

Hedonic Nondurability Revisited: A Case for Two Types

Raegan J. Tennant and Christopher K. Hsee
University of Chicago

Hedonic durability refers to the extent to which the hedonic impact of a change lasts, that is, how long the unhappiness from a loss (or happiness from a gain) will endure over time. The lesson from previous research on this topic has been that the long-term effect of most changes (e.g., larger incomes, bigger houses, shorter commutes) is negligible. The present research shows something different. Consistent with previous research, we observed a pattern of hedonic nondurability in which the impact of a change did not endure over time. However, we also observed a pattern of hedonic durability in which the impact of a change does endure over time. We demonstrate differential rates of hedonic durability for losses, both across variables (Experiment 1) and within different ranges of the same variable (Experiment 2). We also extend our research to show differential rates for gains (Experiment 3). To explain our results, we propose a distinction between preference types, arguing that comparison-independent (i.e., absolute) preference types are hedonically more durable than comparison-dependent (i.e., relative) preference types. This research offers a method for validating preference-type categorization as well as a novel paradigm for testing hedonic durability in the laboratory. Moreover, it yields theoretical insights for affective forecasting and adaptation as well as practical implications for the hedonic treadmill and the joyless economy.

Keywords: adaptation, cognition, evaluability, happiness, preferences

Derived from the Latin verb *durare* (to last), durable means “capable of withstanding wear and tear or decay” (Pickett, 2015, p. 556). For several decades, psychologists have been interested in the durability of hedonic reactions and have broadened the definition of the term as a result of their research. Studies encompass not only the decay of positive reactions, but also the rebound of negative ones (Bonanno, 2004; Brickman, Coates, & Janoff-Bulman, 1978; Gilbert, Pinel, Wilson, Blumberg, & Wheatley, 1998; Kahneman & Snell, 1992; Menzel, Dolan, Richardson, & Olsen, 2002; Schkade & Kahneman, 1998; Zamble, 1992; see Frederick & Loewenstein, 1999, for a review on hedonic adaptation). The most striking example from this body of research compared paraplegics and lottery winners with controls, showing that all reported being similarly happy in the long run (Brickman et al., 1978). Counter to expectations, hedonic reac-

tions in response to events—even extraordinary events, such as losing a limb or winning the lottery—appear to dip or soar for a much shorter period of time (Gilbert, 2006; Gilbert et al., 1998; Gilbert & Wilson, 2000; Riis et al., 2005; Taylor, 1983; Wilson & Gilbert, 2005). Subsequently, hedonic *non*-durability has seemed almost universal.

Yet counterexamples illuminate the boundaries of hedonic non-durability. Common sense suggests not all experiences regress at the same rate. For example, one would expect the psychological adjustment to the loss of one’s vision to take much longer than adjustment to the loss of one’s limb, even though both would be extraordinary losses. Moreover, research has shown that some experiences are hedonically durable. Studies on sensitization have shown people do not adapt to certain physical sensations, such as cold temperatures (e.g., Cabanac, 1981) and noise levels (e.g., Weinstein, 1982). Rather than becoming neutral over time, people find these sensations more intense and intolerable. These counterexamples reveal hedonic nondurability is not ubiquitous (cf. Mancini, Bonanno, & Clark, 2011).

Given the boundaries to hedonic nondurability, the present research investigates the relative rates of durability, and proposes a novel moderator: preference types. In the next section, we provide a context for the idea of preference types and specify our hypothesis. We also identify questions about preference types that are beyond the scope of this article. Finally, before reporting our experiments, we explain our novel experimental paradigm and briefly preview our three experiments.

Preference Types: Absolute Versus Relative

Different areas in the social sciences have put forth taxonomies classifying preferences into two general categories. The simplest classification in economics divides goods into necessities and luxuries, positing that the first type of goods serve fundamental

Raegan J. Tennant and Christopher K. Hsee, Department of Behavioral Science, Booth School of Business, University of Chicago.

Earlier versions of this paper appeared in the Raegan J. Tennant’s University of Chicago doctoral dissertation and on her Academia profile. She also presented this work at several conferences and workshops, including Society for Judgment and Decision Making, Society for Personality and Social Psychology, and Chicago Center for Decision Research.

The authors thank the Templeton Foundation for research support. This research is a part of Raegan J. Tennant’s doctoral dissertation, and she graciously thanks the dissertation committee members for their thoughtful and critical comments at both the proposal and defense. The authors thank the Chicago Decision Research Lab for its assistance with data collection, and many other colleagues at Chicago and beyond who have impacted the project in various ways.

Correspondence concerning this article should be addressed to Raegan J. Tennant, Workbar 45 Prospect St., Cambridge, MA 02139. E-mail: rtennant@uchicago.edu

needs, whereas the second type serve inessential but coveted wants (Scitovsky, 1976). In a similar view, sociologist Ruut Veenhoven distinguishes innate needs from acquired standards (Veenhoven, 1991). More recently, psychologists Christopher Hsee and colleagues have proposed a related categorization: Type A (inherent) versus Type B (learned). This categorization conceptualizes Type A as related to basic psychobiological needs and Type B as standards created by culture and society (Hsee, Xi, & Tang, 2008; Hsee, Yang, Li, & Shen, 2009; Hsee & Zhang, 2010; Tu & Hsee, 2016; cf. Maslow, 1943). One further assumption these taxonomies share is that the first category of preferences (i.e., necessities, innate needs, Type A) is more fundamental to happiness because it is relevant to a person's basic well-being, whereas the second category (i.e., luxuries, acquired standards, Type B) is less fundamental because it is context dependent, learned through trends, customs, and culture.

Drawing on these taxonomies, we propose two types of preference exist: absolute versus relative. A key difference between absolute and relative preference types is the role of standards of comparison on value. The value of absolute preferences is comparison independent, whereas the value of relative preferences is comparison dependent. For example, suppose people prefer one state (call it X) to another state (call it Y). We propose that a preference is comparison independent (i.e., an absolute type) if, even without comparison, people in state X are happier than those in Y. That is, state X is better than state Y independent of comparison. By contrast, a preference is comparison dependent (i.e., a relative type) if, without comparison, people in state X are not happier than those in Y, and only with comparison will people in state X be happier than those in Y. That is, state X is better than state Y dependent on comparison.

Hsee et al. (2009) collected field and laboratory data that provide examples for both preference types. Preference for house temperatures in winter showed an absolute pattern. In both within-city and across-city comparisons, people with warmer house temperatures in the winter were happier than people with colder house temperatures. Similar findings were replicated in the laboratory for water temperature. By contrast, preference for jewelry value showed a relative pattern. Within a city, people with more expensive jewelry were happier than people with less expensive jewelry, but across cities, this effect disappeared. People in cities with more expensive jewelry were no happier than people in cities with less expensive jewelry. Presumably, people within cities can engage in social comparison of jewelry value, but they cannot do this comparison across cities. These findings were replicated in the laboratory using diamond size.

As illustrated by these examples, standards of comparisons have a significant influence on assessments of value for relative preference types (e.g., jewelry value). Without comparison, people become indifferent between state X or Y. Only when they compare X with Y do they prefer X to Y. By contrast, for absolute preferences (e.g., house temperature), value is more inherent to the stimuli, and people react more positively to X than Y regardless of whether direct comparison can be made.

The authors have debated how susceptible absolute preferences are to shifts in value based on the relevant contrast. What is agreed upon is that for absolute preferences, the comparison is not necessary for people to know which level is good and which is not. However, we are not in agreement about the susceptibility of this

preference type to contrast effects. One view is that comparisons can still influence absolute preferences. For example, most people find a 70°F room temperature comfortable, but may initially find that room to be too cool if they came indoors from 100°F heat or too hot if they came indoors from 40°F. The other view is that absolute preferences are less susceptible to shifts in value caused by exogenous factors, such as comparisons, than are relative preference types. The argument is that people register the value of absolute preference types fairly accurately and report them reliably, regardless of the context or comparison. Our experiments provide fodder for this debate but do not completely resolve it.

