

Narrow Focusing: Why the Relative Position of a Good in Its Category Matters More Than It Should

France Leclerc, Christopher K. Hsee

Graduate School of Business, University of Chicago, 5807 S. Woodlawn Avenue, Chicago, Illinois 60637
{france.leclerc@gsb.uchicago.edu, chris.hsee@gsb.uchicago.edu}

Joseph C. Nunes

Marshall School of Business, University of Southern California, Accounting 304, Los Angeles, California 90089-0443,
jnunes@marshall.usc.edu

This research examines whether a low-ranking member in a high-status category (e.g., a low-end model of a high-end brand) or a high-ranking member in a low-status category (e.g., a high-end model of a low-end brand) is favored, holding the objective qualities of the items constant. Brand equity research suggests that the quality of a brand is more important than the ranking of a product within a brand. Our research documents a robust *ranking effect*—whereby a high-ranking product in a low-status category is favored over a low-ranking product in a high-status category even when information on competing categories is made available. We explain this effect in terms of narrow focusing and evaluability, and we identify boundary conditions of the effect.

Key words: brand choice; brand product management; ranking effect; narrow focusing; evaluability

History: This paper was received August 20, 2002, and was with the authors 17 months for 3 revisions; processed by Joel Huber.

Introduction

Imagine two automobiles of nearly identical quality were available on the market at approximately the same price. One is the lowest quality model offered by a prestigious brand; say, Audi. The other is the highest quality model of a less prestigious brand; say, Volkswagen. Which car would be evaluated more favorably? More generally, holding objective quality constant, would consumers favor a low-ranking member of a high-status category or a high-ranking member of a low-status category? The present research attempts to answer this question.

The central issue here entails how consumers use brand information, or more generally, category information to evaluate a specific good within a brand or category. There are two straightforward ways in which category information can be used. First, the good can be evaluated on the basis of the category to which it belongs. In the case where brand serves as the category, consumers could transfer the quality associations of the brand to all offerings under that brand. If we assume that Audi is generally perceived as a higher status brand than Volkswagen, and objectively the lowest quality Audi is identical in quality to the highest quality Volkswagen, we would expect the Audi to be evaluated more favorably than the Volkswagen. We refer to this as the *category effect*, because the option is evaluated based on the category to which it belongs. Alternatively, consumers could

use the category (brand) information as a frame of reference. In doing so, consumers would focus on the rank or position of the good within the category. If this were the case, we would expect the highest quality Volkswagen to be evaluated more favorably than the lowest quality Audi. We refer to this as a within-category ranking effect or simply *ranking effect* hereafter, because the option is evaluated based on its rank within its category.

We know from the extensive literature on brand equity that brand can play a significant role in the evaluation of a product (for a review, see Aaker 1996, Aaker and Joachimsthaler 2000, Keller 2003; see also Bell et al. 1999 for a discussion of the impact of brand on price promotions). Therefore, *ceteris paribus*, one might expect that an option introduced under a high-status brand (i.e., a brand with a mix of offerings that is composed of products/services that offer more attributes, of higher quality and at higher prices) would generally dominate an option introduced under a low-status brand, demonstrating what we call the category effect. We believe a category effect may prevail when two options from different categories are evaluated simultaneously. However, we show that the ranking of an item within a category (e.g., within a brand) has a major influence on its evaluation when items are evaluated independently. Consequently, a high-ranking model (i.e., a model that is near the upper boundary of the offerings in the

set) of a low-status brand may be evaluated more favorably than a low-ranking model of a high-status brand when the target items are evaluated separately. In other words, the ranking effect may prevail.

The rest of this paper is organized as follows. First, we briefly discuss the relevant literature before proposing our conceptual model, which explains when and how people are prone to the ranking effect. We then present five studies demonstrating the ranking effect, two of which test various boundary conditions. More specifically, in Study 1, we demonstrate that the ranking effect occurs when target items are evaluated independently. We show that consumers rely on the ranking of an item within a category even when information on all competing categories is made available. Study 2 provides further evidence that people are prone to use the immediate category as a frame of reference rather than the broader universe of like goods. In Studies 3 and 4, we identify boundary conditions for the ranking effect. We show that information format can facilitate the ranking effect, while encouraging people to engage in intercategory comparisons can dampen the ranking effect. Finally, Study 5 lends external validity, as we provide evidence for the ranking effect utilizing non-fictitious brands.

Literature Review

The decision problem described previously occurs in product categories in which we observe “vertical differentiation,” such that more is generally deemed as better by all consumers (Cremer and Thisse 1991). Such categories have variation in quality levels of products within the category (e.g., resolution for printers) or brand (e.g., BMW 3 series, 5 series, and 7 series). This type of product line structure is commonly found in consumer durables, although some packaged good manufacturers use brand modifiers to signal noticeable, but perhaps not dramatic, quality differences (e.g., Pampers Ultra Dry Thin diapers, Extra Strength Tylenol). This is in contrast to “horizontal differentiation,” where variation is typically in the function of the products (e.g., BMW roadster versus sport wagon) and differences are more a matter of individual preference. Much of the existing research on branding has focused on the evaluation of the horizontal structure of the product line and has been primarily concerned with brand extensions (Aaker and Keller 1990, Loken and Roedder John 1993, Boush and Loken 1991). An elemental finding from this literature is that brand associations carry over to an extension in a similar category, but do not carry over when brands extend their product lines into entirely different categories.

Research on the relationship between the vertical structure of the product line and brand equity

is scarce. What does exist focuses mainly on how the introduction of new items in a product line impacts the equity of a brand. For example, Loken and Roedder John (1993) find some support for the notion that consumers perceive a brand’s quality to be diminished if a low-quality product is added to the product line. Kirmani et al. (1999) show that consumer response to vertical line extensions is a function of brand image, stretch direction, and ownership status. Randall et al. (1998) have shown that a price premium (used as a proxy for brand equity) is significantly and positively correlated with the quality of the *lowest quality* model in a brand’s product line in lower quality segments of the market, while it is the quality of the *highest quality* model in the brand’s product line that matters for the upper quality segments of the market. Perhaps most directly related to our work was an experiment in which participants were presented with catalogs describing two fictitious brands of mountain bikes: a high-end and a low-end brand. Bicycles in the high-end brand ranged in price from \$759 to \$2,799, whereas bicycles in the low-end brand ranged in price from \$199 to \$959. Participants were then told that both companies were planning to introduce a new model of mountain bike with the same basic features priced at about \$800. When asked to choose between the two bikes, 63% of respondents preferred the product offered by the high-end brand. This result suggests a category effect, such that consumers evaluate competing products based on the desirability of the brand rather than on the ranking of the product within the brand.

