

The modifier effect: Default inheritance in complex noun phrases.

James A. Hampton (hampton@city.ac.uk)

Department of Psychology, City University London, Northampton Square
London, EC1V OHB UK

Martin L. Jönsson (Martin.Jonsson@fil.lu.se)

Department of Philosophy, Lund University,
Kungshuset, Lundagård SE-222 22 Lund, Sweden

Alessia Passanisi (alepassa81@hotmail.com)

Department of Psychology, University of Catania,
95124 Catania, Italy

Abstract

The modifier effect is the reduction in judged likelihood of a generic statement (*Apples are sweet*) when the subject is modified (*Chinese apples are sweet*). Connolly, Fodor, Gleitman and Gleitman, (2007) argued that this effect undermines the principle of default property inheritance in conceptual combination. In a series of studies, we replicated the effect and its interaction with modifier typicality. We elicited justifications for the judgments, and found three common accounts were given – pragmatics, knowledge-based reasoning, and uncertainty about attribute inheritance. We also showed that the mutability of a property for the subject concept affected its judged likelihood when the concept was modified, and that likelihood judgments were correlated between modified and unmodified versions of sentences. It is argued that contrary to the claims of Connolly et al., the modifier effect provides clear evidence for the default inheritance of prototypical properties in modified concepts.

Keywords: Prototypes, Conceptual combination, Default inheritance, Generics, Modifier Effect

The Modifier Effect

Compositionality is the doctrine that the meaning of a complex phrase in language should be composed only of the meanings of its components and the syntax by which they are combined. A critical question for the problem of compositionality is the extent to which the complex concept representing a complex noun phrase “inherits” the default prototypical properties of the head noun concept. Prototypical properties are those generic statements such as “birds fly” or “crocodiles are dangerous” that may be considered *generally* or *typically* true, even when there may be known counterexamples (penguins or dead crocodiles). By default inheritance is meant the notion that when a complex concept is generated such as “Albino crocodile”, all the properties associated with crocodile will be inherited by the complex concept, except for those explicitly related to the modifier (such as color in this case).

In a seminal paper, Connolly, Fodor, Gleitman, and Gleitman (2007), (CFGG), showed that the judged likelihood of such sentences was reduced when the subject of the sentence was modified, even when the modifier was carefully chosen as one which would not itself directly affect the property. Moreover the effect was greater for atypical modifiers. CFGG argued from these results that complex concepts do not inherit the default prototypes of

the head noun – since if they did, the generic properties would have to be considered equally true of unmodified and modified noun phrases. Instead they proposed that concepts are combined by simple compositional rules. The prototype of a concept needs to be distinguished from the concept itself, which is an atomic representation pointing to a class in the external world.

We have addressed the basis of CFGG’s interpretation of their result in another paper (Jönsson & Hampton, 2008). The present research set out to examine the modifier effect with further empirical studies, three of which are reported here. Given that this is a new phenomenon, we first replicated the effect, with the same materials and design. In further experiments we then explored the basis of the effect, in order to throw further light on the processes involved.

Experiment 1

The first experiment was a partial replication of CFGG’s study.

Method

Participants. Twenty-nine students at City University London, participated for course credit.

Materials. Each booklet contained 40 target and 90 filler sentences. All were simple declarative sentences, consisting of a noun phrase and a predicate. Four versions of each target sentence were constructed to give a total of 160 sentences (the same sentences used by CFGG). The head noun could either be (a) unmodified, (b) modified by a *typical* modifier, generally true of things denoted by the head, (c) modified by an *atypical* modifier, not typically true of things denoted by the head noun, or (d) modified by *two atypical* modifiers, one of which was that used in condition (c). For instance, for the noun “ducks” the following 4 sentences were constructed:

- a) Ducks have webbed feet.
- b) Quacking ducks have webbed feet.
- c) Baby ducks have webbed feet.
- d) Baby Peruvian ducks have webbed feet.