Also note that certain questions about preference types are beyond the scope of this article. For example, do basic needs determine absolute value? If so, what role does biology play? More broadly, do preference types pivot on a psychobiological versus sociocultural axis? And are preference types perhaps evaluated by different systems of thought? Other scholars (e.g., Bettman, Luce, & Payne, 2008; Dhar & Novemsky, 2008; Hsee et al., 2008, 2009; Hsee & Zhang, 2010; Kivetz, Netzer, & Schrieff, 2008; Simonson, 2008; Veenhoven, 1991) have engaged with these broader questions, and we outline a few perspectives in this article's general discussion.

Preference Types and Hedonic Durability

We hypothesize that preference types moderate hedonic durability. That is, holding constant the initial hedonic impact of a change, the change will have a longer-lasting effect if it affects an absolute-preference type than if it affects a relative-preference type. In our experiments, we use a paradigm designed specifically for the goal of investigating relative rates of durability in the laboratory. In our paradigm, we first collect data for a baseline measure (Time 0) and then introduce an exogenous change (a downgrade or an upgrade) and then collect data for the immediate impact of the change (Time 1) and the lasting impact of the change (Time 2). See Figure 1 for an illustration. As the figure shows, the solid line depicts responses to the absolute (i.e., comparison-independent) preference variable, and the dotted line depicts responses to a relative (i.e., comparison-dependent) preference variable.

To meaningfully compare the two types of responses, we design the stimuli so that the initial impact of the change is roughly the same; in Figure 1 that means the initial drop from Time 0 to Time 1 is roughly the same for the solid line and the dotted line. (Note that, in our experiments, the y axis is a rating scale; we do not know whether responses on the scale linearly represent how people feel. The best we could do is keep the numerical size of the two drops roughly the same.) We predict that after a similar initial drop, the solid line, representing an absolute (i.e., comparison-independent) preference, will last; namely, it will not return to the baseline at Time 2. By contrast, the dotted line, representing a relative (i.e., comparison-dependent) preference, will not last; namely, it will return to the baseline at Time 2.

Our theory for these predictions is the following: when people are evaluating the baseline (Time 0), they are mostly in single evaluation; they may rely on a fuzzy exemplar experience from memory, but mostly they lack a salient comparison and are singularly evaluating the baseline experience. However, immediately after a downgrade occurs (Time 1), people are in joint evaluation,

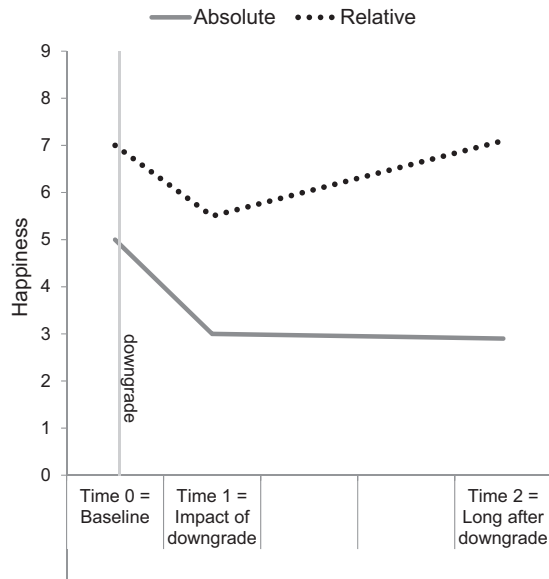


Figure 1. Paradigm and predictions for hedonic durability of preference types. Specifically, prediction of different rates of durability as a function of preference types for a downgrade. Two patterns of durability are predicted: hedonic durability for absolute-preference types (solid line) and hedonic nondurability for relative-preference types (dotted line). The y axis measures the happiness, or the subjective value, of the experience. The x axis measures time at three time points: a baseline measure (Time 0), the immediate impact of the change (Time 1), and the impact of the change long after it has occurred (Time 2). In our experiments, the hedonic trajectory is a between-subjects measure; it is mapped across individuals' assessments of time in order to avoid within-subject consistency effects (see Frederick & Loewenstein, 1999, for a review). This figure depicts predictions for a downgrade; the inverse pattern is predicted for an upgrade.

in which they evaluate the downgraded experience in comparison to how they felt about the baseline experience. Comparisons here are not fuzzy, but quite salient, though not physically direct. As time passes (between Time 1 and Time 2), people gradually forget the baseline experience and attend mostly, or only, to the current state. Thus, when people report their evaluation at Time 2, they are singularly evaluating the present experience.

As explained earlier, for absolute (i.e., comparison-independent) preferences, people will feel worse in a worse state than in a better state regardless of whether they are in joint evaluation or in single evaluation of the two states. By contrast, for relative (i.e., comparison-dependent) preferences, people will feel worse in a worse state than in a better state only in joint evaluation of the two states, and not in single evaluation. Therefore, we predict that for comparison-independent preferences, people will feel worse at both Time 1 and Time 2 than at Time 0; that is, the change (the downgrade) is hedonically durable. By contrast, for comparison-dependent preferences, people will feel worse at Time 1 than at Time 0, but will not feel worse at Time 2 than at Time 0; that is, the change (the downgrade) is not hedonically durable (see Figure 1). Note the preceding analysis concerns a downgrade; the same applies for an upgrade, except the patterns in Figure 1 would be the inverse.

Overview of the Three Experiments

In the upcoming experiments, we used the aforementioned paradigm to investigate relative rates of durability and the preference-types-as-moderator hypothesis. Specifically, we report three experiments. In Experiment 1, we conducted a preliminary test of our hypothesis using different stimuli for the preference types. In Experiments 2 and 3, we employed a cleaner operationalization of preference types by using the same stimuli but manipulating only one of its attributes.

Hsee et al. (2009) used qualitatively different stimuli to represent the two types of preferences; in the laboratory, shower water temperature represented a comparison-independent preference whereas diamond size represented a comparison-dependent preference. We extended this thinking to the selection of stimuli for our Experiment 1 and manipulate different aspects of a female model's appearance, namely, the darkness of her facial hair and/or the stylishness of her sunglasses. We verified in a pretest that the preference for the female model's face with light hair versus dark hair is comparison independent (absolute), and the preference for that female model's face with a pair of stylish sunglasses versus nonstylish ones is comparison dependent (relative). Given these results, we predicted that a downgrade to the female model's facial hair (from light to dark) is hedonically more durable than a downgrade in sunglass style (from fashionable to unfashionable).

Although previous research and our first experiment used different types of stimuli to represent different types of preferences, our other experiments used the same stimulus—but different ranges—to represent the different preference types. In Experiments 2 and 3, we used the screen size of prototypical e-readers as our stimulus. We manipulated absolute versus relative preference types by focusing on different ranges of the stimulus. In one condition, we used a small-size screen range (e.g., comparing a 4×7 with a 3×6 screen). In the other, we used a large-size screen range (e.g., comparing a 7×11 with a 5×9 screen). We verified in a pretest that the preference for a relatively larger screen among small screens (e.g., the preference for 4×7 over 3×6) was comparison independent (i.e., absolute). Whereas, the preference for a relatively larger screen among large screens (e.g., the preference for 7×11 over 5×9) was comparison dependent (i.e., relative). Given these results, we predicted that holding the initial hedonic impact constant, a change in screen size among small screens is hedonically more durable than a change in screen size among large screens. Experiment 2 examined the effects of a downward change in screen size (i.e., a downgrade), and Experiment 3 examined the effects of an upward change in screen size (i.e., an upgrade).

Experiment 1

Experiment 1 aimed to show hedonic durability could vary between different types of changes. In this experiment, we categorized a person's facial-feature appearance, in particular, a female's facial hair, as an absolute-preference type, and categorized the stylishness of a person's fashion accessory, in particular, sunglasses, as a relative-preference type. We found support for our categorizations in a pretest, which is presented next. We predicted that, holding constant the initial hedonic impact of a change, the change will be more hedonically durable if it affects facial-feature appearance than if it affects fashion-accessory style.