Hypotheses and Experiments

The results reported by Randall et al. (1998) are based on respondents performing a direct comparison, evaluating the two bikes in joint evaluation mode, such that multiple options are presented side by side and evaluated simultaneously. What if the two products (i.e., bikes) were evaluated independently or by different consumers? Would consumers still base their evaluation on the desirability of the category (i.e., brand)?

Recent work reveals systematic inconsistencies between evaluation modes; that is, between *joint evaluation* and *separate evaluation* (see Hsee et al. 1999 for a review). Within both evaluation modes it is assumed that when evaluating an option, people compare it to a referent. However, one of the main differences between the two evaluation modes is what referent people use (Hsee and Leclerc 1998). Under joint evaluation mode, people make their evaluations by comparing one option to another. In the Randall et al. (1998) experiment, it is likely that people simply compared the two \$800 bicycles. Given the two options appeared equally attractive on an absolute basis, the

category to which each belongs (brand) was likely to have been used to distinguish between the two, especially as it was the only significant difference between the two options.

Under separate evaluation, options are evaluated independently. Because no alternative is provided explicitly, the option to be evaluated must be compared to some other referent. Normatively, one would expect the search for a referent to be quite broad, such that ideally it would represent the entire universe of goods, such as the distribution of bicycles available or the average bicycle. However, we propose that when evaluating a single good belonging to a category (e.g., a brand), by default, consumers use the immediate category (brand) as their frame of reference relying principally on the good's position (ranking) within the category. They do so unless explicitly instructed or primed by the context to do otherwise. Consequently, a high-ranking member in a low-status category will be evaluated more favorably than a low-ranking member in a high-status category, even though objectively the latter product is equal or superior to the former. This is what we refer to as the "ranking effect."

The ranking effect may occur because of attribute evaluability (Hsee 1996, Hsee et al. 1999). An attribute is said to be easy (difficult) to evaluate if it has (does not have) well-defined reference information, such as range, distribution, and so on. When people evaluate an item in isolation (separate evaluation), they often know only about the category to which the item belongs and are not told about alternative categories. In this type of situation, the ranking of a product within the category is typically easier to evaluate than the actual value of the product or the desirability of the category. This is because there is a clear referent by which to evaluate the product, its position within the category, while there is no clear referent by which to judge the actual value of the good or desirability of the category. This leads to our first hypothesis:

HYPOTHESIS 1. *In separate evaluation mode, people will exhibit the ranking effect when information about other categories is not available.*

There is some evidence for such a ranking effect. In a study conducted by Hsee (1998), a relatively expensive gift in a relatively inexpensive product category was judged more favorably than a relatively inexpensive gift in a relatively expensive product category. Specifically, respondents were asked to judge the generosity of one of two gift givers. One gave a \$45 scarf, a relatively expensive member of a relatively inexpensive category (scarves). The other gave a \$55 wool coat, a relatively inexpensive member of a relatively expensive category (coats). When evaluated separately, the giver of the \$45 scarf was considered

more generous, which is consistent with the evaluability explanation. The ranking of the gift within its category is easier to evaluate independently than the desirability of the category or the actual value of the gift. Similar ranking effects have been reported by Kahneman and Ritov (1994) and Kahneman et al. (2000) in the context of willingness to contribute for public goods. For example, Kahneman et al. (2000) found that in separate evaluation, cyanide fishing in coral reefs (a high-ranking environmental issue) is considered more serious than multiple myeloma among elderly (a low-ranking human health issue).

There are certainly many real-world situations in which people only have information about the category of the good to be evaluated and do not have information about other categories, such as "coats" in the gift study. Another example is when a prospective car buyer visits a car dealership that only carries the brand of cars in which the buyer is interested. In these cases, attribute evaluability is a compelling explanation for the ranking effect as there is no clear referent available other than the category. However, there are also many real-world situations in which information about other categories is readily available. A visit to a car dealership that carries different brands of cars is an example of such situations. Would the ranking effect still occur in the latter situation? If so, then attribute evaluability alone is not sufficient to account for the ranking effect and we need another explanation.

We propose a novel explanation for why the ranking effect may occur in separate evaluation, even when information about alternative categories is available. We call our explanation "narrow focusing"; that is, even if information about items in other categories (broad information) is available, people tend to focus on the category to which the target belongs (narrow information). They act as if information about the other categories did not exist. Our position is consistent with Kahneman and Miller's (1986) norm theory that suggests that when evaluating an object in isolation, people often compare it only with alternatives in the same category. A similar process has also been documented in the social judgment literature. Biernat et al. (1991) conducted a study in which respondents rated the height of different students who were shown in full-length photographs. Despite explicit instructions that stressed a constant judgmental framework, results suggested that the male targets were inadvertently rated in comparison with other men and female targets were compared with other women.

The fact that people focus on a subset of the relevant information has been documented in various streams of research involving diverse behaviors, including

choice bracketing (Read et al. 1999, Read and Loewenstein 1995); the money illusion effect (Shafir et al. 1997); the medium effect (Hsee et al. 2003, van Osse-laer et al. 2004); and the effect of local set consideration on global choice (Simonson et al. 1993, Simonson and Tversky 1992). For example, Simonson et al. (1993) found that people are more likely to choose the lowest quality/lowest price option from a global, trinary choice set (A, B, and C) if they are first asked to choose an option in each of three local binary choice sets, (A, B), (B, C), and (C, A), than if they are directly asked to choose an option from the global trinary set. Presumably, when people choose from a local set, they focus on the local set alone and do not consider the option not in the local set.

The present research extends this literature by demonstrating how people are predisposed to use only the local category to evaluate a product even if other category information is available, and by demonstrating marketing-relevant consequences of such narrow focusing. Thus, we propose the following hypothesis.