The unmodified sentences were all typically true. The atypical modifiers in conditions (c) and (d) were chosen by CFGG such that while they would form relatively novel and unfamiliar phrases, yet they would still be compatible with the predicates. As they put it, “the introduction of the modifier does not necessitate a change in the applicability of

the predicate”. The target sentences were rotated across four booklets so that each booklet contained 10 target sentences with different head nouns for each of the 4 conditions. The 40 target sentences in each booklet were embedded randomly in 75 filler sentences and an additional 15 filler sentences appeared at the front of each list to avoid warm-up effects. Filler sentences could be unlikely, moderately likely, or highly likely.

Design and Procedure. Participants were randomly divided into 4 groups, each receiving one of the four booklets. They indicated the likely truth of each sentence using the numbers 1 through 10 (1 = very unlikely; 10 = very likely).

Results

Mean likelihood ratings. Table 1 shows that our results were largely in line with those obtained by CFGG. The data were submitted to ANOVA. The effect of modifier condition was significant ($\text{Min } F'(3, 198) = 17.83, p < .001$), and post-hoc pairwise comparisons between conditions were all significant ($p < .001$) except for that between the atypical and twice atypical modifier conditions. Modification of the subject noun reduced judged likelihood, and the effect was greater with atypical modifiers.

Table 1: Mean likelihood ratings for Experiment 1

Sentence condition	Mean (SD)
Unmodified	8.31 (3.55)
Typically modified	7.51 (3.31)
Atypically modified	6.59 (2.53)
Twice atypically modified	6.27 (3.05)

Correlations. Having replicated the results of the earlier study, we tested the hypothesis that, although rated likelihood is reduced by a modifier, the *relative* rated likelihood of different properties is still maintained. If a modified concept inherits the default properties of the noun, but with a general decrease in confidence, then there should be a positive correlation between the judged truth of properties for the unmodified noun N, and that for each of the different modified noun phrases. If the correlation were lacking, then that would clearly provide evidence against default inheritance. Mean rated likelihoods for the sentences in each condition were correlated across the 40 concepts. (Estimated pooled reliability was 0.6). The Unmodified, Atypically modified and Twice Atypically modified sentences all correlated with each other with $r(38)$ between .41 and .45 ($p < .001$). Typically modified sentences had lower, non-significant positive correlations with the other three. There was therefore evidence that (with the exception of the Typically modified sentences) the strength of a feature for the noun prototype was predictive of its strength for the modified noun phrase concepts.

Discussion

Experiment 1 demonstrated the modifier effect to be replicable and robust. In addition to the modifier affecting

mean ratings, there was good evidence that the relative strengths of properties for the head noun concepts were inherited by the modified noun concepts – with the exception of the Typically modified concepts. This pattern of correlation is consistent with the hypothesis that modifying the head noun with an atypical/unfamiliar modifier has a general depressive effect on judged likelihood of all properties. In order to explore the basis of the modifier effect further, in Experiment 2 participants were asked first to provide a comparative judgment of whether the two sentences (N and MN) were equally likely to be true or not, and were then asked to explain their judgments for those cases where they said they were not.

Experiment 2

The second experiment had two aims. The first was to test whether the reduction in likelihood would still be found if participants made a direct comparison between the two sentences. In Experiment 1 the same participant never judged both modified and unmodified versions of the same sentence. A stronger test of the modifier effect is to set the two sentences side-by-side and ask people to judge whether or not one was more likely, and if so which. This design has the advantage of drawing to the participants’ attention the explicit possibility that the two sentences are equally likely (as would be predicted by simple default inheritance). The design also allowed us to fulfill our second aim, which was to explore reasons for the effect by asking participants to justify their likelihood judgments. Unlike for single rating judgments it was quite reasonable to ask for a justification of a preference, as in “Why did you think this sentence was more/less likely than that one?” Filler sentences were again used to reduce response bias. So that giving justifications did not influence the original decisions, participants made all their judgments first, and were then unexpectedly asked to revisit them and provide justifications.

Method

Participants. Forty students at City University London participated for course credit or for a small payment.