Pretest

We assumed preferences for fashion accessories were relative whereas preferences for facial features were absolute. To verify these assumptions, we conducted a pretest. One hundred one Mechanical Turk (MTurk) participants (43 females, $M_{\text{age}} = 37.35$, $SD = 10.03$) were assigned to one of three conditions: baseline, facial-feature downgrade, and fashion-accessory downgrade. Participants looked at a photograph of a female model holding a sign in the corresponding condition and then assessed her attractiveness on a 100-point sliding scale in which *ugly* and *good-looking* were the left and right endpoints, respectively. We expected the female model to be rated as less attractive in the facial-feature-downgrade condition than in both the baseline and fashion-accessory-downgrade conditions. The results supported our prediction: attractiveness ratings were significantly lower in the facial-feature-downgrade condition ($M = 37.68$, $SD = 23.60$) than in both the baseline condition ($M = 53.03$, $SD = 18.84$), $t(65) = 2.94$, $p = .005$, and the fashion-accessory downgrade condition ($M = 51.59$, $SD = 21.04$), $t(66) = 2.57$, $p = .013$. Results revealed no significant difference between the latter two conditions, $t(65) = .295$, ns .¹ Supporting our assumption, these results suggested that without direct comparison, the presence or absence of facial hair made a difference to perceived attractiveness, but sunglass style did not.

Method

We assigned 150 participants to a five-level between-participants design. One condition was the baseline (Time 0), and the other four conditions were constructed by a 2 (type: facial-feature appearance vs. fashion-accessory style) $\times$ 2 (rating time: t1 vs. t2). Given statistical guidelines (Agresti & Finlay, 2009), our overall sample size was composed of 30 participants in each of the five conditions.

Participants were recruited from the Chicago Decision Research Lab (71 females, $M_{\text{age}} = 19.83$, $SD = 2.93$) for an experiment on "Memory & Perception." They read the following instructions before beginning the experiment:

This study is about perceptual attention and memory. You will be presented with a series of photographs of a young woman holding a sign. Different letters will appear on the sign every few seconds and an aspect of the woman's appearance may change at some point. Please pay attention, because periodically you will be prompted to recall and/or evaluate what you have observed.

On individual computers, participants watched this long series of photographs. In each photograph, the female model was wearing a pair of sunglasses and holding a sign with a different letter on it. In the baseline condition, the female model's physical appearance remained unchanged throughout the series, whereas in the two type conditions (facial-feature appearance vs. fashion-accessory style), the model's appearance changed a half-minute into the series. The change was negative, but the type of change varied by condition. In the facial-feature-appearance condition, the female model's upper lip hair changed from her normal light hair to slightly darker hair. By contrast, in the fashion-accessory-style condition, the model's sunglasses changed to unstylish sunglasses.

All participants were prompted to answer five questions asking, "What was the last letter you saw?" These five questions occurred at different points throughout the series. The purpose of these questions was to maintain the theme of the experiment on memory and perception and to exclude noncompliant (i.e., inattentive) participants, that is, those who answered all of these questions incorrectly.² The key dependent variable question asked participants to rate the attractiveness of the female model on a 22.5 cm horizontal-line scale from *very unattractive* (on the left) to *very attractive* (on the right). Note this rating scale had two endpoint labels yet no tick marks or numbers on the horizontal line. This attractiveness rating was elicited at one of three time points: Time 0, Time 1, or Time 2. Time 0 was 30 s into viewing photographs of the model holding a sign; Time 1 was 12 s after the negative change to the model's appearance occurred; and Time 2 was 156 s after the negative change to the model's appearance occurred.

To summarize, participants in the baseline condition did not see a change to the female model's appearance, whereas, participants in the type conditions viewed either a change to the stylishness of the female model's sunglasses (fashion-accessory-style condition) or a change to the darkness of the female model's upper lip hair (facial-feature-appearance condition). All participants made one critical attractiveness rating when prompted. Ratings at Time 0 assessed a baseline measure of the female model's appearance; at Time 1, they assessed the experience with the downgraded appearance of the model shortly after the change; and at Time 2, they assessed the experience with the downgraded appearance of the model long after the change.

Results and Discussion

The results, summarized in Figure 2, supported our prediction that holding constant the initial hedonic impact of a negative change, the change in the facial-feature-appearance condition was more durable than the downgrade in the fashion-accessory-style condition.

Specifically, we first ran a 2 (type: facial-feature appearance vs. fashion-accessory style) $\times$ (time: t0 vs. t1) ANOVA to show that the initial hedonic impact of the downgrade (from t0 to t1) was similar between the conditions.³ The analysis yielded no significant main effect of type, $F(1, 119) = 0.00$, ns , $\eta_p^2 = .000$, a significant main effect of time, $F(1, 119) = 4.61$, $p = .034$, $\eta_p^2 = .039$, but most importantly, no significant interaction between type and time, $F(1, 119) = 0.00$, ns , $\eta_p^2 = .000$. The lack of a significant interaction suggests the negative change in the two conditions initially created similar impacts on attractiveness ratings.

We next ran a 2 (type: facial-feature appearance vs. fashion-accessory style) $\times$ 2 (time: t1 vs. t2) ANOVA to test our critical prediction concerning differential rates of durability between t1 and t2. The analysis yielded a significant main effect of type, $F(1,$

¹ Prior to analyses, we dropped two participants from the analyses because they were considered noncompliant; that is, they answered each of the five letter-memory questions incorrectly.

² The statistical abbreviation *ns* stands for not statistically significant.

³ Prior to analyses, the baseline condition's ratings were duplicated so that each of the type-condition's Time 1 ratings could be compared with the baseline condition's Time 0 ratings.

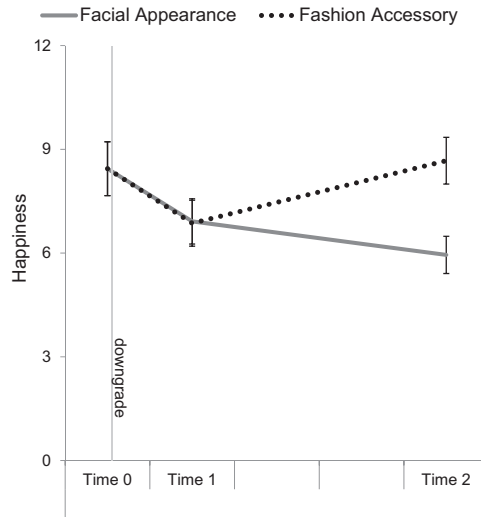


Figure 2. Results of Experiment 1. Hedonic durability of a downgrade as a function of the type of change to a female model's appearance. In the facial-appearance condition, a feature of the female model's facial appearance changed, and in the fashion-accessory condition, the style of the female's fashion accessory changed. Ratings were elicited on a plain horizontal line with two endpoint labels; the line was 22.5 cm long. Coders measured the left endpoint, which had a negative label, as zero, whereas they measured the right endpoint, which had a positive label, as 22.5. Therefore, larger numbers on the graph's y axis indicate higher levels of happiness. The error bars for the point estimates are standard errors.

117) = 4.44, $p = .037$, $\eta_p^2 = .038$, no significant main effect of time, $F(1, 117) = 0.44$, ns , $\eta_p^2 = .004$, and, most importantly, a significant interaction between type and time, $F(1, 117) = 4.79$, $p = .031$, $\eta_p^2 = .041$. This interaction supported our prediction, suggesting that after the initial drop, the attractiveness ratings regressed faster in the fashion-accessory-style condition than in the facial-feature-appearance condition. Planned contrasts found that in the fashion condition, attractiveness ratings increased (i.e., regressed) significantly at the .1 level from t1 to t2, $F(1, 58) = 3.63$, $p = .062$, $\eta_p^2 = .061$, whereas in the appearance condition, attractiveness ratings remained largely unchanged during the same period, $F(1, 59) = 1.32$, ns , $\eta_p^2 = .023$. Moreover, additional planned contrasts found that for the fashion condition, attractiveness ratings at t0 and t2 were not significantly different, $F(1, 58) = 0.05$, ns , $\eta_p^2 = .001$, whereas for the appearance condition, the difference between t0 and t2 was highly significant, $F(1, 60) = 6.95$, $p = .011$, $\eta_p^2 = .107$.

