



LOUVAIN
RESEARCH INSTITUTE
IN MANAGEMENT
AND ORGANIZATIONS

WORKING PAPER

2018-19

Conditioned to Like Brands and
Mindless Models

Iskra Herak,

Nicolas Kervyn,

Louvain Research Institute in Management and
Organizations

Adrien Mierop,

Université catholique de Louvain,

Matthew Thomson,

Western University

LOUVAIN RESEARCH INSTITUTE IN MANAGEMENT AND ORGANIZATIONS

Louvain Research Institute in Management and Organizations

Working Paper Series

Editor: Prof. Valérie Swaen

(president-lourim@uclouvain.be)

Conditioned to like brands and mindless models

Iskra Herak, Louvain Research Institute in Management and Organizations

Adrien Mierop, Institut de recherche en science psychologique, Université catholique de Louvain

Nicolas Kervyn, Louvain Research Institute in Management and Organizations

Matthew Thomson, Western University

Summary

Two popular tactics of enhancing brands' perceptions are: 1) brand anthropomorphism (i.e. thinking of brands as having humanlike mind, MacInnis & Folkes, 2017) and 2) creating associations by presenting a brand with a person (e.g. a model, Shavitt, Swan, Lowrey, & Wänke, 1994). So far, the focus was solely on the consequences of these practices for brands. Here, the aim is to extend the scope on to both brands *and* models, by looking at the transfer of target's non-evaluative (i.e. mind) and evaluative (i.e. likeability) properties. By integrating these ideas, an important yet neglected question is addressed: to what extent attributing mind to objects (e.g. brands) is interwoven with the decreased mind attribution to people (e.g. models, Epley, Schroeder, & Waytz, 2013). Since pairing an entity (e.g. a brand) that is initially neutral in meaning (i.e. mind) and valence (i.e. likeability) with an entity that has higher amount of these attributes (i.e. a model) leads to change in former's perceptions (e.g. Unkelbach, 2011), then, the brand-model (vs brand-brand) pairing should increase brand's perceived mind and likeability. On the contrary, the theory suggests that the brand-model (vs model-model) pairings should decrease perception of model's mind and likeability. In sum, repetitive exposure to mix, brand-model pairs, i.e. entities differing in perceived mind and likeability, compared to pairs not differing in these properties (brand-brand and model-model pairs) should result in increased mind and likeability of brands and decreased mind and likeability of models. As expected, exposure to mix vs same pairs resulted in increased mind and likeability of brands *and* in decreased mind ratings of models. Contrary to expectations, models' likeability increased after their pairing with brands (vs models). Results are discussed in light of the literature on attribution of humanness and conditioning.

Keywords : Anthropomorphism, dehumanization, attribute transfer, advertising.

JEL Classification: C61, etc.

This work is part of Phd Thesis and financial support from FSR Université catholique de Louvain is gratefully acknowledged.

Corresponding author :

Iskra Herak

CERMA

Louvain Research Institute in Management and Organizations / Campus Louvain la Neuve

Université catholique de Louvain

Address

B-1050 Brussels, BELGIUM

Email : herak.iskra@uclouvain.be

The papers in the WP series have undergone only limited review and may be updated, corrected or withdrawn without changing numbering. Please contact the corresponding author directly for any comments or questions regarding the paper.

President-lourim@uclouvain.be, LouRIM, UCL, 1 Place des Doyens, B-1348 Louvain-la-Neuve, BELGIUM

<https://uclouvain.be/en/research-institutes/lourim>

Introduction

Marketers frequently pursue humanization as a tactic since thinking of things as friends has benefits like greater attachment (Chandler & Schwarz, 2010). Another popular practice is employing actual people due to another form of benefits: the transfer of appealing qualities from model to the brand (McCracken, 1989). Although it is believed that people can help with brand humanization (Fleck, Michel, & Zeitoun, 2014), the idea whether this quality is transferred to the brand remained untested.

Humanization, or anthropomorphism implies reasoning about objects as if they have humanlike mind and qualities (Waytz, Gray, Epley, & Wegner, 2010). It is frequently achieved through direct rhetorical and/or visual cues embedded in the presentation of the object itself (Kim & Kramer, 2015). For example, increased mind perception can be attained through the use of pronouns such as “I, him, her” instead of “it” (Ahn, Kim, & Aggarwal, 2014), or by presenting products in a way that is evoking human schemas (e.g. grids on a car arranged as a smile or as a frown, Aggarwal & McGill, 2007). Even though explicit anthropomorphism of brands can sometimes cause negative reactions (e.g. reluctance to disclose sensitive information; Puzakova, Rocereto, & Kwak, 2013), it remains a prevalent marketing technique with overall positive consequences for brands (for a recent overview see MacInnis & Folkes, 2017). Similarly, ever since studies revealed that advertisings with (vs without) models, have more positive effects for products’ quality perceptions (Kanungo & Pang, 1973), the quest for an optimal match for a brand has been pursued. Person’s qualities like attractiveness, credibility, or fame were identified as important (Till, 1998; Till & Busler, 2000). Attractive (vs unattractive) models are better at enhancing perceptions of brands of cosmetics (Till & Busler, 2000). Credibility of a sports brand is greater if it is endorsed by a professional athlete (vs a professional singer, Temperley &

Tangen, 2006). Endorsement of high-tech products is more effective when the endorser is a non-celebrity rather than a celebrity (Biswas, Biswas, & Das, 2006). Some pairings are not advisable like when controversial individuals (e.g. Britney Spears) paired with mature liked brands (e.g. Clinique) damage perceptions of the brand's quality, and lead to its perception as being more trashy and cheap than previously (Miller & Allen, 2012). While research on brand-model pairing provided important insights into the transfer of attributes between two entities that are repeatedly presented together, the idea whether it can lead to anthropomorphism remained untested.