Materials. Each booklet contained 42 target and 58 filler sentence pairs. Pairs consisted of a sentence “N are P” together with a second sentence “MN are P” in which M was one of the 3 possible modifiers used in Experiment 1, i.e. typical, atypical, and double atypical. The 40 sets of materials from Experiment 1 were supplemented by 2 more, to create a number that could be divided evenly over the 3 booklets. Target pairs were counterbalanced across the 3 booklets, and embedded in 58 filler sentence pairs. Fillers used knowledge effects to render the modified sentence as less likely for one third (e.g. “Malfunctioning white radiators are warm”), to make it more likely for another third (e.g. “Prison doors are made of metal”), and to leave the two sentences possibly equally likely (e.g. “Green feathered parrots are noisy”). The fillers thus ensured that there were opportunities to use all three response options.

Design and Procedure. The first and last page of each

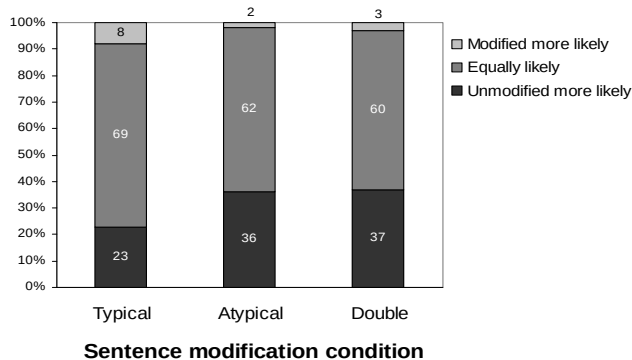
booklet contained instructions, and each page in between contained 4 sentence pairs. Participants circled one of 3 response options, printed to the right of each pair; 1) “the first sentence is more likely to be true”, 2) “the second sentence is more likely to be true”, and 3) “the two sentences are equally likely to be true”. At the end, participants were instructed to go back and justify why they answered in the way that they did, by writing a short statement next to each item. They were asked only to justify odd numbered sentences where they had stated that one of the sentences was more likely than the other (even numbered positions always contained fillers). A decision of “equally likely” was taken to be a default judgment, not requiring any further justification.

Results

Frequencies. Figure 1 shows the percentage of each response by condition. The bottom black bars represent the modifier effect – choosing the unmodified sentence as more likely. The central gray bars represent responses of “equally likely”. “Equally likely” was the most commonly chosen option for all three conditions (60 - 69% of responses).

When not equally likely, the unmodified sentence was selected as more likely 74% of the time for a typical modifier, 95% of the time for a single atypical modifier and 93% for a double atypical modifier, thus replicating the modifier effect and its interaction with modifier typicality.

Figure 1: Percent of responses for each condition



Since the 3 response proportions summed to 1, equally likely responses were omitted, and a 2-way ANOVA was run on response frequencies with factors of condition (3 levels of modifier), and response (2 levels: selecting the unmodified vs. selecting the modified as more likely). The significant main effect of condition ($\text{Min } F'(2, 153) = 3.74, p < .05$) corresponded to the fact that there were fewer preferences (in either direction) expressed in the typical modifier condition than in the other two, meaning that the “equally likely” response was significantly more frequent in the typical modifier condition (69%) than in the others (60-62%). The significant main effect of response ($\text{Min } F'(1, 55) = 41.0, p < .001$) confirmed that when a preference was expressed it was more often for unmodified sentences (32%) than for modified sentences (4%), and the significant

interaction ($\text{Min } F'(2, 158) = 8.73, p < .001$) confirmed the greater modifier effect seen in the atypical (34%) than the typical modifier conditions (15%).

Justifications. Justifications were provided for 87% of the requested cases. They were transcribed and classified by two independent judges. (Any given justification could be classified in more than one class.) Frequency by condition for cases where the unmodified sentence was preferred (as in the standard modifier effect) is shown in Table 2.

