

Transcending Embarrassment: On The Reputational Benefits of Laughing at Yourself

SUPPLEMENTARY MATERIALS

A: SCENARIO WORDINGS2

B: DETAILS ON DEPENDENT VARIABLES5

C: ATTENTION CHECK MEASURES7

D: AI STIMULI STUDY 2.....9

E: DEVIATIONS FROM PRE-REGISTRATIONS & ADDITIONAL ANALYSIS.....11

F: GENDER INTERACTIONS.....32

G: SUPPLEMENTARY STUDY S1.....33

A: SCENARIO WORDINGS

STUDY 1 & POST-TEST

Imagine that you go to a dinner party at a friend's house on Sunday. You start chatting with [name], a cousin of your friend whom you have never met before.

During the conversation, you ask [name] how [his/her] weekend was. [name] then tells you the following story:

Note: See the scenarios presented in the Appendix of the manuscript.

Display Manipulation

EMBARRASSMENT DISPLAY: As [he/she] is telling you the story, [he/she] tends to look away and avoids eye contact with you. You also realize that [his/her] face is turning red.

AMUSEMENT DISPLAY: As [he/she] is telling you the story, [he/she] smiles and laughs at [himself / herself] about what happened. [He/She] also makes fun of [himself / herself] and you realize that [he/she] finds the situation funny.

STUDIES 2 & 5

Imagine that you go to a dinner party at a friend's house on Sunday.

You start chatting with [name], a cousin of your friend whom you have never met before.

During the conversation, you ask [name] if [he/she] did anything fun this weekend.

Note: See the scenarios presented in the Appendix of the manuscript.

STUDY 3

Note: Disclosure Scenarios are the same as those in Studies 2 & 5

No Disclosure Scenarios

Imagine that you go to a dinner party at a friend's house on Sunday.

You start chatting with [name], a cousin of your friend whom you have never met before.

During the conversation, you ask [name] if [he/she] did anything fun this weekend.

HIKING: [name] tells you that [he/she] went hiking with [his/her] sister and a large group of [his/her] sister's friends yesterday. [He/She] says that [he/she] had a lot of fun during the hike.

THEATRE: [name] tells you that [he/she] went to see a famous show at a large theatre yesterday. [He/She] says that [he/she] had a lot of fun during the theatre.

PARTY: [name] tells you that [he/she] went to a colleague's housewarming party yesterday. The party was outside on the patio. [He/She] says that [he/she] had a lot of fun during the party.

YOGA: [name] tells you that [he/she] went to a one-day yoga retreat yesterday. During lunch break, [he/she] met a group of people, and they joined the workshop together in the afternoon. [He/She] says that [he/she] had a lot of fun during the retreat.

Note: Display manipulation is the same as in Study 1.

STUDIES 4 & S1

Scenarios

FALL: Imagine that you are walking on a crowded public street. Suddenly, you notice that the person walking in front of you catches [his/her] foot on the sidewalk. [He/She] then awkwardly flings [his/her] body forward barely preventing [himself / herself] from falling over. Many people see [him/her] trip.

SNORE: Imagine that you are sitting in a crowded theater. You notice that the person next to you accidentally falls asleep. Before [he/she] realizes it, [he/she] snores loudly which causes [him/her] to wake [himself / herself] up and see that many people are looking at [him/her].

BEGGAR: Imagine that you are taking a walk in a park. You notice that the person in front of you approaches a shabby-looking man who is sitting on a bench and holding out a cup. [He/She] walks up to this man and puts some coins in his cup; but it turns out that this man was not begging, he was just a man sitting on the bench drinking coffee.

GLASS: Imagine that you are in a restaurant having dinner. You notice that the person in the next table reaches to grab some bread, but [he/she] accidentally knocks over the water glass causing it to break on the floor with a crash.

ZIPPER: Imagine that you attend a presentation that has a large audience. You see that the zipper of the presenter is down during the presentation. The presentation goes well, but at the end the presenter looks down and notices that [his/her] zipper has been down the whole time.

GLASS DOOR: Imagine that you are at a friend's house for a party. The party is outside on the patio. In the middle of the party, you notice that the person in front of you attempts to go inside. [He/She] doesn't notice that the sliding glass door is closed and walks right into it causing it to make a very loud rumbling noise. Everyone at the party looks over and sees what [he/she] has done.

Display Manipulation

EMBARRASSMENT DISPLAY: Just at that moment, you catch each other's eye with this person. You see that [he/she] tends to look away and avoid eye contact with you. You also realize that [his/her] face is turning red.

AMUSEMENT DISPLAY: Just at that moment, you catch each other's eye with this person. You see that [he/she] smiles and laughs at [himself / herself] about what happened. You also realize that [he/she] finds the situation funny.

NEUTRAL DISPLAY: Just at that moment, you catch each other's eye with this person. You see that [he/she] does not show any reaction. You also realize that [he/she] acts as if nothing has happened.

STUDY 6

Imagine that you go to a friend's birthday dinner at a restaurant in a large city. You sit next to [name], a cousin of your friend whom you have never met before.

