Steyvers, M., & Tenenbaum, J. (2005). The large-scale structure of semantic networks. *Cognitive Science* 29(1), 41-78.
- Toglia, M., & Battig, W. (1978). *Handbook of semantic word norms*. Hillsdale, N.J: Erlbaum.
- Tversky, A. (1977). Features of similarity. *Psychological Review*, 84, 327-352.
- Tyler, L., Moss, H., & Jennings, F. (1995). Abstract word deficits in aphasia: Evidence from semantic priming. *Neuropsychology*, Vol 9(3), 354-363.
- Warrington, E. (1975). The selective impairment of semantic memory. *Quarterly Journal of Experimental Psychology* 27(4), 635-657.
- Wiemer-Hastings, K., & Xu, X. (2005). Content differences for abstract and concrete concepts. *Cognitive Science* 29(5), 719-736.