Table 2: Percent Justifications for preferring the unmodified sentence in Experiment 2

Justification	Typical	Atypical	2x Atypical
Pragmatic	58	42	40
Knowledge	8	20	23
Uncertainty	1	18	19
Other	33	20	18

The key to the classification is as follows:

Pragmatic. N was preferred as more general, while MN was seen as redundant. Example:

Flightless penguins live in cold climates

“All penguins live in cold climates and all penguins are flightless so to make a distinction is arbitrary, just say penguins live in cold climates”.

Pragmatic justifications, on the face of it, actually provided a reason for selecting the “equally likely” response – both flightless penguins and penguins in general live in cold climates. In fact, participants sometimes added “so I could also have said they were equally likely”. The unmodified sentence was chosen solely on the grounds of relevance. Pragmatic justifications were particularly common for typically modified pairs (58%), but were also very common for the other conditions (about 40%).

Knowledge. Knowledge of individuals in the modified noun category led people to doubt the truth of the MN sentence. Example:

Edible catfish have whiskers

“Edible catfish probably do not have whiskers still attached, as they could not be eaten like this”.

Knowledge-based justifications were the second most frequent type and indicated either that some of the combinations were not sufficiently novel, or that people had chosen to draw inferences from broader background knowledge. They were more common in the conditions using atypically modified sentences.

Uncertainty. General doubt was expressed about the modified sentence. Example:

Storage shacks are made of wood

“Shacks tend to be made of wood but storage shacks may not be”.

Justifications based on Uncertainty were of particular theoretical interest in that they could be interpreted as implying that people were not applying default inheritance to the properties of the modified concept, as suggested by CFGG. There were some 18-19% of these justifications in

the two atypical conditions, and hardly any in the typical condition. They were distributed evenly across most items.

Discussion

As in Experiment 1 there was a modifier effect which was greater for atypical modifiers. A striking difference from Experiment 1 was that when participants directly compared the relative likelihood of N and MN sentences, over 60% of the time they judged them equally likely – even when the modification involved two atypical modifiers. It is therefore by no means automatic that modifiers must affect the likelihood of properties. On most occasions people considered the property's likelihood to be unchanged, consistent with default property inheritance. When a preference was expressed however, it was almost always the modified sentence that had the lower likelihood.

The most common justification for selecting the unmodified sentence as more likely was on the basis of the pragmatic implications of uttering each statement, in line with Grice's (1975) maxims of cooperative communication. To utter the modified statement when one knew that the more general one was also true would be to violate Grice's maxim of quantity ("Be as informative as you can"). Participants were apparently sensitive to this kind of consideration when they judged the sentences for relative likelihood. What was striking here was that the pragmatic explanation was the most frequently offered justification for both typically and atypically modified sentences.

For the typically modified sentences, the pragmatic justifications more or less exhausted the cases showing the modifier effect. There was no evidence that properties of the modified concept differed substantially from those of the unmodified concept when the modifier was typical. For typically modified sentences, there was also a number (8%) of preferences expressed for the *modified* sentence as more likely – usually justified in terms of knowledge. These cases, working in the reverse direction from the normal modifier effect, explain the low correlation between the Typical modifier and the other conditions in Experiment 1.

For atypically modified sentences, two additional types of justification were given. First, even though the modifiers were chosen to be independent of the properties, around 20% of the justifications showed that participants had thought of ways in which they might be related. Around 20% of justifications also referred to uncertainty regarding the modified concept. This latter justification is consistent with the notion of a general reduction in judged likelihood applying to unfamiliar novel noun phrases.

To provide further evidence for default attribute inheritance, Experiment 3 aimed to show that people construct representations of modified concepts using information derived from the prototypes for the corresponding noun concepts (that is, that prototypes play an important role in judgments concerning complex linguistic expressions). Experiment 3 aimed to demonstrate that a structural dimension of the properties of a prototype, namely mutability, also affects judgments of likelihood of

the properties of modified concepts. Specifically, it was predicted that the more *immutable* a feature is for a concept, then the more likely it is to be true of modified versions of that concept.