The transfer of non-evaluative and evaluative properties are robust phenomena, especially when there is a contingency awareness (De Houwer, Thomas, & Baeyens, 2001; Pleyers, Corneille, Luminet, & Yzerbyt, 2007; Unkelbach, 2011). They are recognized as distinct effects, but have similar basis: the formation of a conditioned response (being non-evaluative or evaluative) through repetitive exposure to pairs of unconditional stimuli (US) and conditional stimuli (CS) (De Houwer, 2007; Meersmans, Houwer, Baeyens, Randell, & Eelen, 2005). While the US is perceived as high in meaning and valence the CS is initially neutral on these dimensions or lower compared to the US (De Houwer, 2007). Typical outcome is that CS acquires more of the non-evaluative attribute after repeatedly being paired with the US and becomes more valenced (more positive or negative depending on the valence of US (De Houwer et al., 2001). Hence, after being paired with photos of racing cars brands are perceived as “faster” (Kim, Allen, & Kardes, 1996), “softer” and more liked after being paired with photos of kitties (Meersmans et al., 2005) , or cheaper and disliked after being paired with controversial celebrities (Miller & Allen, 2012). Hence, pairing with models that are inherently more mindful than brands and frequently chosen to be likeable should result in transfer of these non-evaluative (i.e. mind) and evaluative properties (i.e. likeability):

H1a: repetitive exposure to brand-model pairs compared to brand-brand pairs should lead to increase in mind perception of brand.

H1b: repetitive exposure to brand-model pairs compared to brand-brand pairs should lead to increase likability of brand.

While implications for brands are straightforward, what happens to the model? So far, diminished mind perceptions are found in advertisements where the focus is on models' bodies rather than faces, or when the portrayal of the model is sexualizing (Gray, Knobe, Sheskin, Bloom, & Barrett, 2011; Heflick, Goldenberg, Cooper, & Puvia, 2011; Puvia & Vaes, 2013). However, what is the contribution of brands to this phenomenon? The reflection of anthropomorphism is theorized to be the lack of mind attribution to people (Epley et al., 2013). Although the empirical support is still very limited, first insights show that even in simplest ads where people are displayed in a neutral way alongside brands, diminishes mind perceptions and leads to drop in evaluations, regardless of the person's gender or fame (Herak, Kervyn, Thomson, under review). It represents the first support of the idea that there is a collateral effect of brand-model pairing on people promoting brands. However, it is still unclear to what extent increased mind attribution to brands and decreased mind attribution to models are interwoven phenomena.

Implications stemming from theory of anthropomorphism (Epley et al., 2013) coupled with previous overview suggesting that pairing entities differing (vs not differing) in properties leads to greater transfer of attributes (Dimofte & Yalch, 2011), it is hypothesized that:

H2a: repeated pairings of brands and models (compared to model-model pairings) will decrease mind ratings of models

H2b: repeated pairings of brands and models (compared to model-model pairings) will decrease the liking of models.

In two studies, these ideas are tested by comparing the pairing of a brand with a model (i.e. mix-pairs) to the pairing of two brands together and two models together (i.e. same-pairs) while measuring mind and likeability of both targets before and after exposure to pairs (see Figure 1 for hypothesized pattern of results). By doing so, we were able to isolate the impact of the mix-pairing on mind and likeability for each target by keeping constant the frequency of exposure to brands and models.

Study 1

We report how we determined our sample size, all data exclusions, manipulations, and measures in the study (see preregistration). The data from the experiment are available on Open Science Framework (OSF link).

Material

Eight photos of fictional brands in which the chosen brands do not differ on perceived measures of mind and liking were made (see Supplementary material). Likewise, 8 photos in which models within that sample do not differ on perceived measures of mind and likeability were pretested and chosen (see Supplementary material)

Design and Procedure

The 2 (Pairing: Mix-pairs, Same-pairs) $\times$ 2 (Target: Person, Brand) $\times$ 2 (Ratings: Before, After) mixed factorial design, with Pairing as between participants factor, and Target and Ratings as within participants factors was applied. Participants were first asked to rate the 16 targets (8 of

brands and 8 photos of models) on two dependent variables (mind attribution and likeability) using a slide bar (0-“Definitely not” to 100-“Definitely yes”). In a second, exposure phase, depending on the ‘Pairing’ factor, participants saw either 8 mix-pairs (4 brand-model and 4 model-brand pairs) or 8 same-pairs photos (4 model-model pairs and 4 brand-brand). Stimuli order within each phase was randomized as well as the assignment of participants to conditions (see Appendix A). Participants saw a total of 16 pairs, and each pair was presented 6 times during 3000ms. Finally, participants were again asked to rate the targets on mind attribution and likeability.

Results and Discussion

Participants. One hundred seventy-eight participants (113 female, $M_{age}=32.86$) from an online panel completed the survey for a small payment. Participants’ age and gender did not have a significant effect and were excluded from analysis. The distribution of participants across conditions was not significantly different ($\chi^2 = .56, p = .45$; 94 in mix-pair conditions vs. 84 in same-pair conditions).

Dependent Measures and Analysis. Mind scores were averaged by participant and the difference in ratings before and after exposure to pairs was calculated to reflect the change ($D-Mind = Mind_{after_exposure} - Mind_{before_exposure}$). The same was applied for the likeability measure ($D-Like = Like_{after_exposure} - Like_{before_exposure}$). Positive scores indicate an increase in ratings after the exposure to pairs. Negative scores indicate a decrease in ratings after the exposure to pairs. Repeated measures (RM) MANOVA was performed, with Pairing as a between factor (-1: mix-pair and 1: same-pair) and Target (-1: brand and 1: model) and Ratings ($D-Mind$ and $D-Like$) as within factors. The correlation of the two dependent variables was small, but significant ($r(178) = .26, p < .001$).

Hypotheses specify the Pairing $\times$ Target interaction, which was not significant ($F(1, 176) = 1.34, p = .25, \eta^2 = .02$). However, almost all the trends were in the hypothesized direction (see Figure 2). For brands, changes in mind ratings (H1a) and likeability (H1b) were positive but non significant. For models there was a drop in mind ratings (H2a), but there was an increase in likeability (contrary to H2b). Means and standard deviations are represented in Table 1.