Experiment 1 lent support to our prediction, suggesting that even though the two negative changes created similar initial effects, a change in the absolute-preference-type condition created a longer-lasting effect on attractiveness ratings than did an initially comparable change in the relative-preference-type condition. A limitation of this experiment, however, is that it involved two attributes of physical attraction (i.e., facial appearance and fashion accessories), making de-confounding all the possible factors in these attributes difficult. Next we therefore focus on a cleaner manipulation for our independent variable by studying the different rates of durability within different ranges of a single attribute, namely, the screen size of e-readers.

Experiment 2

In both this and the next experiment, we manipulated the size of prototypical e-reader display screens. Holding everything else constant, larger e-reader screens are typically considered more desirable than smaller sizes. However, we argue the preferences for a relatively large screen in a small-size-screen range (e.g., between a small and medium screen) is comparison independent (i.e., absolute) and therefore hedonically durable. By contrast, the preference for a relatively larger screen in a large-size-screen range (e.g., between a medium and large screen) is comparison dependent (i.e., relative) and therefore hedonically nondurable. We predicted that holding constant the initial impact of a change, the change will have a longer-lasting hedonic effect if it occurs in the small-size-screen range than if it occurs in the large-size-screen range. Both Experiments 2 and 3 tested this prediction. Similar to Experiment 1, Experiment 2 investigated a negative change (a downgrade), whereas Experiment 3 investigated a positive change (an upgrade).

Pretest

We assumed preferences for larger screens in the small-size range were absolute and preferences for larger screens in the large-size range were relative. To verify these assumptions, we conducted a pretest. One hundred seventeen MTurk participants (53 females, $M_{\text{age}} = 37.63$, $SD = 10.04$) were assigned to one of four e-reader screen-size conditions: 3×6 , 4×7 , 5×9 , and 7×11 . Participants read a neutral text on a prototypical e-reader in the corresponding condition for 30 s and then assessed the comfort of the reading experience on a 100-point sliding scale in which *very uncomfortable* and *very comfortable* were the left and right endpoints, respectively. The pretest results showed that, in single evaluation, participants viewing the 3×6 screen felt worse ($M = 53.38$, $SD = 27.16$) than those viewing the 4×7 screen ($M = 71.14$, $SD = 24.89$), $t(56) = 2.60$, $p = .012$. By contrast, in single evaluation, participants viewing the 5×9 screen did not feel worse ($M = 77.36$, $SD = 21.84$) than those viewing the 7×11 screen ($M = 79.10$, $SD = 19.16$), $t(55) = 0.32$, ns . Supporting our assumption, the results indicated that without direct comparison, the smaller size (3×6) and the larger size (4×7) in the small-size range made a difference to hedonic experience, but the smaller size (5×9) and the larger size (7×11) in the large-size range made no difference.

Method

We assigned 179 participants from the Chicago Decision Research Lab (129 females, $M_{\text{age}} = 19.47$, $SD = 1.48$) to a 2 (range: small size vs. large size) $\times$ 3 (rating time: t0, t1, t2) between-participants design.⁴ Note that unlike Experiment 1, Experiment 2 was a fully crossed design. Given statistical guidelines (Agresti & Finlay, 2009), our overall sample size was composed of 30 participants in each of the six conditions. However, in this experiment, one data point was not recorded due to a clerical error.

⁴ Prior to analyses, we dropped four participants from the analyses because they were outliers, defined here as extreme values between 1.5 and 3 times the interquartile range of the boxplot.

Each participant read from two prototypical e-readers—the first e-reader for a brief period of time (2 min), followed by a second, downgraded one for a much longer period (10 min and 15 s). In the large-size-range condition, the sizes of both e-readers were relatively large: the first was 7×11 in. and the second was 5×9 in. By contrast, in the small-size-range condition, the sizes of both e-readers were relatively small: the first was 4×7 in. and the second was 3×5 in. Participants were prompted to answer two questions at one of three time points: Time 0, Time 1, or Time 2. Time 0 was 90 s into reading on the first e-reader; Time 1 was 30 s into reading on the second e-reader; and Time 2 was 615 s into reading on the second e-reader. The first question participants answered was nonessential and was about their attitude toward the content of the reading (excerpts from the student handbook of their university). The second question participants answered was our key dependent variable, which was about their happiness with their e-reader reading experience. Participants answered the question by marking a 22.5 cm horizontal-line scale from *very uncomfortable* (on the left) to *very comfortable* (on the right). Note that this rating scale had two endpoint labels yet no tick marks or numbers on the horizontal line.

In sum, participants read either on a large display and then a medium display (the large-size range) or a medium display and then a small display (the small-size range) and made one happiness rating when prompted. Happiness ratings at Time 0 assessed experience with the first display; at Time 1, they assessed the experience with the second (downgraded) display shortly after the change; and at Time 2, they assessed the experience with the second (downgraded) display long after the change.

Results and Discussion

The results, summarized in Figure 3, supported our prediction that holding constant the initial hedonic impact of the downgrades, the downgrade in the small-size range was more durable than the downgrade in the large-size range. Specifically, we first ran a 2 (range: small size vs. large size) $\times$ (time: t0 vs. t1) ANOVA to show the initial hedonic impact of the downgrade (from t0 to t1) was similar between the small and large conditions. The analysis yielded a significant main effect of range, $F(1, 117) = 12.72, p = .001, \eta_p^2 = .101$, a significant main effect of time, $F(1, 117) = 36.34, p < .001, \eta_p^2 = .243$, but most importantly, no significant interaction between range and time, $F(1, 117) = 2.37, ns, \eta_p^2 = .021$. The lack of a significant interaction suggests the downgrades in screen size in the two ranges initially created similar impacts on happiness ratings.

We next ran a 2 (range: small size vs. large size) $\times$ 2 (time: t1 vs. t2) ANOVA to test our critical prediction concerning differential rates of durability between t1 and t2. The analysis yielded a significant main effect of range, $F(1, 115) = 57.14, p < .001, \eta_p^2 = .340$, a marginally significant main effect of time, $F(1, 115) = 3.13, p = .080, \eta_p^2 = .027$, and, most importantly, a significant interaction between range and time, $F(1, 115) = 4.21, p = .043, \eta_p^2 = .036$. This interaction supported our prediction, suggesting that after the initial drop, the happiness ratings regressed faster in the large-size range than in the small-size range. Planned contrasts found that in the large-size range, happiness ratings increased (i.e., regressed) significantly from t1 to t2, $F(1, 59) = 4.63, p = .036, \eta_p^2 = .075$, whereas in the small-size range, happiness ratings remained largely unchanged dur-

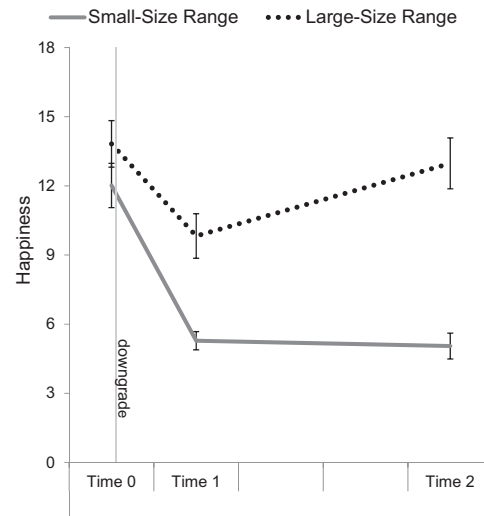


Figure 3. Results of Experiment 2. Hedonic durability of a downgrade as a function of the range of the size change. The small-size range represents relatively small-size e-reader screens. The large-size range represents relatively large-size e-reader screens. Ratings were elicited on a plain horizontal line with two endpoint labels; the line was 22.5 cm long. Coders measured the left endpoint, which had a negative label, as zero, whereas they measured the right endpoint, which had a positive label, as 22.5. Therefore, larger numbers on the graph's y axis indicate higher levels of happiness. The error bars for the point estimates are standard errors.

ing the same period, $F(1, 56) = 0.11, ns, \eta_p^2 = .002$. Moreover, additional planned contrasts found that for the large-size range, happiness ratings at t0 and t2 were not significantly different, $F(1, 59) = 0.34, ns, \eta_p^2 = .006$, whereas for the small-size range, the difference between t0 and t2 was highly significant, $F(1, 59) = 38.61, p < .001, \eta_p^2 = .404$.