HYPOTHESIS 2. *In separate evaluation mode, people will exhibit the ranking effect even when information about other categories is available.*

The two hypotheses posited thus far only concern separate evaluation situations. What would happen in a joint evaluation situation in which the same people are asked to compare and evaluate both a low-ranking member in a high-status category and a high-ranking member in a low-status category? Purposefully directing people to compare across categories would be expected to attenuate or eliminate the ranking effect. This leads us to our third hypothesis.

HYPOTHESIS 3. *People will be less likely to exhibit the ranking effect in joint evaluation mode than in separate evaluation mode.*

Study 1 is designed to test Hypotheses 1–3. It demonstrates that, when performing a separate evaluation, the rank of the item within the category plays a more prominent role in the evaluation, even when category information involving the universe of goods is available. However, we show that while performing a joint evaluation, the two items will serve as referents for one another.

Study 1

Method

Participants were 175 students from a Midwestern university who were told that a friend (or friends) had given them a music dictionary. They were also told that there were four publishers (Feine, Chime, Likert, and Huine) and each had recently released 12 music dictionaries, some more comprehensive

(more entries) than others. In this study, the publisher served as the category. The task was to evaluate two specific music dictionaries: (1) the 76,000-entry (the second most comprehensive) dictionary published by Huine and (2) the 78,000-entry (the eleventh most comprehensive) dictionary published by Chime (see Appendix A). It is important to note that the Huine dictionary is a high-ranking member in a low-status category, while the Chime dictionary is a low-ranking member in a high-status category, and objectively, the low-ranking dictionary is more comprehensive than the high-ranking one. We should also point out that neither dictionary is the best (highest ranking) or worst (lowest ranking) in its category (brand).

Participants were randomly assigned to one of five conditions: one joint evaluation condition, two separate evaluations without information about other categories (publishers), and two separate evaluation conditions with information about other categories (publishers). In the joint evaluation condition, respondents were asked to imagine that two friends each gave them a music dictionary. They were presented with information about all of the publishers and their dictionaries as in Appendix A and were asked to evaluate *both* target dictionaries. To ensure respondents recognized which dictionaries to evaluate, we first asked them to identify the target dictionaries by circling them in the tables. Twenty-six respondents did not correctly identify the target(s) in the table and their responses were excluded. Participants then were asked to indicate their happiness with the two dictionaries on a scale ranging from 1 (very unhappy) to 7 (very happy).

The two separate evaluation conditions where respondents did not get other publishers' information were similar to the joint evaluation condition except for the following. The respondents were told about only one friend and evaluated only one dictionary, either the 76,000-entry Huine dictionary or the 78,000-entry Chime dictionary. They were shown only the table for the publisher of the target dictionary (Chime or Huine) and evaluated only that dictionary. The separate evaluation conditions with information were identical, except that respondents were shown the tables of all the publishers as shown in Appendix A.

Our three primary predictions can be summarized as follows. According to Hypothesis 1, we predict that those who received the 76,000-entry dictionary, the second most comprehensive dictionary from Huine, would be happier than those who received the 78,000-entry dictionary, the eleventh most comprehensive dictionary from Chime, despite the fact that the second is objectively superior. First, this would occur in the separate evaluation conditions without information about other categories (publishers). Second, in line with Hypothesis 2, we predict that even in

the separate evaluation conditions with information about other publishers, the ranking effect would still exist, albeit perhaps in a weaker form. This is the ranking effect. Finally, corresponding to Hypothesis 3, we predict that in the joint evaluation condition, this ranking effect would be weakened or eliminated.

Results and Discussion

Our results, summarized in Figure 1, are consistent with all three predictions. First, in the two separate evaluation conditions *without* information on other publishers, people receiving the 76,000-entry Huine dictionary reported being happier than people receiving the 78,000-entry Chime dictionary, even though, objectively, the latter dictionary dominates the former. This difference is statistically significant ($\mu_{\text{Huine}} = 6.11$ versus $\mu_{\text{Chime}} = 3.96$, $t(51) = 6.8$, one-tailed $p < 0.001$) and illustrates the ranking effect. It is a replication of the gift study mentioned earlier (Hsee 1998) and supports the evaluability hypothesis.

Second, in the two separate evaluation conditions *with* information on other publishers, people who received the less comprehensive, yet second ranked Huine dictionary still reported being happier than people who received the more comprehensive, eleventh-ranked Chime dictionary and the difference is statistically significant ($\mu_{\text{Huine}} = 4.95$ versus $\mu_{\text{Chime}} = 4.34$, $t(70) = 2.10$, one-tailed $p < 0.05$). This ranking effect supports our narrow focusing argument, suggesting that even when information about other categories is available, people still use the “local” category as their frame of reference to evaluate the target.

Finally, in the joint evaluation condition, the ranking effect not only dissipated, but actually reversed. Respondents reported being happier with the eleventh-ranked Chime dictionary than with the second-ranked Huine dictionary and the difference is statistically significant ($\mu_{\text{Huine}} = 5.00$ versus $\mu_{\text{Chime}} = 5.54$, paired $t(24) = 2.6$, one-tailed $p < 0.01$). Presumably, in the joint evaluation conditions, respondents

compared the two dictionaries directly and realized that the Huine dictionary was less comprehensive than the Chime dictionary after all.

In summary, the results support our predictions (Hypotheses 1–3) entirely. When performing a joint evaluation between two items differing in comprehensiveness, the two items appear to serve as referents for one another. However, when performing a separate evaluation, the rank of the item within the category plays a more prominent role in the evaluation. When this occurs and the only information provided is the information about the category to which the item belongs, the ranking effect can be explained by attribute evaluability. More importantly, even with full information that allowed respondents to compare the category to which the item belongs to other categories, we still observe the ranking effect. This supports our proposed explanation—narrow focusing.¹ Furthermore, the target items are neither at the top nor at the bottom of their category (either extreme), and the ranking effect persists. Thus, extremity is not a necessary condition for our findings.