To understand lack of significant results for both brands and models, remaining main effects and interactions were inspected from the output. The main effect of Pairing was marginally significant suggesting that regardless of the Target or the Ratings factors, a greater change in ratings occurred for the mix-pairs rather than for the same-pairs ($M_{\text{mix-pair}} = 4.07, SD_{\text{mix-pair}} = 13.98$ vs $M_{\text{same-pair}} = 1.76, SD_{\text{same-pair}} = 14.76, F(1, 176) = 2.93, p = .08, \eta^2 = .02$). The main effect of Target was significant with the higher change for brands compared to models ($M_{\text{Brands}} = 5.19, SD_{\text{Brands}} = 16.46$ vs $M_{\text{Models}} = 0.64, SD_{\text{Models}} = 12.28, F(1, 176) = 15.97, p < .001, \eta^2 = .08$) regardless of the Ratings or the Pairing factors. The main effect of Ratings was not significant ($F(1, 176) = 1.06, p = .24$).

Unexpectedly, the Target $\times$ Ratings interaction was significant ($F(1, 176) = 4.29, p = .04, \eta^2 = .02$). Analysis revealed that for brands, the change in before and after ratings were in the same, positive direction and the difference between D-Mind and D-Like was not significant ($M_{\text{D-Like}} = 5.58, SD_{\text{D-Like}} = 20.37$ vs $M_{\text{D-Mind}} = 5, SD_{\text{D-Like}} = 12.91, F(1, 176) = .221, p = .63$). However, for the models, changes of ratings were operating in diverging directions unlike postulated. While there was a drop in perceived mind there was an increase in likeability and the difference between two measures was significant ($M_{\text{D-Mind}} = -.82, SD_{\text{D-Mind}} = 15.48$ vs $M_{\text{D-Like}} = 2.14, SD_{\text{D-Like}} = 9.04, F(1, 176) = 5.68, p = .02, \eta^2 = .03$). The Pairing $\times$ Ratings interaction was not significant ($F(1, 176) = 1.83, p = .18, \eta^2 = .01$), nor was the triple Pairing $\times$ Target $\times$ Ratings interaction ($F(1, 176) = 1.14, p = .28, \eta^2 = .00$).

The patterns of results show that transfer of non-evaluative and evaluative attributes, i.e., the changes of mind and likeability ratings, are similar after exposure to mix-pairs and same-pairs contrary to expectations. However, this could lie in the procedure itself. The order of presentation of pairs was fully randomized in both conditions, the mix (brand-model) and the same-pairs (brand-brand and model-model pairs). Since the same-pairs appeared one after another randomly, it could potentially mislead participants to create links between pairs of brands and pairs of models in the same-pair condition and lead to the sequential conditioning (Lascelles & Davey, 2006). While this issue was directly address in Study 2, another unexpected result emerged. The transfer of mind and likeability were in the opposite directions when the Target was a model, contrary to H2b. This could mean that the relationship between the transfer of non-evaluative and evaluative qualities is more complex than literature suggests. It could also mean that the likeability is serving as a protective belt against the significant drop in mind ratings. However, before improving the procedure and replicating the results, conclusions regarding robustness and meaning of this finding is limited. Hence in Study 2 the primacy was on establishing an unambiguous procedure for the same-pairing condition.

Study 2

In Study 2 a more distinguishing difference between the mix-pairs and the same-pairs was built. Two clearly distinct blocks of same-pairs with counterbalanced order of presentation was made. Some participants were first exposed to all brand-brand pairs, and then to all model-model pairs, and others first saw model-model pairs followed by brand-brand pairs. Stimuli order within each block remained fully randomized. Everything else remained the same as in Study 1 (see Study 1 Design and Procedure for details).

Results and Discussion

Participants. Three hundred eleven participants (159 “female”, 150 “male”, 2 “other”, $M_{age}=34.82$) completed the survey from the online pool for a small remuneration. Participants’ age and gender did not have a significant effect and were excluded from further analysis. The distribution of participants across pairing conditions was not statistically different (168 in mix-pair and 143 in same-pairs, $\chi^2=2.10$, $p = .16$).

Dependent Measures and Analysis. As previously mind scores were averaged by participants and D-Mind was calculated. Likewise, likeability scores were averaged by participants then D-Like was computed. The correlation between the two dependent measures was moderate and significant ($r(311)= 0.45$, $p < .01$). RM-MANOVA with Pairing as between factor and Target and Ratings as within factors was performed as previously.

Hypothesized Pairing $\times$ Target interaction emerged as significant ($F(1,309) = 16.89$, $p < .001$, $\eta^2 = .05$, see Table 2). With respect to brands the support for both H1a and H1b hypothesis was obtained: the perceived change in mind was higher in mix-pairs compared to same-pairs ($F_{D_Mind}(1, 309) = 12.52$, $p < .001$, $\eta^2 = .04$). Likeability of brands also increased more after exposure to mix-pairs vs same-pairs ($F_{D_Like}(1, 309) = 16.79$, $p < .001$, $\eta^2 = .05$). The difference in ratings of brands was significantly higher in the mix compared to the same-pair condition ($M_{mix-pair} = 6.61$ $SD_{mix-pair} = 16.55$ vs $M_{same-pair} = .05$ $SD_{same-pair} = 10.62$; $F(1,309) = 18.56$, $p < .001$, $\eta^2 = .06$). With regard to models, the change in mind ratings was not statistically significant across the pairing conditions, although it was in hypothesized negative direction (H2a) ($F_{D_Mind}(1, 309) = 1.95$, $p = .16$, $\eta^2 = .01$, see Table 2). The difference in likeability of models was also not different between the pairing conditions ($F_{D_Like}(1, 309) = 2.79$, $p = .09$, $\eta^2 = .01$). Contrary to H2b, as in Study 1, although non significant, the change of likeability of models was positive after exposure to mix-pairs (see Figure 3). The difference in ratings of models did not vary as a function of the

Pairing factor ($M_{\text{mix-pair}} = -.04$ $SD_{\text{mix-pair}} = 14.55$ vs $M_{\text{same-pair}} = -.02$ $SD_{\text{same-pair}} = 12.35$; $F(1,309) = 0.28$, $p = .86$, $\eta^2 = .00$).