Experiment 3

Mutability of a property for a concept relates to whether that property can be readily changed without further consequences for the concept (Sloman, Love & Ahn, 1998). For example, it is easy to imagine a world in which all ravens were white, but otherwise they were just the same. It is harder to imagine a world in which ravens had no wings, but were otherwise unchanged. Being black is thus more mutable for ravens than having wings. According to most accounts, mutability is based on a set of dependency relations defined over the properties in a concept's prototype. More dependencies lead to less mutable properties. If modified concepts do in fact inherit the default properties of the concept prototype, then we can expect that the mutability of properties would be the same for the modified concept as for the unmodified concept. In addition we expected that mutable properties would show a greater modifier effect, having less support from the dependency relations to other properties in the prototype.

Method

Participants. Seventy-two students at City University London participated for a lottery ticket for a small prize.

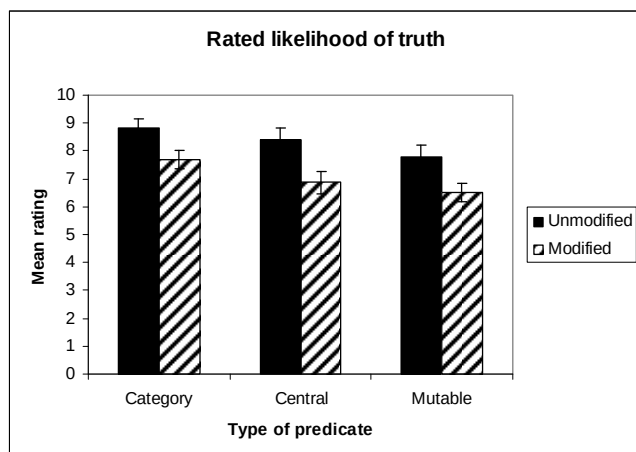
Materials and Design. Based on pre-testing, 3 properties for each of 33 concepts were selected, corresponding to three levels of mutability. Mutable sentences (e.g. lambs are white) were judged as easily imagined to be negated, while central sentences (e.g. lambs are warm-blooded) were harder to imagine as untrue. Categorical sentences (e.g. lambs are mammals) were chosen as the strongest level of immutability, on the assumption that moving a concept out of its superordinate class would affect many of its properties. Six sentences were constructed for each head noun, with the noun either being modified with an atypical modifier (e.g. inedible lambs) or left unmodified, and with the property being mutable, central or categorical. Each booklet contained 11 sentences at each level of mutability (mutable, central or categorical). Items were rotated across conditions, and each subject noun occurred only once in each booklet. The 33 target sentences in each booklet were randomly embedded in 33 filler sentences. Filler sentences were chosen so that some were clearly true, others clearly false, and others unclear. About half of the fillers had modified subject nouns.

Procedure. Each participant was given a booklet, comprising instructions and 3 pages of sentences. Each sentence was rated on a likelihood scale from 1 to 10.

Results and Discussion

The mean likelihood judgments for modified and unmodified sentences are shown in Figure 2.

Figure 2: Mean ratings of truth in Experiment 3



Comparing the filled and the striped bars, it is evident that there was a modifier effect in each condition (confirmed with $\text{Min } F'(1, 70) = 34.4, p < .001$). Comparing conditions, as expected the mutable sentences were considered less true than central, which were in turn less true than categorical ($\text{Min } F'(2, 140) = 38.9, p < .005$). Correlation between rated truth for a sentence when unmodified, and the same sentence when modified was .47, indicating further evidence for default inheritance of prototype structure in a modified concept. Interestingly, there was no interaction between the modifier effect and mutability ($F < 1$).

We had expected that if the modifier was casting doubt on a property, then the more secure people were in their belief about the property, the less effect would be shown by modification. Thus we expected people to be equally confident that lambs are white and that lambs are mammals, but we expected the atypically modified “inedible lambs” to be more certainly mammals than white. In the event, this interaction was not seen. (Two other experiments, not reported here, have confirmed this lack of interaction.)