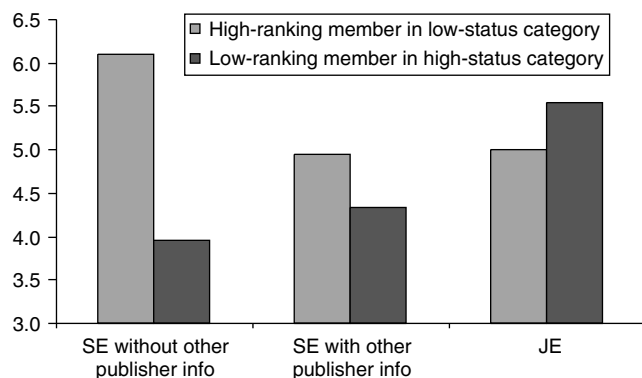
Our next study provides a more direct test of the narrow focusing idea. According to narrow focusing, people tend to use a narrow reference (e.g., a local category) rather than a broader reference (e.g., other categories) to make their judgment, even though the latter information is available. This leads us to our Hypothesis 4.

¹ One could argue that the ranking effect occurred in the separate evaluation conditions even when information on other publishers was provided, simply because participants did not pay enough attention to the information. To rule out lack of motivation as an explanation, participants completed a need-for-cognition (NFC) scale, an individual difference measure reflecting one's tendency to engage in and enjoy effortful cognitive processing (Cacioppo and Petty 1982). Considerable evidence suggests that people high in NFC exert more effort on cognitive tasks than do people low in NFC (see Cacioppo et al. 1996 for a review).

Participants in this study were categorized as either “high” or “low” in NFC based on a median split on the scale. A 2 (Rank: low/high) $\times$ 2 (NFC: low/high) ANOVA was performed for each of the three types of evaluation. The main effect for Rank was significant for all three types (all $ps < 0.01$), whereas a main effect for NFC was never present. The interaction between Rank and NFC reached significance in the joint evaluation condition ($F_{1,44} = 4.38$, $p < 0.05$): the difference between the high- and low-rank options was significant for individuals who scored high on NFC ($t_{1,27} = 2.80$, $p < 0.05$), but not for individuals who scored low on NFC ($t_{1,19} = 0.27$, $p > 0.5$). This suggests that as individuals with a greater need for cognition put more effort into the evaluation task, they were more likely to notice the difference between the two options. This provides evidence that the scale indeed identified participants who exerted more cognitive effort in the task.

However, in the separate evaluation conditions, the interaction did not reach significance. Participants with a high NFC were not significantly less likely to evaluate an item based on its position within a category, suggesting that the ranking effect is not a function of insufficient motivation to process information.

Figure 1 The Effect of Item Ranking and Category Status Under Joint (JE) and Separate (SE) Evaluation Mode



HYPOTHESIS 4. *When evaluating an item, given the option between information about the immediate category to which it belongs (narrow referent) and broader, normatively more relevant information (broad referent), people are more likely to seek out and utilize information about the immediate category.*

Study 2 tests this hypothesis directly.

Study 2

Method

Participants were 102 undergraduate students from a West Coast university. They were given a questionnaire describing a situation in which they, as an admissions officer at a college in the United States, must evaluate the test scores of a student from a foreign country. More specifically, respondents read the following:

Imagine the following scenario. You work in the admissions office of a college in the United States. You have received an application from an African country. The applicant, whose first name is Ann, has not taken any standard tests you are familiar with, but has taken a standard scholastic aptitude test administered by the government of her country to all high school seniors in that country. The test is called SBT. Ann's SBT score is 225. To assess how good her score is, you have obtained the following pieces of information: The name of Ann's high school is Yaas; it is one of 500 high schools in Ann's country. There are 800 students, including 200 seniors, at Yaas. Ann is a senior at Yaas and will graduate in a few months.

You could find out one of the following additional pieces of information:

- (1) The average SBT score of all seniors at an average quality high school in Ann's country.
- (2) The SBT score of Ann's friend, Mary.
- (3) The average SBT score of all seniors at Yaas, Ann's high school.

If you could obtain only one of the above pieces of information, which one would you obtain? Circle one above.

Of the three additional pieces of information, the first one is the grand average of all the students in that country, and is the broadest and normatively speaking, the most diagnostic of the three. The second piece of information is merely filler and is least diagnostic. The last piece of information is the average of the narrow category—the applicant's own high school, and is therefore less diagnostic than the first. Nevertheless, narrow focusing suggests that a significant proportion of people would select the third choice, which provides narrower information than the first choice, which provides information from a broader set (i.e., the universe of seniors in that country).

Results and Discussion

As predicted, 74% of respondents opted for the narrow information (this proportion is significantly different than 50%; $\chi^2 = 22.59$, $p < 0.01$). The remaining 26% of respondents chose the broader information. None chose the filler. This finding supports Hypothesis 4 directly and reinforces our belief that the ranking effect in the separate evaluation conditions when category information is provided (e.g., Study 1) arises from narrow focusing.

In the following two studies, we explore boundary conditions of the ranking effect. As already proposed in Hypothesis 1 and tested in Study 1, one boundary condition is evaluation mode; the ranking effect is less likely to emerge in joint evaluation than in separate evaluation. This is because joint evaluation facilitates cross-category comparisons. Generally speaking, the ranking effect is more likely to occur if the local category information is salient (e.g., Abele and Petzold 1998), and cross-category comparison is not facilitated. More specifically, we offer the following additional boundary-condition hypotheses.

HYPOTHESIS 5. *The ranking effect will be more likely to occur if the local category is salient than if the local category is not salient.*

HYPOTHESIS 6. *The ranking effect will be less likely to occur if intercategory comparisons are encouraged.*

These two hypotheses are tested in Studies 3 and 4, respectively.

Study 3

Method

Respondents were 225 students from a Midwestern university who were provided information on the comprehensiveness of 16 music dictionaries as well as the names of their publishers. They were told that the list was inclusive, encompassing all of the music dictionaries available on the market. They were instructed to imagine that they had been given one of these dictionaries as a birthday gift and asked to indicate their happiness with the gift on a scale ranging from 1 (very unhappy) to 9 (very happy).