After you introduce yourself to each other, you ask [name] how [his/her] day was at work today. [name] then tells you the following story:

Note: See the scenarios presented in the Appendix of the manuscript.

Note: Display manipulation is the same as in Studies 1 and 3.

B: DETAILS ON DEPENDENT VARIABLES

Authenticity (Studies 1 – 6)

Please indicate the extent to which you agree or disagree with the following statements about [name].

- [name] is out of touch with his/her real self.
- [name] is true to himself/herself in most situations.
- [name] lives in accordance with his/her values and beliefs.
- [name] feels alienated from himself/herself.
- [name] would prefer to be himself/herself than to be popular.
- [name] doesn't know what he/she really feels inside.
- [name] is strongly influenced by the opinions of others.
- [name] would usually do what other people tell him/her to do.
- [name] always feels the need to do what others expect him/her to do.
- Other people influence [name] greatly.
- [name] knows himself/herself very well.
- [name] always stands by what [he/she] believes in.

Note that the measure of authenticity includes three sub-dimensions: Authentic Living (items 2, 3, 5, and 12), Accepting External Influence (items 7, 8, 9, and 10), and Self-Alienation (items 1, 4, 6, and 11).

Resilience (Studies 1 and 4)

Please indicate the extent to which you agree or disagree with the following statements about [name].

- [name] is able to adapt to change.
- [name] can deal with whatever comes.
- [name] tends to bounce back after illness or hardship.
- [name] can achieve goals despite obstacles.
- [name] can stay focused under pressure.
- [name] is not easily discouraged by failure.
- [name] is a strong person.
- [name] can handle unpleasant feelings.
- [name] can cope with stress.

Social Competence (Studies 1 and 4)

Please indicate the extent to which you agree or disagree with the following statements about [name].

- [name] behaves in a way that is considerate of others.

- [name] has the capacity for close relationships.
- [name] is liked by most people.
- [name] is sought out for advice and reassurance.
- [name] appears comfortable in social situations.
- [name] enjoys being with others.
- [name] is sociable.
- [name] is distrustful of people.
- [name] keeps people at a distance.
- [name] is personally charming.

C: ATTENTION CHECK MEASURES

PRE-SCREEN QUESTIONS

At the beginning of each study, participants were presented with an image of a sentence and asked to write down one word from that sentence into a text box. As per preregistration, those who did not answer this question correctly were not allowed to participate in the study. Below is the pre-screen question presented to participants.

Please type in the SECOND word ONLY from the following sentence.

The cat slept all day on the red couch.

ATTENTION CHECK QUESTIONS

At the end of each study, participants were asked to answer one or two multiple choice attention check questions. As per preregistration, those who did not answer the question(s) correctly were excluded from data analysis. Below is the list of attention check questions presented to participants across studies.

Study 1 & Post-test. “What was the name of the person you met in the scenario?”
[options: Randomly presented name; Mark; Ashley; Tim]

Studies 2 & 3. “Where did [name] say that [he/she] went yesterday?” [options: Yoga retreat; Housewarming party; Hiking; The theatre]

Studies 4 & S1. “What did we say that this person did in the scenario you read at the beginning of the study?” [options: Tripped and fell in the street; Fell asleep and snored in the theatre; Mistook someone for a beggar in a park; Knocked over and broke a glass in a restaurant; Gave a presentation with the zipper undone; Walked into a glass door at a party]

Study 5. “What was the situation that [name] experienced yesterday?” [options: Farting in a yoga class; Walking into a glass door at a party; Falling into horse poop during hiking; Waving at the wrong person at the theater]

Study 6. “What did we say that [name] did at work?” [options: Spilled a cup of hot tea; Tripped and [fell / knocked a colleague over] in the conference room; Walked into a glass door]