General Discussion

What are the implications of the modifier effect for theories of prototype combination? CFGG argued that the modifier effect was evidence that people do not take the prototypical properties of the head noun concept as a default for the complex concept. This claim is clearly at odds with the findings presented here which showed very systematic relations between the head noun prototype and the properties judged likely to be true for the modified noun phrase. Experiment 1 replicated CFGG’s study but showed additionally that judgments of properties were correlated between the modified and unmodified versions of sentences. What is more or less likely for a noun concept is also more (or less) likely for the modified noun concept, as would be predicted if the modified noun inherits properties (and their relative likelihoods) from the unmodified noun prototype.

In order to get a clearer picture of the basis of the effects, Experiment 2 provided qualitative data on the ways in

which participants justify their judgments. The first important result to note from Experiment 2 was that the majority (60% or more) of pairs of sentences in all conditions were judged to be equally true. If one additionally were to discount those responses where the justification was either pragmatic or based on unintended knowledge-based reasoning, the number of judgments that showed a “pure” modifier effect was very small. For typical modifiers, 8% of responses showed a modifier effect that was not pragmatic or knowledge based. The equivalent figure for the atypical modifier conditions was 14%. We should therefore be wary of drawing strong conclusions about the failure of default inheritance in modified concepts.

In support of default inheritance, we showed that the degree to which likelihood judgments were affected by the modifier was influenced by information contained in the concept prototype in several ways. The studies showed that the more atypical the modifier was for the prototype then the stronger was the modifier effect. In addition the more central the property was for the prototype, then the more likely it was also to be considered true of the subset class.

Systematic patterns of attribute inheritance have been reported elsewhere. For example, Hampton (1987) demonstrated that in explicit conjunctions formed from relative clause constructions (birds that are also pets, sports that are also games), the judged importance of properties for the conjunctive phrase was predictable from their importance for each of the concepts separately. Our results here generalize this notion of “importance” to judgments of property likelihood.

In spite of the evidence for default inheritance, the modifier effect still requires some explanation. Across all the experiments, adding a modifier had a remarkably similar effect of generally reducing likelihood judgments. And as described above, there was still a residual 14% of cases in Experiment 2 where the effect with atypical modifiers was not explained in either pragmatic or knowledge-based terms. How should the effect be explained?

The clear evidence for the modifier effect across all experiments is inconsistent with a simple model in which properties are inherited by complex concepts with their likelihood unchanged. The best theoretical account of the data is probably therefore one in which attributes are inherited by default, but other factors can come into play that, overall, tend to reduce their judged likelihood.

The modifier effect and models of conceptual combination

Having established that the modifier effect follows systematic patterns consistent with the inheritance of prototypical properties, but subject to additional constraints, we turn in the final section to consider whether models of conceptual combination might account for the effect. That is, we reject the negative conclusions drawn by CFGG, (that people do not use default inheritance), in favor of an attempt to find an explanation for the systematic patterns of data that have been shown. Unfortunately, current models of

conceptual modification (e.g. Hampton, 1987, Smith et al, 1988) make no predictions about likelihood judgments. They make predictions about how modification affects the weight of different properties for judgments of typicality or category membership, but it is not obvious how such weights would translate into judgments of likelihood. While for both models, addition of a modifier will attract weight *away from* other properties in the prototype, this process is not sensitive to typicality, so that the models do not provide a good basis for explaining the effect.

Explaining the modifier effect. It is therefore necessary to look elsewhere for an explanation of the effect, and its dependence on typicality. We suggest three (speculative) possibilities, based respectively on similarity, familiarity and sample size. Developing a full account will need further experiments beyond the scope of the present study, designed to separate out these possible explanations.

One possibility would be to stipulate that confidence in all properties is reduced as a function of the similarity between the modified and unmodified concepts, borrowing the same general similarity principle that has been applied to explain inductive reasoning (Sloman, 1993). Atypical modifiers would generate complex concepts with less similarity to the original concept, and so the strength of the argument from concept to sub-concept would be correspondingly weaker. Calvillo and Revlin (2005) found that when properties are projected to an atypical subset, there is also generally less confidence in the concept being a subset of the concept. In other words, not only is it less likely that Giant Namibian zebras have stripes, but confidence in Giant Namibian zebras actually being a kind of zebra is also reduced, as a function of the atypicality of the modifier. The results of Experiment 3 appear to support this account.