As previously the remaining main effects and interactions were inspected from the output. Main effect of Pairing was significant: greater changes occurred in the mix-pair compared to the same-pair conditions ($M_{\text{mix-pair}} = 3.11$ $SD_{\text{mix-pair}} = 15.55$ vs $M_{\text{same-pair}} = 0.17$ $SD_{\text{same-pair}} = 11.49$; $F(1, 309) = 7.1$, $p < .01$, $\eta^2 = .02$) regardless of Ratings and Target factors as in Study 1. Likewise, the main effect of Target was significant and the change in ratings was higher for brands compared to models ($M_{\text{Brands}} = 3.56$ $SD_{\text{Brands}} = 13.59$ vs $M_{\text{Models}} = -0.28$ $SD_{\text{Models}} = 13.45$; $F(1, 309) = 25.10$, $p < .001$, $\eta^2 = .08$) regardless of the Ratings or the Pairing factors. The main effect of Ratings was not significant ($F(1, 309) = 2.74$, $p = .09$).

As before, the Target $\times$ Ratings interaction was significant ($F(1, 309) = 7.49$, $p = .01$, $\eta^2 = .02$). The significant difference in perceived change of mind and perceived change in likeability occurred for models ($M_{\text{D-Mind}} = -1.62$ $SD_{\text{D-Mind}} = 15.79$ vs $M_{\text{D-Like}} = 1.05$ $SD_{\text{D-Like}} = 11.11$; $F(1, 309) = 8.2$, $p < .01$, $\eta^2 = .03$). For the brands this difference was not significant ($F(1, 309) = 7.49$, $p = .01$, $\eta^2 = .02$). Again the Pairing $\times$ Ratings interaction was not significant ($F(1, 309) = 2.75$, $p = .10$, $\eta^2 = .01$). The triple Pairing $\times$ Target $\times$ Ratings interaction emerged significant ($F(1, 309) = 4.72$, $p = .03$, $\eta^2 = .02$). For brands the difference between increases of mind and likeability ratings did not vary within the pairing conditions ($F_{\text{mix-pair}}(1, 309) = .62$, $p = .43$, $\eta^2 = .00$ and $F_{\text{same-type}}(1, 309) = .09$, $p = .76$, $\eta^2 = .00$, respectively).

For models only in the mix-pairing condition the changes in mind ratings and the changes in likeability operated in different directions: while there was a drop in the mind, the likeability of models increased ($F(1, 309) = 15.71$, $p < .001$, $\eta^2 = .05$). This difference was not significant within the same-pairing ($F(1, 309) = .06$, $p = .81$, $\eta^2 = .00$).

Additional analysis

So far in both studies the pattern of changes in perceptions of models is more complex than initially hypothesized. The drop in both mind (H2a) and likeability ratings (H2b) was expected to occur after their pairing with a brand (i.e. mix-pairing) vs pairing with another model (i.e. same-pairing). However, as in Study 1 these measures operated in opposing direction: the drop in mind was followed by an increase in likeability. In Study 1 the fully randomized order of presenting same-pairs in one block (brand-brand & model-model) was potentially misleading participants to create links between brand-brand and model-model pairs unlike wanted. This limitation was successfully addressed in Study 2. Nonetheless, even though the procedure was improved in Study 2, ratings of models operated in opposed directions in mix-pairs. One might argue that the increase in likeability is protecting or masking the drop in mind ratings of models in mix-pairs.

To test this assumption, an ANOVA with D-Like as a covariate, D-Mind as a dependent variable (DV) and Pairing as between factors was performed. When, statistically controlling for change in likeability the observed drop in mind was significantly lower for models in mix-pairs compared to models in same-pairs as initially hypothesized (H2a) ($F(1, 308) = 4.10, p = .04, \eta^2 = .02$). This drop was significantly different from zero ($t_{D-Mind}(168) = -2.14, p = .03$). The same was applied with D-Like as the DV and D-Mind as covariate. The change in likeability when statistically controlling for change in mind ratings, was positive and higher for models in mix-pair conditions compared to models in same-pair conditions ($F(1, 308) = 4.94, p = .03, \eta^2 = .02$) contrary to hypothesis (H2b). Another possibility is that the increase in likeability is a result of the exposure effect (Zajonc, 2001). If it was true, then in both pairing conditions the observed changes would be significantly different from the value of zero. However, the increase in

likeability of models was significantly different than zero only in mix-pairs ($t_{D-Like}(168) = 2.36, p = .02$), and not in same-pairs ($t_{D-Like}(142) = -.01, p = .98$). This suggests that significant increase in likeability cannot be accounted by the mere exposure but rather to the pairing with a brand.

Discussion

Although it is known that anthropomorphism and pairing brands with people enhance brands' perception (Aggarwal & McGill, 2007; Shavitt et al., 1994), these ideas have been pursued separately. Since transfer of properties is possible when two entities initially differing in non-evaluative and evaluative attributes are repeatedly paired together (Dimofte & Yalch, 2011) it follows that the brand-model (vs brand-brand) pairing can lead to both anthropomorphism (H1a) and increased likeability (H1b). Additionally, suggested flip side of anthropomorphism-decreased mind perception of people (Epley et al., 2013) and consequences for people after brand-model pairing have been neglected issues in marketing. In an effort to integrate these ideas with previous findings, the second hypothesis was that the brand-model pairing (vs model-model pairing) will have inverse, negative consequences for perception of models i.e. decreased mind (H2a) and likeability (H2b).