The study utilized a 2 (target: low-ranking member in a high-status category versus high-ranking member in a low-status category) $\times$ 2 (category: salient versus nonsalient) factorial design. As in Study 1, publishers serve as categories. The two target dictionaries included one published by Wilson & Son (55,000 entries) and one published by Elton Hills (56,000 entries). See Appendix B for the entire list of dictionaries. Note that the Wilson & Son dictionary is a high-ranking member in a generally low-status category, while the Elton Hills dictionary is a low-ranking member in a generally high-status category, and objectively the latter is better than the former.

Our second independent manipulation is salience of category information. In the salient category condition, the books were first organized by publishers and then by their comprehensiveness (see Appendix B). In the nonsalient category condition, the books were just sorted by their comprehensiveness without first being categorized by their publishers (see Appendix C). Notice that the information provided in the two conditions is identical and the only difference is that in the first case, publishers are presented in such a way that they serve as natural categories, while in the second case they are not.

Following from Hypotheses 2 and 5, our predictions were twofold. First, in the salient category condition, we expected the most comprehensive Wilson & Son dictionary would generate greater happiness than the least comprehensive Elton Hills dictionary, even though the latter was objectively more comprehensive. Conversely, in the nonsalient category condition, we predicted that this ranking effect would disappear.

Results and Discussion

In the salient category condition, the 55,000-entry most comprehensive Wilson & Son dictionary indeed evoked greater happiness than the 56,000-entry least comprehensive Elton Hills dictionary ($\mu_{55,000} = 5.18$ and $\mu_{56,000} = 4.02$, $p < 0.05$). This result supports our first prediction and replicates the ranking effect reported in Study 1. Moreover, in the nonsalient category condition, the ranking effect disappeared ($\mu_{55,000} = 4.81$ and $\mu_{56,000} = 5.04$, $p > 0.50$). A 2 (category) $\times$ 2 (target) analysis of variance (ANOVA) reveals the predicted two-way interaction ($F_{1,222} = 7.26$, $p < 0.01$). The results suggest that *ceteris paribus*, a manipulation of information display can make a category more or less salient, thereby turning the ranking effect on or off. These findings support many people's lay intuition that anyone can claim to be the best in some strategically selected category. For example, a good basketball team may claim that it is the best in the west conference, the best in its state, the best in its city, or taken to the extreme, the best among teams with a coach whose last name starts with a "Z."

Study 3 provides further evidence for the proposition that when items are partitioned into categories, people will use the category as a frame of reference within which to evaluate the target item. In other words, by sorting the music dictionaries into what may be specious classifications (e.g., imaginary publishers), we have prompted respondents to judge an individual dictionary by its rank within the category and not its overall position among the universe of music dictionaries. Furthermore, by doing so, participants evaluated an objectively less valuable product more favorably than an objectively more valuable good.

Study 4

We seek to achieve two objectives in this study. First, we attempt to replicate the ranking effect using purchase likelihood rather than happiness as the dependent variable. Second, we test another boundary condition of the ranking effect, as proposed in Hypothesis 6, concerning the type of comparison consumers make. We predict that the ranking effect is less likely to occur if an intercategory comparison is encouraged.

Method

The respondents were 176 students from a Midwestern university who were given information on 25 fictitious personal digital assistants (PDAs), as presented in Appendix D. The respondents' task was to indicate the likelihood that they would buy a specified model by indicating a number on a scale ranging from 1 (definitely not) to 9 (definitely yes).

The study utilized a 2 (target: low-ranking member in a high-status category and high-ranking member in a low-status category) $\times$ 3 (type of comparison: intracategory, intercategory, and control) between-subjects design. The target was manipulated by varying whether respondents were asked to indicate their purchase likelihood for the worst model of a higher quality brand or the best model of a lower quality brand. Both models had 35 MB of memory, the only objective attribute available. The type of comparison was manipulated by instructing the respondents to think about the position of the target model within its brand, the position of the target brand among other brands, or neither. More specifically, before being asked their likelihood of purchase, students in the intracategory comparison condition read the following:

"Compared with the other models within Tversky (i.e., Tracy, Trish, Ted, and Terry), how good is the Tom Model?"

Those in the intercategory comparison condition read the following instead:

"Compared with the other brands (i.e., Darley, Skinner, Bandura, and Cooper), how good are Tversky Brand PDAs *on average*?"

Those in the control condition received neither set of instructions.

We had two predictions for this study. First, by default, people would use the local category as the frame of reference. Therefore, in the control condition, where neither intracategory nor intercategory comparison is encouraged, people would exhibit a ranking effect as in the intracategory comparison. Second, the ranking effect would be attenuated in the intercategory comparison condition.

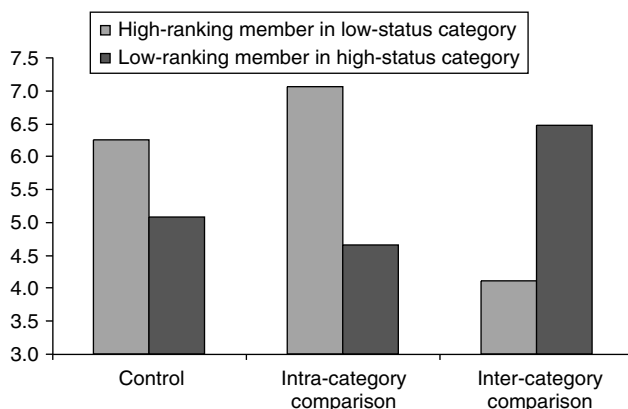
Results and Discussion

The results, summarized in Figure 2, support all of our predictions. As expected, a 2 (target) $\times$ 3 (type of comparison) ANOVA revealed a significant interaction effect ($F_{2,171} = 17.68$, $p < 0.001$). Neither main effect was significant. Subsequent analyses focus on the three type of comparison conditions separately. We examine the control condition first. In that condition, the respondents were significantly more likely to purchase a high-quality good from a lower quality brand (Sean of Skinner) than the low-quality good from a higher quality brand (Tom of Tversky) ($\mu_{\text{Tom}} = 5.07$ and $\mu_{\text{Sean}} = 6.24$, $t_{56} = 2.15$, $p < 0.05$). This result replicates the ranking effect found in the other studies for happiness and extends the effect to purchase likelihood. In the intracategory comparison condition, the respondents exhibited a ranking effect as well ($\mu_{\text{Tom}} = 4.65$ and $\mu_{\text{Sean}} = 7.06$, $t_{57} = 3.90$, $p < 0.001$).