In sum, Experiment 2 lent support to our prediction, suggesting that even though the two downgrades created similar initial effects, a downgrade in the small-size-screen range created a long-lasting effect on happiness ratings, whereas a downgrade in the large-size-screen range did not. The next experiment tested whether this prediction holds for a positive change, namely, an upgrade in e-reader screen sizes.

Experiment 3

Pretest

We assumed preferences for larger screens in the small-size range were absolute and preferences for larger screens in the large-size range were relative. Given that the e-reader screen sizes in Experiments 2 and 3 were similar, we used the pretest data from Experiment 2 to verify our assumptions for Experiment 3. The pretest results showed that, in single evaluation, participants viewing the 4×7 screen felt better ($M = 71.14, SD = 24.89$) than those viewing the 3×6 screen ($M = 53.38, SD = 27.16$), $t(56) = 2.60, p < .012$. By contrast, in single evaluation, participants viewing the 7×11 screen did not feel better ($M = 79.10, SD = 19.16$) than those viewing the 4×7 screen ($M = 71.14, SD =$

24.89), $t(56) = 1.37$, *ns*. Supporting our assumption, the results indicated that without direct comparison, the larger size (4×7) and the smaller size (3×6) in the small-size range made a difference in happiness ratings, but the larger size (7×11) and the smaller size (4×7) in the large-size range did not.

Method

Whereas Experiment 2 tested the hedonic durability of the impact of a downgrade in screen size, Experiment 3 tested that of an upgrade. We assigned 180 participants from the Chicago Decision Research Lab (74 females, $M_{\text{age}} = 20.41$, $SD = 5.11$) to a 2 (range: small size vs. large size) $\times$ 3 (rating time: t0, t1, t2) between-participants design.⁵ Given statistical guidelines (Agresti & Finlay, 2009), our overall sample size was composed of 30 participants in each of the six conditions. The procedure of Experiment 3 was similar to that of Experiment 2 except the size change was reversed and the time sequences between Time 0 and Time 1 and Time 1 and Time 2 were shorter.

Each participant read from two prototypical e-readers—the first e-reader for a brief period of time (2 min), followed by the second (upgraded) one for a longer period (5 min and 45 s). In the large-size-range condition, the sizes of both e-readers were relatively large: the first was 4×7 in. and the second was 7×11 in. By contrast, in the small-size-range condition, the sizes of both e-readers were relatively small: the first was 3×6 in. and the second was 4×7 in. Participants were prompted to answer two questions at one of three time points: Time 0, Time 1, or Time 2. Time 0 was 90 s into reading on the first e-reader; Time 1 was 15 s into reading on the second e-reader; and Time 2 was 345 s into reading on the second e-reader. The first question participants answered was nonessential and was about their attitude toward the content of the reading (excerpts from an encyclopedia entry on ancient civilization). The second question participants answered was our key dependent variable, which was about their happiness with their e-reader reading experience. Participants answered the question by marking a 22.5 cm horizontal-line scale from *very uncomfortable* (on the left) to *very comfortable* (on the right). Note this rating scale had two endpoint labels yet no tick marks or numbers on the horizontal line.

In sum, participants read either on a medium display and then a large display (the large-size range) or a small display and then a medium display (the small-size range) and made one happiness rating when prompted. Happiness ratings at Time 0 assessed experience with the first display; at Time 1, they assessed the experience with the second (upgraded) display shortly after the change; and at Time 2, they assessed the experience with the second (upgraded) display long after the change.

Results and Discussion

The results, summarized in Figure 4, supported our prediction that holding constant the initial hedonic impact of the upgrades, the upgrade in the small-size range was more durable than the upgrade in the large-size range. Specifically, we first ran a 2 (range: small size vs. large size) $\times$ (time: t0 vs. t1) ANOVA to show the initial hedonic impact of the upgrade (from t0 to t1) was similar between the small and large conditions. The analysis yielded a significant

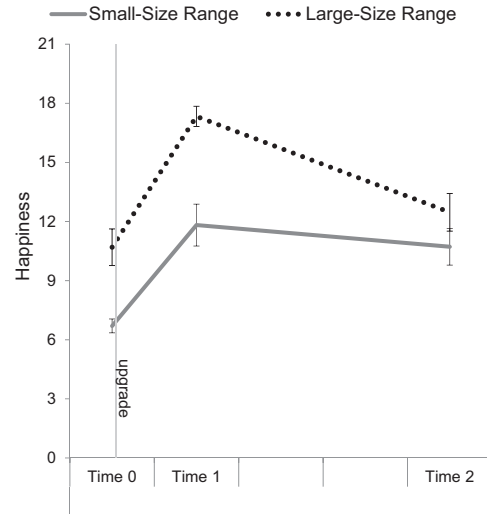


Figure 4. Results of Experiment 3. Hedonic durability of an upgrade as a function of the range of the size change. The small-size range represents relatively small-size e-reader screens. The large-size range represents relatively large-size e-reader screens. Ratings were elicited on a plain horizontal line with two endpoint labels; the line was 22.5 cm long. Coders measured the left endpoint, which had a negative label, as zero, whereas they measured the right endpoint, which had a positive label, as 22.5. Therefore, larger numbers on the graph's y axis indicate higher levels of happiness. The error bars for the point estimates are standard errors.

main effect of range, $F(1, 109) = 32.55$, $p < .001$, $\eta_p^2 = .237$, a significant main effect of time, $F(1, 109) = 49.75$, $p < .001$, $\eta_p^2 = .321$, but most importantly, no significant interaction between range and time, $F(1, 109) = 0.83$, *ns*, $\eta_p^2 = .008$. The lack of a significant interaction suggests the upgrades in the two ranges initially created similar impacts on the happiness ratings.

We next ran a 2 (range: small size vs. large size) $\times$ 2 (time: t1 vs. t2) ANOVA to test our critical prediction concerning differential rates of durability between t1 and t2. The analysis yielded a significant main effect of range, $F(1, 116) = 15.76$, $p < .001$, $\eta_p^2 = .123$, a significant main effect of time, $F(1, 116) = 10.63$, $p = .001$, $\eta_p^2 = .087$, and, most importantly, a significant interaction between range and time, $F(1, 116) = 4.24$, $p = .042$, $\eta_p^2 = .037$. This interaction supported our prediction, suggesting that after the initial increase, the happiness ratings regressed faster in the large-size range than in the small-size range. Planned contrasts found that in the large-size range, happiness ratings decreased (i.e., regressed) significantly between t1 and t2, $F(1, 56) = 18.44$, $p < .001$, $\eta_p^2 = .255$, whereas in the small-size range, happiness ratings remained largely unchanged during the same period, $F(1, 60) = 0.60$, *ns*, $\eta_p^2 = .010$. Moreover, additional planned contrasts found that for the large-size range, happiness ratings at t0 and t2 were not significantly different, $F(1, 60) = 1.77$, *ns*, $\eta_p^2 = .030$, whereas for the small-size range, the difference between t0 and t2 was highly significant, $F(1, 53) = 13.05$, $p = .001$, $\eta_p^2 = .204$.