When it comes to brands our findings suggest that the brand-model pairings are indeed more beneficial than brand-brand pairing, leading to both anthropomorphism (i.e. increase in mind) and greater appreciation (i.e. increase in likeability). This is the first empirical support of the idea that “attaching a human” humanizes the brand. Furthermore, repeated association with brands does lead to decreased mind (H2a), but unexpectedly increase likeability of models. Contrary to expectations, findings suggests that decreased mind perception do not undermine likeability. Moreover, brand-model pairing increases likeability of models as well. Within the broader context of research on dehumanization a study showed that denial of human uniqueness

or human nature isn't always detrimental for the target (Martínez, Rodríguez-Bailón, Moya, & Vaes, 2015). That is, respondents do not reject to interact with dehumanized target in some contexts (e.g. approaching animal-like group in a party, or working with robot-like groups Martínez et al., 2015). It could be that the context of advertisement leads to decreased mind perception of models but paradoxically renders them more likeable. Future research should address whether in advertisement people gain some quality that potentially increases the likeability despite the lack of perceived mind? Recently it has been shown that the transfer of properties is possible even when the targets are people: perception of person's athleticism increases after being paired with a photo of another person doing sport (Förderer & Unkelbach, 2012). It is plausible that models gained in object-likeness, which in advertisement context represents an appealing quality. Although mind is the first step toward recognizing someone's humanness and acknowledging moral treatment (Epley et al., 2013), it is restrictive to understanding what other aspect of humanness are also potentially decreased. Despite the limitation of chosen measures, the finding that advertisement seemingly harmless (simple presentation of a brand next to the model) decreases mind perception but increases likeability is a novel one and represents a contribution to conditioning literature as well.

From the perspective of non-evaluative and evaluative transfer, it is also interesting to explore what happens to the US after the pairing procedure, since this question is also quite neglected (for a notable exception, see Mierop, Bret, Yzerbyt, Dumas, & Corneille, under review). To our best knowledge, the studies constitute the first evidence (1) that not only the conditioned response but also the (non-evaluative) unconditioned response is changed through a conditioning procedure, and (2) that two responses can be changed in a dissociative pattern.

Willingly or not, the public is frequently exposed to faces and logos on billboards, print ads, or more recently online banners. Eager to understand what traits a person has to have in order to benefit the brand, researchers neglected to focus on consequences such pairing has on a person. Today when everyone is a potential model, it looks like a question worthy of exploring. Social media such as YouTube or Instagram, allow people to gain their living by testing and recommending products and brands, and this profession of “social influencers” has a steady rise¹. Hence our interest to study what consequences such pairing has on perception of people comes from both theory and everyday life. Yet, very little is known how such practices influence our perception of what is common for all these people: humanness. We hope to raise interest among researchers regarding this topic.

¹ Typing « social influencer » in Google Trends shows that in last five years this phrase has been search from 2 times per week to 79 times per week making it a real global trend
<https://trends.google.fr/trends/explore?date=today%205-y&q=social%20influencer>

References

- Aggarwal, P., & McGill, A. L. (2007). Is that car smiling at me? Schema congruity as a basis for evaluating anthropomorphized products. *Journal of Consumer Research*, 34(4), 468–479.
- Ahn, H.-K., Kim, H. J., & Aggarwal, P. (2014). Helping Fellow Beings Anthropomorphized Social Causes and the Role of Anticipatory Guilt. *Psychological Science*, 25(1), 224–229. <https://doi.org/10.1177/0956797613496823>
- Biswas, D., Biswas, A., & Das, N. (2006). The Differential Effects of Celebrity and Expert Endorsements on Consumer Risk Perceptions. The Role of Consumer Knowledge, Perceived Congruency, and Product Technology Orientation. *Journal of Advertising*, 35(2), 17–31. <https://doi.org/10.1080/00913367.2006.10639231>
- Chandler, J., & Schwarz, N. (2010). Use does not wear ragged the fabric of friendship: Thinking of objects as alive makes people less willing to replace them. *Journal of Consumer Psychology*, 20(2), 138–145. <https://doi.org/10.1016/j.jcps.2009.12.008>
- De Houwer, J. (2007). A Conceptual and Theoretical Analysis of Evaluative Conditioning. *The Spanish Journal of Psychology*, 10(02), 230–241. <https://doi.org/10.1017/S1138741600006491>
- De Houwer, J., Thomas, S., & Baeyens, F. (2001). Association learning of likes and dislikes: A review of 25 years of research on human evaluative conditioning. *Psychological Bulletin*, 127(6), 853.
- Dimofte, C. V., & Yalch, R. F. (2011). The mere association effect and brand evaluations. *Journal of Consumer Psychology*, 21(1), 24–37. <https://doi.org/10.1016/j.jcps.2010.09.005>
- Epley, N., Schroeder, J., & Waytz, A. (2013). Motivated Mind Perception: Treating Pets as People and People as Animals. In S. J. Gervais (Ed.), *Objectification and*

(De)Humanization (Vol. 60, pp. 127–152). New York, NY: Springer New York.

Retrieved from http://link.springer.com/10.1007/978-1-4614-6959-9_6

Fabricant, S. M., & Gould, S. J. (1993). Women's Makeup Careers: An Interpretive Study of Color Cosmetic Use and "Face Value." *Psychology and Marketing*, 10(6). Retrieved from <https://search.proquest.com/docview/1308067689/citation/FA3F9B2C68124C45PQ/1>

Fleck, N., Michel, G., & Zeitoun, V. (2014). Brand Personification through the Use of Spokespeople: An Exploratory Study of Ordinary Employees, CEOs, and Celebrities Featured in Advertising. *Psychology & Marketing*, 31(1), 84–92.
<https://doi.org/10.1002/mar.20677>

Förderer, S., & Unkelbach, C. (2012). Hating the cute kitten or loving the aggressive pit-bull: EC effects depend on CS–US relations. *Cognition & Emotion*, 26(3), 534–540.