A second possibility is that the *familiarity* of the modified concept determines the confidence with which the property is judged to be inherited. The moderation of the modifier effect by typicality can be readily explained this way. Typically modified concepts such as striped zebras are very familiar, while atypically modified concepts (Giant Namibian zebras) are not. Familiar modified concepts inevitably introduce knowledge effects, with consequent unpredictable changes in property likelihood (Rips, 1995). Typical modifiers may therefore generate no modifier effect (beyond pragmatic considerations) simply because they are already familiar concepts whose properties are known. It is not possible to create a typically modified concept that is not at least as familiar as the concept itself. In Experiment 2 we found that typical modifiers would also at times produce sentences that were *more* likely than the unmodified forms. Familiarity could also explain why the modifier effect appears to be insensitive to mutability. If the atypical modifier simply casts general doubt on the modified concept, then one's confidence in the inheritance of all properties may be diminished, regardless of whether they are mutable, central or categorical.

A third account of the typicality interaction would be to suppose that participants were employing intuitions of

sampling reliability. Suppose that the likelihood judgment is made by assessing the probability that any randomly selected member of the class will have the property, (for example the probability that a randomly sampled zebra will be fast). The modifier then restricts the broader class to a smaller subset (e.g. Giant Namibian zebras). One reasonable intuition might then be that although one's best guess is that Giant Namibian zebras should be just as likely to be fast as zebras in general, the likelihood of that guess being wrong is greater, the smaller the subset selected. Thus if there are very few Giant Namibian zebras it is more likely that Namibian zebras are not fast, than if most zebras are in fact Giant Namibian. Put another way, the larger the subclass, the more likely that the probability of finding the property in the subclass is just the same as the probability of finding it in the class itself, in the absence of any other knowledge.

In conclusion, the modifier effect presents a challenge to models of prototype combination. We have provided evidence that modified concepts do inherit prototypical properties from their constituent concepts, and in ways that are moderated by information within the concept prototype. In the absence of models of conceptual combination that speak directly to the question of property likelihood, three suggestions have been made of how the modifier effect and its interaction with typicality interaction may work – through a generalized similarity principle, through familiarity of typical modifiers introducing knowledge-based effects, or through reduced confidence when a numerically small subset is identified by the modifier. We suspect that more than one of these effects may in fact be at work across different examples.

References

- Calvillo, D. P. & Revlin, R. (2005). The role of similarity in deductive categorical inference. *Psychonomic Bulletin and Review*, 12, 938-944.
- Connolly, A. C., Fodor, J. A., Gleitman, L. R., & Gleitman, H. (2007). Why stereotypes don't even make good defaults. *Cognition*, 103, 1-22.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. L. Morgan (Eds.), *Syntax and Semantics* (3rd edition, pp. 41-58). New York: Academic Press.
- Hampton, J. A. (1987). Inheritance of Attributes in Natural Concept Conjunctions. *Memory & Cognition*, 15, 55-71.
- Jönsson, M. L. & Hampton, J. A. (2008). On prototypes as defaults. (Comment on Connolly, Fodor, Gleitman and Gleitman, 2007). *Cognition*, 106, 913-923.
- Rips, L. J. (1995). The current status of research on concept combination. *Mind and Language*, 10, 72-104.
- Sloman, S. A. (1993). Feature-based induction. *Cognitive Psychology*, 25, 231-280.
- Sloman, S. A., Love, B. C., & Ahn, W. (1998). Feature centrality and conceptual coherence, *Cognitive Science*, 22, 189-228.
- Smith, E. E., Osherson, D. N., Rips, L. J., & Keane, M. (1988). Combining Prototypes: A Selective Modification Model. *Cognitive Science*, 12, 485-527.