⁵ Prior to analyses, we dropped 11 participants from the analyses because they were outliers, defined here as extreme values between 1.5 and 3 times the interquartile range of the boxplot.

Together, Experiment 2 and Experiment 3 supported our hypothesis. That is, a change (whether it is a downgrade or an upgrade) in the small-size range created a longer-lasting effect on happiness ratings than did a comparable hedonic change in the large-size range.

General Discussion

The present research investigated the hedonic durability of different changes. We posited that, holding constant the initial impact of a change, the change would have a longer-lasting hedonic effect if it concerned an absolute-preference type than if it concerned a relative-preference type. Operationally, we defined the construct in terms of an absolute (i.e., comparison-independent) preference for facial-feature appearance versus relative (i.e., comparison-dependent) preference for fashion-accessory style in Experiment 1, and defined the construct in terms of an absolute (i.e., comparison-independent) preference for values in the small-size e-reader screen range versus relative (i.e., comparison-dependent) preference for values in the large-size e-reader screen range in Experiments 2 and 3. The results of all three experiments supported our hypothesis: holding constant the initial impact of a change, the change, whether it is a downgrade or an upgrade, is hedonically less durable when it affects a relative-preference type than if it affects an absolute-preference type.

Hedonic Nondurability: Discrepancy Between Absolute and Relative

Early investigations on this topic (e.g., [Brickman et al., 1978](#)), which concluded the impact of changes (good or bad) on well-being and happiness in the long run are negligible, hinted that cognitive processes may underlie the apparent hedonic nondurability of most changes. Specifically, one theory implicates the influence of contrast on short-term but not long-term evaluations of the change ([Brickman et al., 1978](#)). This theory proposes that when a change occurs, a contrast between the initial state and the new state is available ([Gilbert et al., 1998](#); [Smith, Loewenstein, Jankovic, & Ubel, 2009](#)). When asked to judge the value of the new state, while the contrast is activated in the short run, people make a comparative assessment. But when asked to judge the value of the new state, while the contrast is no longer activated in the long run, people make an absolute judgment. This theory concludes hedonic nondurability arises because many changes are better or worse in relative but not absolute terms ([Brickman et al., 1978](#); [Smith et al., 2009](#); [Taylor, 1983](#); [Wilson & Gilbert, 2005](#)).

This proposal is a plausible account for a set of our findings. We observed a pattern of nondurability for the fashion-accessory-style downgrade in Experiment 1 and large-size-range downgrade and upgrade in Experiments 2 and 3. In these cases, consistent with previous research, the negative or positive change initially affected happiness (at Time 1), but eventually happiness returned to baseline levels (at Time 2). In our experimental paradigm, assessments at Time 1 were made seconds after the change, such that a contrast between the initial state and the new state was possible. This could have induced subjects to judge the value of the new state relative to the value of initial state. On the other hand, assessments at Time 2 were made many minutes after the change occurred, such that the

contrast may no longer have been salient. Without a contrast available, subjects would have defaulted to making an absolute assessment of the state.

The relative-absolute discrepancy theory has explanatory merit for nondurability but requires further thinking for it also to be an account of hedonic durability. We observed a pattern of hedonic durability for the facial-feature-appearance downgrade in Experiment 1 and small-size-range downgrade and upgrade in Experiments 2 and 3. Here we observed that the negative or positive change initially affected happiness (at Time 1) and continued to have a significant impact on happiness in the long run (at Time 2). Even though our experimental paradigm was the same, and therefore relative assessments at Time 1 and absolute assessments at Time 2 were as likely, we did not observe a discrepancy. Time 1 and Time 2 assessments are well calibrated, meaning judgments made from joint evaluation were also identical to those made in single evaluation.

One way to reframe the relative-absolute discrepancy theory is to incorporate preference types as a moderator of hedonic durability versus nondurability. Hedonic nondurability can arise because some preferences derive value from relative standing but lack a meaningful absolute value, and a discrepancy between relative and absolute assessments will emerge when comparing the short-run impact of the change to its long-run value. By contrast, hedonic durability can arise because some preferences have significant absolute value and are less affected by comparative shifts caused by contrasts, and this facilitates a nondiscrepancy, or a calibration, of the short-run impact of the change and its long-run value. Although this reframing may push the original theory beyond its limits, we believe the reframing is quite compelling. It does, however, trigger two important questions: What drives the relative assessments in the short run? Why does one set of experiences have significant absolute value whereas another set lacks it? We address these questions next.

Forms of Relative Assessments and Potential Source of Preference Types

Relative assessments arise from the availability of a contrast, a standard of comparison, and this standard of comparison can take different psychological forms. The form might be merely perceptual. For instance, research has shown pairing a neutral photograph with an attractive photograph can lead to a perceptual contrast that causes the neutral photograph to be judged as less attractive than it would be judged on its own ([Melamed & Moss, 1975](#)). Other research has shown similar effects when stimuli appear sequentially, as they did in our experiments ([Arkes, Hirshleifer, Jiang, & Lim, 2008](#); [Schwarz & Bless, 2007](#)). A perceptual contrast could have driven the relative assessment (at Time 1). Yet the perceptual contrast could also possibly activate certain cognitions that denote a rule about relative standing, and people might use this rule when making the relative assessment. For instance, in Experiment 3, when participants were upgraded to a larger version of the e-reader, they might have thought of the rule "bigger is better." This rule, not merely the perceptual contrast, could have driven the relative assessment.

The implication of these different forms of relative assessments is important for the interpretation of the effects. Up until this point, the more basic form, the perceptual contrast, has been assumed to

be the mechanism of the relative assessment. The second type of relative assessment changes the account of the effects. It implies that a rule about relative standing (e.g., bigger is better) is applied in the short run, yet in the long run, this rule is accurate in the absolute-preference-type case (e.g., small to medium) but not in the relative-preference-type case (e.g., medium to large). This interpretation of the mechanism of the nondiscrepancy observed for the absolute-preference types is slightly different. Rather than a nondiscrepancy arising from accurately registering the absolute value in the short run, it could arise because the relative-standing rule (e.g., bigger is better) is more predictive of what will happen in the long run for absolute but not relative preference types. If this is the case, it may complicate the preference-types-as-moderator account.⁶ Future research could examine the accessible cognitions at the short-run impact of the change (i.e., Time 1). Is a relative-standing rule accessible for both preference types at the time of the change?

Regardless of whether the perception or a rule is driving the relative assessment in the short run, the question of why one set of preferences has significant absolute value whereas another set lacks it remains. Hsee and colleagues have proposed that psychobiology versus sociocultural factors explain the difference between preference types (Hsee et al., 2008, 2009; Hsee & Zhang, 2010; Tu & Hsee, 2016; cf. Maslow, 1943). That is, one set of preferences (Type A), the more absolute and inherent kind, reflects evolutionarily formed or psycho-biological needs, whereas the other set of preferences (Type B), the more relative and noninherent category, reflects culturally or socially learned tastes.

Type A versus Type B theory would interpret, for example, the effects observed in Experiments 2 and 3 (the e-reader studies) as occurring because changes in the small-size range affect a person's basic needs, such as the need for visual clarity and freedom from eye strain, whereas changes in the large-size range reflect socially formed standards ("The bigger, the better!"). Another related idea suggests preference types are registered differently in the mind; an intuitive and experiential apparatus registers absolute preferences, whereas an analytical and conceptual apparatus registers relative preferences (cf. Morewedge, Gilbert, Myrseth, Kassam, & Wilson, 2010). Understanding the source of preference types is another fruitful avenue for future research that relates to themes of nature versus nurture and intuition versus analysis.