Gray, K., Knobe, J., Sheskin, M., Bloom, P., & Barrett, L. F. (n.d.). More than a body: Mind perception and the nature of objectification. *Journal of Personality and Social Psychology*, 101(6), 1207–1220. <http://dx.doi.org/10.1037/a0025883>

Heflick, N. A., Goldenberg, J. L., Cooper, D. P., & Puvia, E. (2011). From women to objects: Appearance focus, target gender, and perceptions of warmth, morality and competence. *Journal of Experimental Social Psychology*, 47(3), 572–581.
<https://doi.org/10.1016/j.jesp.2010.12.020>

Kanungo, R. N., & Pang, S. (n.d.). Effects of human models on perceived product quality. *Journal of Applied Psychology*, 57(2), 172–178. <http://dx.doi.org/10.1037/h0037042>

Kim, H. C., & Kramer, T. (2015). Do materialists prefer the "brand-as-servant"? The interactive effect of anthropomorphized brand roles and materialism on consumer responses. *Journal of Consumer Research*, 42(2), 284–299.

- Kim, J., Allen, C. T., & Kardes, F. R. (1996). An investigation of the mediational mechanisms underlying attitudinal conditioning. *Journal of Marketing Research*, 318–328.
- Lascelles, K. R. R., & Davey, G. C. L. (2006). Successful differential evaluative conditioning using simultaneous and trace conditioning procedures in the picture–picture paradigm. *Quarterly Journal of Experimental Psychology*, 59(3), 482–492.
<https://doi.org/10.1080/02724990444000140>
- MacInnis, D. J., & Folkes, V. S. (2017). Humanizing brands: When brands seem to be like me, part of me, and in a relationship with me. *Journal of Consumer Psychology*, 27(3), 355–374. <https://doi.org/10.1016/j.jcps.2016.12.003>
- Martínez, R., Rodríguez-Bailón, R., Moya, M., & Vaes, J. (2015). Interacting with dehumanized others? Only if they are objectified. *Group Processes & Intergroup Relations*, 1368430215612219.
- McCracken, G. (1989). Who is the celebrity endorser? Cultural foundations of the endorsement process. *Journal of Consumer Research*, 310–321.
- Meersmans, T., Houwer, J. D., Baeyens, F., Randell, T., & Eelen, P. (2005). Beyond evaluative conditioning? Searching for associative transfer of nonevaluative stimulus properties. *Cognition & Emotion*, 19(2), 283–306. <https://doi.org/10.1080/02699930441000328>
- Miller, F. M., & Allen, C. T. (2012). How does celebrity meaning transfer? Investigating the process of meaning transfer with celebrity affiliates and mature brands. *Journal of Consumer Psychology*, 22(3), 443–452. <https://doi.org/10.1016/j.jcps.2011.11.001>
- Morewedge, C. K., Preston, J., & Wegner, D. M. (2007). Timescale bias in the attribution of mind. *Journal of Personality and Social Psychology*, 93(1), 1–11.
<http://dx.doi.org/10.1037/0022-3514.93.1.1>

- Pleyers, G., Corneille, O., Luminet, O., & Yzerbyt, V. (2007). Aware and (Dis)Liking: Item-Based Analyses Reveal That Valence Acquisition via Evaluative Conditioning Emerges Only When There Is Contingency Awareness. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 33(1), 130–144. <https://doi.org/10.1037/0278-7393.33.1.130>
- Puvia, E., & Vaes, J. (2013). Being a Body: Women's Appearance Related Self-Views and their Dehumanization of Sexually Objectified Female Targets. *Sex Roles*, 68(7–8), 484–495. <https://doi.org/10.1007/s11199-012-0255-y>
- Puzakova, M., Rocereto, J. F., & Kwak, H. (2013). Ads are watching me. *International Journal of Advertising*, 32(4), 513–538. <https://doi.org/10.2501/IJA-32-4-513-538>
- Shavitt, S., Swan, S., Lowrey, T. M., & Wänke, M. (1994). The interaction of endorser attractiveness and involvement in persuasion depends on the goal that guides message processing. *Journal of Consumer Psychology*, 3(2), 137–162. [https://doi.org/10.1016/S1057-7408\(08\)80002-2](https://doi.org/10.1016/S1057-7408(08)80002-2)
- Temperley, J., & Tangen, D. (2006). The Pinocchio factor In Consumer Attitudes Towards Celebrity Endorsement: Celebrity Endorsement. The Reebok Brand, And An Examination Of A Recent Campaign. *Innovative Marketing*, 2(3), 97–111.
- Till, B. D. (1998). Using celebrity endorsers effectively: lessons from associative learning. *Journal of Product & Brand Management*, 7(5), 400–409.
- Till, B. D., & Busler, M. (2000). The match-up hypothesis: Physical attractiveness, expertise, and the role of fit on brand attitude, purchase intent and brand beliefs. *Journal of Advertising*, 29(3), 1–13.
- Unkelbach, C. (2011). Beyond evaluative conditioning! Evidence for transfer of non-evaluative attributes. *Social Psychological and Personality Science*, 479–486.

Waytz, A., Gray, K., Epley, N., & Wegner, D. M. (2010). Causes and consequences of mind perception. *Trends in Cognitive Sciences*, 14(8), 383–388.

<https://doi.org/10.1016/j.tics.2010.05.006>

Zajonc, R. B. (2001). Mere Exposure: A Gateway to the Subliminal. *Current Directions in Psychological Science*, 10(6), 224–228.

Appendix

Table 1. Study 1: means and standard deviations of change in perceived liking (D-Like) and mind (D-Mind) after the pairing procedure, with positive scores indicating an improvement, and negative a decrease.

Target	Pairing	Ratings			
		D-Like		D-Mind	
		M	SD	M	SD
Model	mix-pair	2.86	7.91	-0.59	14.21
	same-pair	1.35	10.14	-1.07	16.87
Brands	mix-pair	7.85	13.50	6.17	20.31
	same-pair	1.82	11.49	4.94	20.54

Note: Heteroskedasticity was examined with Q-Q normality plots and Test Glejser; and both indicators suggest that there is no significant violation of the normality of residuals, hence allowing analysis of variance (see Garson, 2012).