Finally, in the intercategory comparison condition, the ranking effect was not only absent, but respondents' preferences reversed: Respondents were more likely to buy the low-quality product from the higher quality brand (Tom of Tversky) than to buy the high-quality product from the lower quality brand (Sean of Skinner) ($\mu_{\text{Tom}} = 6.47$ and $\mu_{\text{Sean}} = 4.1$, $t_{58} = 3.87$, $p < 0.001$). This is a category effect.

This study yields two important conclusions. First, when the focus of attention is not manipulated, people are inclined to pay more attention to the ranking of the product within its category rather than to the desirability of the category. Second, it is possible to turn off, in fact, even reverse the ranking effect by steering the respondents' focus toward the relative desirability of the category. The ability for marketers to steer consumers toward making category-level or within-category assessments suggests numerous practical implications, some of which we will discuss later on.

Figure 2 The Effect of Item Ranking and Category Status When Comparisons Are Made Intracategory and Intercategory



Study 5

In the studies reported thus far, the categories have all been fictitious. The advantage of using fictitious categories is that they do not carry extraexperimental meanings and can be manipulated cleanly. In the real world, on the other hand, categories are often entities that people already know, such as an existing brand, an existing publisher, or an existing university. In this study, we sought to replicate the ranking effect, as well as the joint-separate reversal effect, with such naturally occurring and familiar categories: car brands.

Specifically, we asked consumers to evaluate one of two automobiles presented along with a variety of models from the same brand either separately or jointly. The study illustrates how evaluations are based on the car's rank within the brand, and reveals how one brand's best model can be evaluated more favorably than a competing, more prestigious brand's worst model, even when the two models are objectively comparable.

Method

Participants were 92 students from a Midwestern university who were instructed to imagine that they were in the market for a new car. All respondents were also told that a key factor that distinguishes an engine's performance is horsepower and that horsepower can range from 100 to 300. Participants were allocated to one of three experimental conditions. In the joint evaluation condition, participants were told that they were at a dealership that carries two brands, Audi and Volkswagen (VW), and four models under each brand. They were given a list of the four models available as well as information on the horsepower of the engine for each model. The horsepower for the Audi cars varied from 190 to 300, whereas for the VW cars, it varied from 100 to 190. This is consistent with the fact that for years Audi has been billed around the globe as Volkswagen's up-market division, the maker of the company's "faster, swankier vehicles" (Frank 2002). The respondents were told that they were interested in two models, the Audi 4 (the lowest model in the higher status category) and the VW Passat V6 (the highest ranked model in the lower status category). They were also told that both cars retail at about \$30,000. They were then asked how much they would be willing to pay for each one of the two cars.

The two remaining experimental conditions were two separate evaluation conditions. In one of these two conditions, respondents were told that they were currently at an Audi dealer. They were given a list of Audi cars (a relatively high-ranking category) and were told that they were interested in the model at the bottom of the list (the lowest ranking model), the Audi A4. In the other separate evaluation condition, participants were told that they were at the VW

dealer. They were given a list of VW cars (a relatively low-ranking category) and were told that they were considering the model at the top of the list (the highest ranking model), the Passat V6.

We predicted that (1) in the separate evaluation conditions, participants would exhibit a ranking effect, favoring the best VW over the worst Audi and (2) in the joint evaluation condition, their preference would be reversed.

Results and Discussion

As expected, in the separate evaluation conditions respondents were willing to pay significantly more for the best VW than for the worst Audi ($\mu_{\text{VW}} = \$24,855$ versus $\mu_{\text{Audi}} = \$22,879$, $t(59) = 1.46$, one-tailed $p < 0.10$). In other words, participants based their willingness to pay more on the position of a model within its brand than on the quality of the brand. This result replicates the ranking effect when brands are categories, and is in stark contrast with the work of Randall et al. (1998) showing the primacy of brands.

On the other hand, in the joint evaluation condition, when the best VW and the worst Audi are juxtaposed, respondents were willing to pay more for the worst Audi than the best VW ($\mu_{\text{VW}} = \$23,935$ versus $\mu_{\text{Audi}} = \$24,742$, paired $t(31) = 1.60$, one-tailed $p < 0.10$). This effect replicates the work of Randall et al. (1998), and supports our hypothesis on joint-separate evaluation reversals. Furthermore, we obtain these effects even though people are told the range of the focal attribute. Together, our results suggest that in separate evaluations, people use brand as a frame of reference and rely on the rank of a model within the brand to evaluate the target. However, in joint evaluations, people may use brands as a cue for the desirability of a product.

General Discussion

In this paper, we have shown that in separate evaluations, people are predisposed to use category information (how the item ranks within the category) as a frame of reference to evaluate the quality of that item. People do so even if full information is provided, information that would allow them to evaluate the desirability of the category or the desirability of the item based on the universe of items. As a result, people may evaluate a high-status member of a low-status category better than a low-status member of a high-status category, holding the objective quality of the items constant. We refer to this as the “ranking effect.” Furthermore, as demonstrated in the first and third studies, they may even do so when the quality of the high-status member in a low-status category is objectively worse than the quality of the low-status member of a high-status category.

We have proposed narrow focusing as an explanation for the ranking effect under full information. This notion suggests that when asked to evaluate an item,

people naturally use a narrow referent such as the category to which the item belongs. We provide support for this view by showing that, when asked which information they would like to evaluate an item, people request a narrow as opposed to broad referent. Furthermore, we have proposed and found empirical support for boundary conditions of the effect; the ranking effect is less likely to occur if (1) the local category information is not salient and (2) the items are evaluated under a joint evaluation mode, or more generally if cross-category comparison is encouraged. Thus, we believe that any situation that explicitly calls for a comparison between categories (such as a joint evaluation or a choice task) would moderate the ranking effect, and may accentuate the category effect. For instance, in the Randall et al. (1998) experiment previously described, asking people to choose between two bikes from two different companies (rather than choose from the entire set) focused their attention on differences between the companies, and thereby accentuated the category effect.