Prescriptive Implications for the Hedonic Durability of Preference Types

The prescriptive lesson from previous research on this topic has been that the long-term effect of most changes (e.g., bigger house, shorter commute) is negligible. Preference types as a moderator of hedonic durability changes this prescriptive outlook and accounts for the full range of findings and observations in this area, namely, sensitization and the other counterexamples discussed in this article's introduction.

Although generalizing from our controlled laboratory studies might be difficult, our results suggest people will be more prone to mispredicting, upgrading, and overinvesting in relative-preference types. People make many choices in a comparative frame of mind (Hsee & Zhang, 2004; Simonson, Bettman, Kramer, & Payne, 2013) that often lead to an evaluation of the difference between two alternatives. Our research implies this strategy is not always

optimal, because when choices are about relative-preference types, contrast effects can overweight a short-term difference in value that will not persist over time (unless the contrast is chronically accessible). This mechanism is likely to be the very one that leads people to get caught on the hedonic treadmill, a constant and insatiable cycle of upgrading one's consumption experiences (Kahneman, 1999; Scitovsky, 1976).

However, preference types as a moderator of hedonic durability suggests a potential prescriptive rule of thumb exists that might help guide predictions and choices about longer-lasting happiness. That rule is defined here as follows: when the initial effect is matched, an upgrade or downgrade to an absolute (i.e., comparison-independent) preference type will have a longer-lasting effect than a relative (i.e., comparison-dependent) preference type. This rule appears to imply people should invest in upgrades and avoid downgrades of absolute-preference types because the happiness or unhappiness derived from a gain or a loss will persist over time.

Yet researchers have cautioned against deriving prescriptions from research that examines only one type of utility (Lichtenstein & Slovic, 2006). The present research focuses exclusively on experience utility, but happiness can come from many sources, including decision utility, transaction utility, experience utility, and remembered utility. These sources can be compensatory, meaning utility from one source can mitigate disutility from another or vice versa (Kahneman, 2006). For example, a consumer who receives a free upgrade to a larger e-reader might experience high transaction utility, yet the comfort of reading on the larger e-reader may be no different from a medium one, if the sizes come from the relative-preference-type range. The consumer could have a net positive assessment of the larger e-reader if she draws on both sources of utility: the positive transaction utility and the neutral experience utility. This dynamic would lead to a different type of hedonic durability pattern for relative preference types than the one demonstrated in the present research.

However, because transaction utility is likely transient, this utility source most likely matters to happiness assessments at the time of the upgrade and would not affect assessments long after the upgrade. Moreover, because transaction utility is not likely unique to one preference type, its effect would be a main effect at the time of the upgrade, simply amplifying, for both

⁶ Another alternative account of the mechanism, which we sought to rule out, is the chronically accessible salient-reference-point story. The idea is that the transition between Time 0 and Time 1 (e.g., between the large and medium e-reader) could serve as a reference point that remains salient not only at the time of the transition, but also long after the transition occurs. The reference point could affect evaluations at Time 1 and Time 2 and could account for the calibration between these time periods. One possibility is that a chronically accessible salient reference point is affecting judgments at Time 2 for the absolute-preference condition but not the relative-preference condition. We sought to rule out this possibility by controlling for the match drop across the conditions; that is, the impact of the transition was not significantly different for the absolute- versus the relative-preference types in any of our studies. However, the transition between the medium to small e-reader, for example, might still have a longer-lasting reference point than the transition between the large and medium e-readers. This version of the mechanism for hedonic durability would be different from the one outlined in the article's general discussion.

types of preferences, the difference between utility at baseline and the upgrade. Transaction utility therefore would not likely interfere with hedonic durability patterns. This analysis of how transaction and experience utility might interact highlights that the compensatory nature of utility assessments would be problematic for our prescription in two cases: one, if a compensatory utility source affects the preference types differently at the time of the change, and/or two, if the compensatory process continues to affect assessments after the change has occurred. Testing these boundary conditions in the laboratory is an important direction for future research.

Another important observation, and perhaps caveat, about the relative and absolute preference types in our research is that on average, the former is the best-liked category and the latter is the least-liked category. That is, the average happiness rating across the time periods is higher for relative-preference types than for its absolute counterpart. For example, in Experiment 2, the large-size e-readers have an average rating of 12.20 ($SD = 4.95$) whereas the small-size e-readers have an average rating of 7.55 ($SD = 5.80$), $t(173) = 5.69$, $p < .001$. A similar pattern holds for Experiment 3. Although this main effect is not discussed in this article, it reveals something interesting about these preference categories: happiness may be more stable for absolute-preference types, but on average, it is lower than its relative-preference-type counterpart. Empirical evidence of an absolute-preference type for which happiness is not only more stable but also on average higher is needed. Such a phenomenon might not exist. Nevertheless, for preference-type research to make useful recommendations for the pursuit of happiness, we need to understand such complexities (cf. Gruber, Kogan, Quoidbach, & Mauss, 2013).