Table 2. Study 2: means and standard deviations of change in perceived liking (D-Like) and mind (D-Mind) after the pairing procedure, with positive scores indicating an improvement, and negative a decrease

Target	Pairing	Rating			
		D-Like		D-Mind	
		M	SD	M	SD
Model	mix-pair	2.11	11.56	-2.89	17.53
	same-pair	-0.01	10.66	-0.34	14.04
Brands	mix-pair	6.18	14.65	7.03	18.46
	same-pair	0.34	9.46	0.69	11.79

Note: Heteroskedasticity was examined with Q-Q normality plots and Test Glejser; and both indicators suggest that there is no significant violation of the normality of residuals, hence allowing the analysis of variance (see Garson, 2012).

Figure 1. Hypothesized pattern of results, indicating an increase in Likeability and Mind ratings of brands after being paired with models (vs brands) and a decrease for models after being paired with brands (vs models).

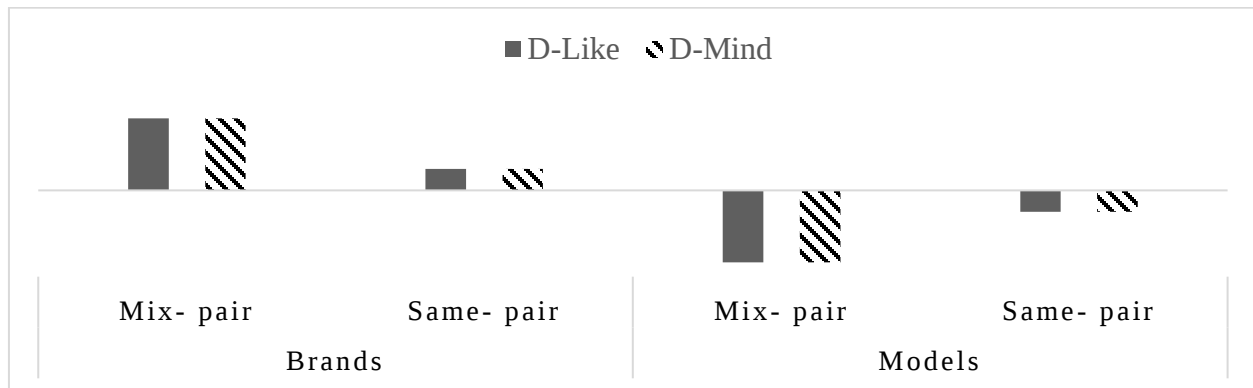
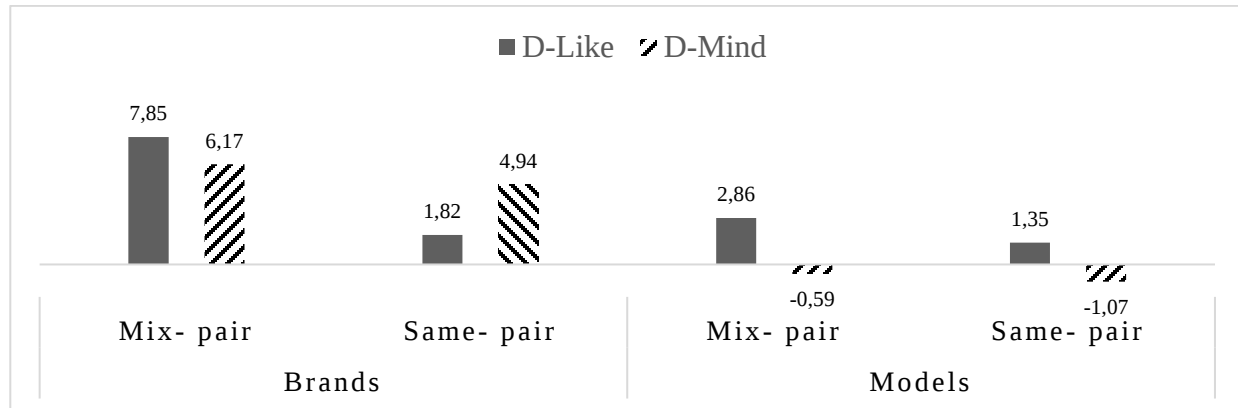


Figure 2. Change in ratings (Likeability and Mind scores) after the pairing procedure suggesting greater change in ratings in mix vs same-pairings.



Note: Although not significant, pattern of means indicates changes in hypothesized direction for brands (an increase in both ratings), while for models the expected decrease in mind is followed by unexpected increase in liking.

Figure 3. Change in Likeability and Mind score after Pairing procedure indicating patterns in hypothesized direction for Brands and unexpected divergence in ratings for Models.



Appendix A

Instrument

Directions:		The experiment follows three parts. You will receive instructions before starting each of these parts. It is normal that some parts seem repetitive.			
Phase 1 Before Treatment Evaluations	We are interested in your perceptions of the photos that will be displayed on the screen. For each photo slide the bar below.				
	To what extent would you say that this person[brand] appears to have a MIND From 0 "definitely not" to 100 "definitely yes"				
	To what extent would you say that you LIKE this person [brand] From 0 "definitely not" to 100 "definitely yes"				
Phase 2 Treatment	The next component of this study deals with visual perception. For this purpose you will be presented with pairs of pictures that will appear on the screen. Please look at them carefully.				
	The visual perception task will take approximately 2.5 minutes. It is VERY IMPORTANT that you pay close attention to the picture pairs throughout the entire task. When you are ready to start please click on START.				
Phase 3 After Treatment Evaluations	We are interested in your perceptions of the photos that will be displayed on the screen. For each photo slide the bar below.				
	To what extent would you say that this person[brand] appears to have a MIND From 0 "definitely not" to 100 "definitely yes"				
	To what extent would you say that you LIKE this person [brand] From 0 "definitely not" to 100 "definitely yes"				
Stimuli 8 Photos Of Models & 8 Photos Of Brands	1		1		
	2		2		
	3		3		
	4		4		
	5		5		
	6		6		
	7		7		
	8		8		

Note. All photos were pre-tested see supplementary material.

Supplementary material.