As mentioned earlier, there is a second mechanism that can account for the ranking effect; namely, attribute evaluability. We believe attribute evaluability and narrow focusing are complementary explanations, and may, at times, both be contributory factors. Both assume that people compare the target to some referent (e.g., Kahneman and Tversky 1979, Tversky and Kahneman 1981), and both posit that the ranking effect is generally greater than the category effect in separate evaluation. The two explanations differ, however, in the situations to which they apply. In situations where there is no well-defined reference to evaluate the desirability of the local category, evaluability is sufficient to explain the ranking effect. In situations where there is alternative category information with which to evaluate the desirability of the local category, evaluability alone is insufficient to explain the ranking effect, and narrow focusing is necessary. Narrow focusing indicates that even in the latter situations, people still use only (or primarily) the local category as their frame of reference, as if the alternative category information were unavailable.

Work by Abele and Petzold (1998) suggests another moderating factor for the use of within-category comparison, which would result in the ranking effect. These authors suggest that display format—whether the information is organized by category or whether it is presented under a mixed format—serves as a metainformational cue. They argue that with a blocked presentation (items organized by category) of the information, participants will infer that their primary task is to differentiate within category, whereas when targets belonging to two different categories are presented in a mixed presentation mode, participants will infer that their primary task is to differentiate

between the categories. We certainly agree that the format of the information can impact the likelihood of within-category comparison. However, metainformational cues provided by the format of the information do not seem to be the only mechanism at play in our research. For one thing, information format is only manipulated in one of our studies (Study 4). In Studies 1, 3, and 5, the information format is kept constant. It also seems unlikely that participants can infer metainformational cues from our manipulations when we test for boundary conditions. For instance, in the joint evaluation mode in Studies 1 and 5, it seems unlikely that people infer that when they are asked to evaluate two items, they have to compare the items with each other as opposed to with any other in the display. Finally, the fact that in Study 2, participants ask for information about the category rather than the universe suggests that the ranking effect is likely to be caused by individuals relying on something more basic than inferential processes.

To conclude, in this paper, we have shown that when evaluating an item under separate evaluation, people will use the category that the item belongs as a frame of reference. They do so because they naturally focus their attention on only a subset of the relevant information. An interesting question beyond the scope of this paper is why people use a narrow referent in such an evaluation task. One possibility is that when multiple pieces of information are available providing many potential frames of reference, people will select a frame of reference that they perceive as distinctive and relevant (see Helson 1964 and Stapel et al. 1997 for a similar argument). In the context of the present research, one can speculate that a specific brand (or more generally, the category that a target item belongs

to) constitutes a distinct and separate entity with relatively clear boundaries.

One limitation of this work is that we utilize only those categories we believe are well defined and in which the worst item in the better category is at least as good, if not better than the best category in the lower ranked category. Perhaps categories with overlap (i.e., where the best of the worst is actually as good as the midpoint in the better category) would exacerbate the effect. It would be interesting to investigate how different degrees of overlap would affect evaluations.

Finally, managers should concern themselves with the results of this research. For virtually any given product, no matter how good or bad it is in the universe of products, one can always find a category in which it is the best member. For example, Amstel Light calls itself the best beer in its class, defining its category as low-calorie imports, of which only Kirin Light comes to mind as a competitor. The maker of a mediocre resolution digital camera could honestly call it “the highest resolution camera under 0.5 pound,” if it were true, or “the highest resolution camera that uses AA batteries.” Or, as many sellers do, they may simply call the product “the highest resolution camera in its class” and leave it to the reader to determine in which class the product falls.

Acknowledgments

The first author thanks the Kilts Center for Marketing of the University of Chicago for research support. The second author thanks the University of Chicago Graduate School of Business for summer support and China Europe International Business School, where he visited during part of this project, for research support.

Appendix A. The Stimuli Used in Study 1

Dictionary Questionnaire

Suppose that you major in musicology. On your birthday, a friend gave you a music dictionary. The table below lists all the music dictionaries recently published by publishers, Feine, Chime, Likert, and Huine. As you can see, some dictionaries are more comprehensive (have more entries) than others. The dictionary your friend gave you is the eleventh most comprehensive dictionary published by Chime Ltd. First, identify the dictionary your friend gave you by circling it in the table below.

Publisher: Feine Ltd.		Publisher: Chime Ltd.		Publisher: Likert Ltd.		Publisher: Huine Ltd.	
Dictionary	Entries	Dictionary	Entries	Dictionary	Entries	Dictionary	Entries
1	99,000	1	98,000	1	97,000	1	78,000
2	96,000	2	96,000	2	93,000	2	76,000
3	87,000	3	94,000	3	90,000	3	74,000
4	83,000	4	92,000	4	89,000	4	72,000
5	79,000	5	90,000	5	85,000	5	70,000
6	76,000	6	88,000	6	82,000	6	68,000
7	73,000	7	86,000	7	77,000	7	66,000
8	69,000	8	84,000	8	71,000	8	64,000
9	65,000	9	82,000	9	67,000	9	62,000
10	61,000	10	80,000	10	63,000	10	60,000
11	58,000	11	78,000	11	57,000	11	58,000
12	55,000	12	76,000	12	54,000	12	56,000

Second, circle a number in the following scale to indicate how happy you are with this dictionary.

1-----2-----3-----4-----5-----6-----7
very unhappy neutral very happy

Publisher: Miles Hall		Publisher: Sterns Inc.	
Dictionary	Comprehensiveness (entries)	Dictionary	Comprehensiveness (entries)
1	86,000	1	98,000
2	62,000	2	68,000
3	38,000	3	44,000
4	10,000	4	20,000
Publisher: Elton Hills		Publisher: Wilson & Son	
1	92,000	1	55,000
2	80,000	2	32,000
3	74,000	3	26,000
4	56,000	4	14,000

Dictionary	Comprehensiveness (entries)	Publisher
1	98,000	Sterns Inc.
2	92,000	Elton Hills
3	86,000	Miles Hall
4	80,000	Elton Hills
5	74,000	Elton Hills
6	68,000	Sterns Inc.
7	62,000	Miles Hall
8	56,000	Elton Hills
9	55,000	Wilson & Sons
10	44,000	Sterns Inc.
11	38,000	Miles Hall
12	32,000	Wilson & Sons
13	26,000	Wilson & Sons
14	20,000	Sterns Inc.
15	14,000	Wilson & Sons
16	10,000	Miles Hall

Suppose that you are shopping for a PDA (e.g., Palm Pilot). A key factor that distinguishes a good one from a mediocre one is memory size. The more memory, the better. Suppose that there are only 25 models on the market, and that they are identical in all aspects except for memory size. These models are manufactured by five different companies and each company has its own brand, see tables below for details.