References

- Agresti, A., & Finlay, B. (2009). *Statistical methods for the social sciences* (4th ed.). Upper Saddle River, NJ: Pearson Hall.
- Arkes, H. R., Hirshleifer, D., Jiang, D., & Lim, S. (2008). Reference point adaptation: Tests in the domain of security trading. *Organizational Behavior and Human Decision Processes*, 105, 67–81. <http://dx.doi.org/10.1016/j.obhdp.2007.04.005>
- Bettman, J. R., Luce, M. F., & Payne, J. W. (2008). Preference construction and preference stability: Putting the pillow to rest. *Journal of Consumer Research*, 18, 170–174. <http://dx.doi.org/10.1016/j.jcps.2008.04.003>
- Bonanno, G. A. (2004). Loss, trauma, and human resilience: Have we underestimated the human capacity to thrive after extremely aversive events? *American Psychologist*, 59, 20–28. <http://dx.doi.org/10.1037/0003-066X.59.1.20>
- Brickman, P., Coates, D., & Janoff-Bulman, R. (1978). Lottery winners and accident victims: Is happiness relative? *Journal of Personality and Social Psychology*, 36, 917–927. <http://dx.doi.org/10.1037/0022-3514.36.8.917>
- Cabanac, M. (1981). Physiological signals for thermal comfort. In K. Cena & J. Clark (Eds.), *Bioengineering, thermal physiology and comfort* (pp. 181–192). Amsterdam, the Netherlands: Elsevier. [http://dx.doi.org/10.1016/S0166-1116\(08\)71089-6](http://dx.doi.org/10.1016/S0166-1116(08)71089-6)
- Dhar, R., & Novemsky, N. (2008). Beyond rationality: The content of preferences. *Journal of Consumer Research*, 18, 175–178. <http://dx.doi.org/10.1016/j.jcps.2008.04.004>
- Frederick, S., & Loewenstein, G. (1999). Hedonic adaptation. In D. Kahneman, E. Diener, & N. Schwarz (Eds.), *Well-being: The foundations of hedonic psychology* (pp. 302–329). New York, NY: Russell Sage Foundation Press.
- Gilbert, D. T. (2006). *Stumbling on happiness*. New York, NY: Random House, Inc.
- Gilbert, D. T., Pinel, E. C., Wilson, T. D., Blumberg, S. J., & Wheatley, T. P. (1998). Immune neglect: A source of durability bias in affective forecasting. *Journal of Personality and Social Psychology*, 75, 617–638. <http://dx.doi.org/10.1037/0022-3514.75.3.617>
- Gilbert, D. T., & Wilson, T. D. (2000). Miswanting: Some problems in the forecasting of future affective states. In J. Forgas (Ed.), *Feeling and thinking: The role of affect in social cognition* (pp. 178–197). Cambridge, UK: Cambridge University Press.
- Gruber, J., Kogan, A., Quoidbach, J., & Mauss, I. B. (2013). Happiness is best kept stable: Positive emotion variability is associated with poorer psychological health. *Emotion*, 13, 1–6. <http://dx.doi.org/10.1037/a0030262>
- Hsee, C. K., Xi, F., & Tang, N. (2008). Two recommendations on the pursuit of happiness. *The Journal of Legal Studies*, 37, 115–132. <http://dx.doi.org/10.1086/592259>
- Hsee, C. K., Yang, Y., Li, N., & Shen, L. (2009). Wealth, warmth, and wellbeing: Whether happiness is relative or absolute depends on whether it is about money, acquisition, or consumption. *Journal of Marketing Research*, 46, 396–409. <http://dx.doi.org/10.1509/jmkr.46.3.396>
- Hsee, C. K., & Zhang, J. (2004). Distinction bias: Misprediction and mischoice due to joint evaluation. *Journal of Personality and Social Psychology*, 86, 680–695. <http://dx.doi.org/10.1037/0022-3514.86.5.680>
- Hsee, C. K., & Zhang, J. (2010). General evaluability theory. *Perspectives on Psychological Science*, 5, 343–355. <http://dx.doi.org/10.1177/1745691610374586>
- Kahneman, D. (1999). Objective happiness. In D. Kahneman, E. Diener, & N. Schwarz (Eds.), *Well-being: The foundations of hedonic psychology* (pp. 3–25). New York, NY: Russell Sage Foundation Press.
- Kahneman, D. (2006). New challenges to the rationality assumption. In S. Lichtenstein & P. Slovic (Eds.), *The construction of preference* (pp. 487–503). Cambridge, UK: Cambridge University Press. <http://dx.doi.org/10.1017/CBO9780511618031.028>
- Kahneman, D., & Snell, J. (1992). Predicting a changing taste: Do people know what they will like? *Journal of Behavioral Decision Making*, 5, 187–200. <http://dx.doi.org/10.1002/bdm.3960050304>
- Kivetz, R., Netzer, O., & Schrift, R. (2008). The synthesis of preference: Bridging behavioral decision research and marketing science. *Journal of Consumer Research*, 18, 179–186. <http://dx.doi.org/10.1016/j.jcps.2008.04.005>
- Lichtenstein, S., & Slovic, P. (Eds.). (2006). *The construction of preference*. Cambridge, UK: Cambridge University Press. <http://dx.doi.org/10.1017/CBO9780511618031>
- Mancini, A. D., Bonanno, G. A., & Clark, A. E. (2011). Stepping off the hedonic treadmill. *Journal of Individual Differences*, 32, 144–152. <http://dx.doi.org/10.1027/1614-0001/a000047>
- Maslow, A. H. (1943). A theory of human motivation. *Psychological Review*, 50, 370–396. <http://dx.doi.org/10.1037/h0054346>
- Melamed, L., & Moss, M. K. (1975). The effect of context on ratings of attractiveness of photographs. *The Journal of Psychology: Interdisciplinary and Applied*, 90, 129–136. <http://dx.doi.org/10.1080/00223980.1975.9923935>
- Menzel, P., Dolan, P., Richardson, J., & Olsen, J. A. (2002). The role of adaptation to disability and disease in health state valuation: A preliminary normative analysis. *Social Science & Medicine*, 55, 2149–2158. [http://dx.doi.org/10.1016/S0277-9536\(01\)00358-6](http://dx.doi.org/10.1016/S0277-9536(01)00358-6)
- Morewedge, C. K., Gilbert, D. T., Myrseth, K. O. R., Kassam, K. S., & Wilson, T. D. (2010). Consuming experience: Why affective forecasters overestimate comparative value. *Journal of Experimental Social Psychology*, 46, 986–992. <http://dx.doi.org/10.1016/j.jesp.2010.07.010>
- Pickett, J. (Ed.). (2015). *The American heritage dictionary of the English language* (5th ed.). Boston, MA: Houghton Mifflin Harcourt.

- Riis, J., Loewenstein, G., Baron, J., Jepson, C., Fagerlin, A., & Ubel, P. A. (2005). Ignorance of hedonic adaptation to hemodialysis: A study using ecological momentary assessment. *Journal of Experimental Psychology: General*, 134, 3–9. <http://dx.doi.org/10.1037/0096-3445.134.1.3>
- Schkade, D. A., & Kahneman, D. (1998). Does living in California make people happy? A focusing illusion in judgments of life satisfaction. *Psychological Science*, 9, 340–346. <http://dx.doi.org/10.1111/1467-9280.00066>
- Schwarz, N., & Bless, H. (2007). Mental construal processes: The inclusion/exclusion model. In D. A. Stapel & J. Suls (Eds.), *Assimilation and contrast in social psychology* (pp. 119–141). New York, NY: Psychology Press.
- Scitovsky, T. (1976). *The joyless economy: An inquiry into human satisfaction*. New York, NY: Oxford University Press.
- Simonson, I. (2008). Will I like a medium pillow? Another look at constructed and inherent preferences. *Journal of Consumer Psychology*, 18, 155–169. <http://dx.doi.org/10.1016/j.jcps.2008.04.002>
- Simonson, I., Bettman, J. R., Kramer, T., & Payne, J. W. (2013). Comparison selection: An approach to the study of consumer judgment and choice. *Journal of Consumer Psychology*, 23, 137–149. <http://dx.doi.org/10.1016/j.jcps.2012.10.002>
- Smith, D. M., Loewenstein, G., Jankovic, A., & Ubel, P. A. (2009). Happily hopeless: Adaptation to a permanent, but not to a temporary, disability. *Health Psychology*, 28, 787–791. <http://dx.doi.org/10.1037/a0016624>
- Taylor, S. E. (1983). Adjustment to threatening events: A theory of cognitive adaptation. *American Psychologist*, 38, 1161–1173. <http://dx.doi.org/10.1037/0003-066X.38.11.1161>
- Tu, Y., & Hsee, C. K. (2016). Consumer happiness derived from inherent preferences versus learned preferences. *Current Opinion in Psychology*, 10, 83–86.
- Veenhoven, R. (1991). Is happiness relative? *Social Indicators Research*, 24, 1–34. <http://dx.doi.org/10.1007/BF00292648>
- Weinstein, N. D. (1982). Community noise problems: Evidence against adaptation. *Journal of Environmental Psychology*, 2, 87–97. [http://dx.doi.org/10.1016/S0272-4944\(82\)80041-8](http://dx.doi.org/10.1016/S0272-4944(82)80041-8)
- Wilson, T. D., & Gilbert, D. T. (2005). Affective forecasting: Knowing what to want. *Current Directions in Psychological Science*, 14, 131–134. <http://dx.doi.org/10.1111/j.0963-7214.2005.00355.x>
- Zamble, E. (1992). Behavior and adaptation in long-term prison inmates: Descriptive longitudinal results. *Criminal Justice and Behavior*, 19, 409–425. <http://dx.doi.org/10.1177/0093854892019004005>

Appendix

Future Research

A reviewer suggested this study idea. It proposes a way to manipulate perceived size through chronically accessible contrasts. Figure A1 illustrates the predictions, and its caption describes the study idea in greater detail. Although this experimental manipulation may lead to similar predictions for hedonic durability, the manipulation falls outside the scope of

our paradigm and theory. For our purposes, a contrast must be salient only at the Time 1 rating (and not at Time 2). However, chronically accessible contrasts during consumption experience likely exist in the real world, so this study's results would be interesting and potentially provide another moderator of hedonic durability.

(Appendix continues)



This document is copyrighted by the American Psychological Association or one of its allied publishers. This article is intended solely for the personal use of the individual user and is not to be disseminated broadly.

This document is copyrighted by the American Psychological Association or one of its allied publishers. This article is intended solely for the personal use of the individual user and is not to be disseminated broadly.

This document is copyrighted by the American Psychological Association or one of its allied publishers. This article is intended solely for the personal use of the individual user and is not to be disseminated broadly.

This document is copyrighted by the American Psychological Association or one of its allied publishers. This article is intended solely for the personal use of the individual user and is not to be disseminated broadly.

This document is copyrighted by the American Psychological Association or one of its allied publishers. This article is intended solely for the personal use of the individual user and is not to be disseminated broadly.