Pre-test

Stimuli selection

To ensure that any observed change in mind and liking ratings is due to the pairing procedure and not due to *a priori* variations from the stimuli, we pre-tested and selected photos that correspond to needed criteria. Specifically, we created two sub-samples of photos (i.e. one of brands and one of models). We ensured that within the sample of brands, there is no differences in mind and liking ratings before the procedure. Likewise, we ensured that models within the sample of models do not differ on mind and liking measure before the conditioning procedure. Initial sample of 16 fictive brands was created with use of free websites that generate random brand names (5-7 letters) and random logos. Initial sample of 16 photos of female models was chosen from the free database of models looking at the camera with a neutral/natural looking make-up (because strong makeup can be perceived negatively Fabricant & Gould, 1993).

Participants

Twenty-nine respondents (13 female, $M_{age}=33.38$, $SD=12.65$) recruited from the online pool completed the questionnaire for a small remuneration. Age and gender were not significant thus were dropped from the analysis.

Design and procedure

We used a 2 (Photo: Brands, Persons) by 3 (Ratings: Mind, Like, Appeal,) within participants design. Participants were randomly exposed to 32 photos and were asked to move a slide bar (0 "definitely not" to 100 "definitely yes") and rate each photo on three measures (To what extent would you say that this person[brand] appears to have a MIND; To what extent

would you say that you LIKE this brand [person]; To what extent would you say that you find this brand [person] APPEALING). The scale was adapted from Morewedge, Preston, and Wegner (2007).

In the pre-test we included a measure of the appeal to additionally ensure that brands and models within their designated subsamples do not vary with regards to appeal. Because, the liking and appeal measures were almost redundant ($r(29) = .83, p < .01$), and liking has more evaluative connotation than appeal, for studies we kept only the liking measure. Liking and mind were moderately correlated ($r(29) = .64, p < .01$), as well as mind and appeal ($r(29) = .54, p < .01$). We run repeated measures ANOVA with Photos (16 of Brands or 16 of Persons) and 3 measures (Appeal, Like, Mind) as within factors.

Brands. The main effect of Photos was not significant $F(1, 28) = 0.39, p = .84, \eta^2 = .001$, the main effect of Measure was marginally significant ($F(2, 27) = 3.02, p = .07, \eta^2 = .18$) as well as the Photos by Measure interaction ($F(1, 28) = 5.06, p = .03, \eta^2 = .15$). Post hoc tests using the Bonferroni correction yielded significant differences in liking and appeal ratings of several brands (B1, B8, B12, B15) so they were excluded from the sub-sample (see Table 3). There were no differences in mind perception of the brands.

Table 3. Means and standard deviations on three measures of initial sample of 16 photos of brands.

Photo Brands	Measure					
	Mind		Like		Appeal	
	M	SD	M	SD	M	SD
B1*	20.27	20.56	24.37*	20.53	31.59	26.16
B2	25.79	25.62	32.72	23.82	39.90	25.06
B3	30.55	28.34	40.52	26.21	38.93	27.11
B4	32.37	28.74	33.90	26.25	39.00	28.92
B5	33.41	28.74	36.17	23.91	38.31	22.87
B6	29.59	28.62	34.72	27.43	32.93	29.03
B7	36.07	32.73	40.93	28.74	45.21	32.03
B8*	33.76	29.67	39.69*	23.93	45.21	25.40
B9	25.59	25.89	28.41	24.25	27.86	24.49
B10	32.21	25.53	35.69	22.00	35.97	24.59
B11	30.21	28.33	36.93	24.65	37.76	26.77
B12*	27.00	26.73	26.10	24.17	24.72*	21.70
B13	29.03	27.88	28.31	22.04	32.07	23.76
B14	27.66	25.41	34.31	25.97	32.03	26.74
B15*	33.38	26.48	42.62*	25.18	51.31	28.27
B16	30.21	28.57	32.79	23.28	31.41	21.89

Note: *indicates photos of brands that were excluded from consideration due to significant difference at the level of $p=.05$ with another brand from the sample on one or more measures.

Models. The main effect of Photos was not significant $F(1, 28) = 0.39, p = .84, \eta^2 = .001$, the main effect of Measure was significant ($F(2, 27) = 8.16, p = .03, \eta^2 = .37$). Post hoc Bonferroni correction revealed that mind scores (73.11 ± 23.5) compared to liking (62.18 ± 20.74) and appeal ratings (65.32 ± 20.95) were significantly higher ($p=.02$, and $p<.01$ respectively). The Photos x Measure interaction was significant as well ($F(1, 28) = 3.89, p = .05, \eta^2 = .12$). Several models (P5, P6, P14, P15, P16) differed from others on measures of liking and appeal so they were excluded from the sub-sample (see Table 4). There were no differences in mind ratings among models.

Table 4. Means and standard deviations on three measures of initial sample of 16 photos of models.

Photos of Models	Measure					
	Mind		Like		Appeal	
	M	SD	M	SD	M	SD
P1	68.97	28.58	58.66	22.43	63.41	22.33
P2	77.07	19.05	63.31	22.24	68.24	22.00
P3	74.21	25.15	64.41	20.26	70.66	20.56
P4	71.34	23.68	64.72	21.10	65.90	22.47
P5*	68.45	26.09	54.79*	18.91	48.93*	19.26
P6*	69.72	29.81	64.03	24.89	71.79*	22.53
P7	76.79	18.11	66.69	19.20	68.34	19.06
P8	72.00	23.85	64.52	20.71	66.72	21.91
P9	73.79	21.83	60.83	19.73	64.31	18.03
P10	75.79	24.61	60.83	24.99	62.14	22.50
P11	72.83	22.69	62.90	20.88	65.21	19.79
P12	77.24	18.89	56.38	20.16	59.21	27.98
P13	69.93	25.09	58.93	23.25	62.41	21.70
P14*	67.59	26.36	55.28*	21.01	52.52*	23.76
P15*	76.66	22.77	68.17*	16.81	67.28	17.14
P16*	77.38	19.46	68.93*	15.30	74.34*	14.22

Note: *indicates photos of models that were excluded from consideration due to significant difference at the level of $p=.05$, from another model from the sample on one or more measures.

After excluding those photos that were significantly different from others within the subsample, we randomly chose 16 (8 of brands and 8 of persons). Final sample used in studies is in the Appendix A.