You are now in a store that sells electronics. It carries only one model of PDA—the Tversky Brand Tom Model. Compared with the other models within Tversky (i.e., Tracy, Trish, Ted, and Terry), how good is the Tom Model? Please examine the above information carefully and circle an answer below.

Suppose that its price is close to what you originally planned to pay for a PDA. Will you buy this PDA?

Aaker, D. A. 1996. *Building Strong Brands*. Free Press, New York.

Aaker, D. A., E. Joachimsthaler. 2000. *Brand Leadership*. Simon & Schuster Inc., New York, 98.

Aaker, D. A., K. L. Keller. 1990. Consumer evaluations of brand extensions. *J. Marketing* **54**(1) 27–42.

Abele, A. E., P. Petzold. 1998. Pragmatic use of categorical information in impression formation. *J. Personality Soc. Psych.* **75**(2) 347–358.

Bell, D. R., J. Chiang, V. Padmanabhan. 1999. The decomposition of promotional response: An empirical generalization. *Marketing Sci.* **18**(4) 504–526.

Biernat, M., M. Manis, T. E. Nelson. 1991. Stereotypes and standards of judgment. *J. Personality Soc. Psych.* **60** 485–499.

Boush, D. M., B. Loken. 1991. A process-tracing study of brand extension evaluation. *J. Marketing Res.* **28**(1) 16–22.

Cacioppo, J. T., R. E. Petty. 1982. The need for cognition. *J. Personality Soc. Psych.* **42** 116–131.

Cacioppo, J. T., R. E. Petty, J. A. Feinstein, W. B. G. Jarvis. 1996. Dispositional differences in cognitive motivation: The life and times of individuals varying in need for cognition. *Psych. Bull.* **119**(2) 197–253.

Cremer, H., J.-F. Thisse. 1991. Location models of horizontal differentiation: A special case of vertical differentiation models. *J. Indust. Econom.* **39**(4) 383–390.

Frank, M. 2002. VW buffs up. *Forbes* online ed., Vehicle Feature, <http://www.forbes.com/2002/05/13/0513feat.html>.

Helson, H. 1964. *Adaptation-Level Theory*. Harper and Row, New York.

Hsee, C. K. 1996. The evaluability hypothesis: An explanation for preference reversals between joint and separate evaluations of alternatives. *Organ. Behavior Human Decision Processes* **67** 247–257.

Hsee, C. K. 1998. Less is better: When low-value options are valued more highly than high-value options. *J. Behavioral Decision-Making* **11** 107–121.

Hsee, C. K., F. Leclerc. 1998. Will products look more attractive when presented separately or together? *J. Consumer Res.* **25** 175–186.

Hsee, C. K., G. F. Loewenstein, S. Blount, M. H. Bazerman. 1999. Preference reversals between joint and separate evaluations of options: A review and theoretical analysis. *Psych. Bull.* **125**(5) 576–590.

Hsee, C. K., F. You, J. Zhang, Y. Zhang. 2003. Medium maximization. *J. Consumer Res.* **30**(1) 1–15.

Kahneman, D., D. T. Miller. 1986. Norm theory: Comparing reality to its alternative. *Psych. Rev.* **93**(2) 136–153.

- Kahneman, D., I. Ritov. 1994. Determinants of willingness to pay for public goods: A study in the headline method. *J. Risk Uncertainty* 9 5–38.
- Kahneman, D., I. Ritov, D. Schkade. 2000. Economists have preferences; psychologists have attitudes: An analysis of dollar responses to public issues. D. Kahneman, A. Tversky, eds. *Choices, Values and Frames*. Cambridge University Press, New York.
- Kahneman, Daniel, Amos Tversky. 1979. Prospect theory: An analysis of decision under risk. *Econometrica* 47 263–292.
- Keller, K. L. 2003. *Strategic Brand Management: Building, Measuring and Managing Brand Equity*, ed. 2. Prentice Hall, Upper Saddle River, NJ.
- Kirmani, A., S. Sood, S. Bridges. 1999. The ownership effect in consumer responses to brand line stretches. *J. Marketing* 63(1) 88–100.
- Loken, B., D. Roedder John. 1993. Diluting brand beliefs: When do brand extensions have a negative impact? *J. Marketing* 57(July) 71–84.
- Randall, T., K. Ulrich, D. Reibstein. 1998. Brand equity and vertical product line extent. *Marketing Sci.* 17(4) 356–379.
- Read, D., G. Loewenstein. 1995. Diversification bias: Explaining the discrepancy in variety between combined and separate choices. *J. Experiment. Psych.: Appl.* 1(March) 34–49.
- Read, D., G. Loewenstein, M. Rabin. 1999. Choice bracketing. *J. Risk Uncertainty* 19 171–197.
- Shafir, E., P. Diamond, A. Tversky. 1997. Money illusions. *Quart. J. Econom.* 112(May) 341–373.
- Simonson, I., A. Tversky. 1992. Choice in context: Tradeoff contrast and extremeness aversion. *J. Marketing Res.* 29(August) 281–295.
- Simonson, I., S. Nowlis, K. Lemon. 1993. The effect of local consideration sets on global choice between lower price and higher quality. *Marketing Sci.* 12(4) 357–378.
- Tversky, A., D. Kahneman. 1981. The framing of decisions and the psychology of choice. *Science* 211 453–458.
- van Osselaer, S. M. J., J. W. Alba, P. Manchanda. 2004. Irrelevant information and mediated intertemporal choice. *J. Consumer Psych.* 14(3) 257–270.