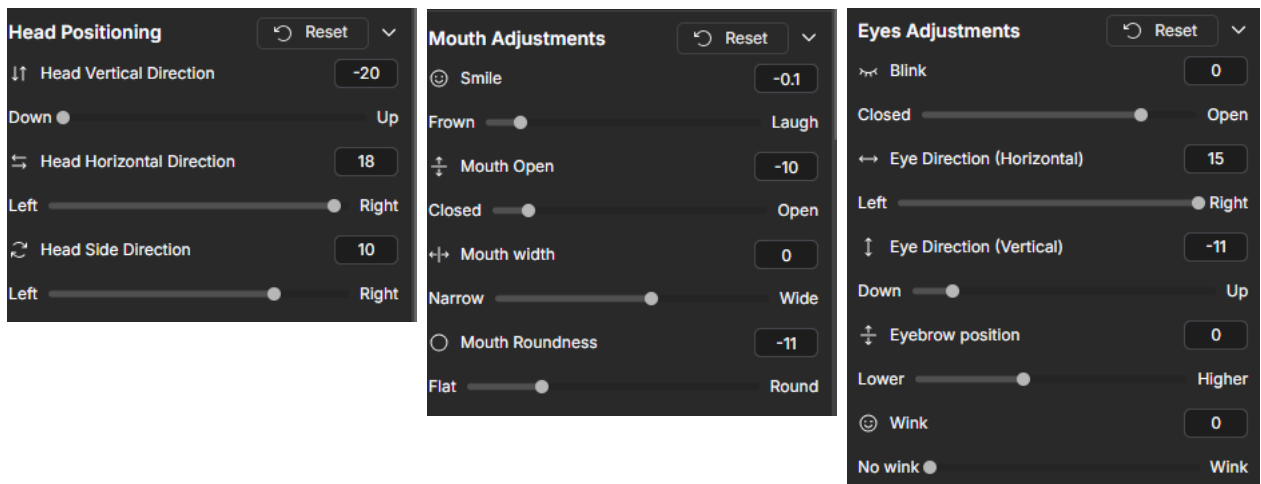
D: AI STIMULI STUDY 2

We used OpenArt (www.openart.ai) to generate photos of an individual displaying amusement or embarrassment. First, we used the image generator function and prompted the tool to generate photos of “an average female in her 30s with a neutral expression” and “an average male in his 30s with a neutral expression.” The AI created the following photos:



Next, we used the image editing function to change the facial expressions by adjusting the position of the actor’s head, mouth, and eyes, according to following configurations:

Embarrassment Display



These adjustments left us with the male and female photos with embarrassment display presented in the paper.

Amusement Display

Head Positioning

↺ Reset ↻

↕ Head Vertical Direction 12

Down ————— Up

↔ Head Horizontal Direction 0

Left ————— Right

↻ Head Side Direction 0

Left ————— Right

Mouth Adjustments

↺ Reset ↻

😊 Smile 1.3

Frown ————— Laugh

⬆ Mouth Open 29

Closed ————— Open

↔ Mouth width 7

Narrow ————— Wide

○ Mouth Roundness 3

Flat ————— Round

Eyes Adjustments

↺ Reset ↻

👁 Blink -5

Closed ————— Open

↔ Eye Direction (Horizontal) 0

Left ————— Right

↕ Eye Direction (Vertical) -7

Down ————— Up

⬆ Eyebrow position 4

Lower ————— Higher

😊 Wink 0

No wink ● Wink

These adjustments left us with the following male and female photos:



Next, we inputted these photos again into the image editor and conducted one additional of adjustment on the smile of the actor, as presented below. None of the other settings were changed.

Mouth Adjustments

↺ Reset ↻

😊 Smile 0.8

Frown ————— Laugh

These adjustments left us with the male and female photos with amusement display presented in the paper.

E: DEVIATIONS FROM PRE-REGISTRATIONS & ADDITIONAL ANALYSIS

STUDY 1

Post-Test. Two separate independent sample t-tests revealed that participants perceived their interaction partner to be less embarrassed ($M = 4.16$, $SD = 1.56$) but more amused ($M = 5.26$, $SD = 1.26$) when they displayed amusement relative to when they displayed embarrassment ($M_{Embarrassed} = 6.09$, $SD = 1.13$, $t(193) = 9.86$, $p < .001$, $d = 1.33$, 95% $CI_d = [1.02, 1.64]$; $M_{Amused} = 2.80$, $SD = 1.62$, $t(590) = 25.68$, $p < .001$, $d = 1.70$, 95% $CI_d = [1.37, 2.02]$).

Maximal Model on Warmth. In addition to the random effects model reported in the main manuscript, we fitted a series of maximal models that included fixed effects for Display (0 = Embarrassment Display; 1 = Amusement Display), Perceived Harm (continuous), and their interaction, and random intercepts for Scenario. In addition, the maximal model also included a random slope for Display, allowing the effect of Display to vary across scenarios. As noted in footnote 5, a model comparison test showed that this model fit the data better than a model that did not include a random slope for Display only for the warmth DV ($\chi^2(2) = 6.76$, $p = .03$).

The maximal model on warmth replicates our existing results. Specifically, we find a significant Display x Perceived Harm interaction ($B = -.22$, $SE = .09$, 95% $CI_B = [-.40, -.04]$, $t(9.73) = -2.50$, $p = .032$). There was a negative association between perceived harm and amusement: when an actor displayed amusement, they were rated less warm the more harm was caused to others ($B = -.24$, $SE = .10$, 95% $CI_B = [-.44, -.04]$, $t(10.21) = -2.52$, $p = .030$). In contrast, there was no significant relationship between perceived harm and character

judgments when the actor displayed embarrassment ($B_{\text{warmth}} = -.02$, $SE = .05$, $95\% CI_B = [-.12, -.08]$, $t(8.86) = -0.43$, $p = .67$).

A floodlight analysis using the Johnson–Neyman technique identified the range(s) of Perceived Harm for which the effect of Display was statistically significant. This analysis revealed that there was a significant positive effect of displaying amusement over displaying embarrassment on perceptions of warmth, when the perceived harm fell below 2.79. There were no values for warmth for which there was a significant positive effect of embarrassment over amusement. Figure S1 displays these results.

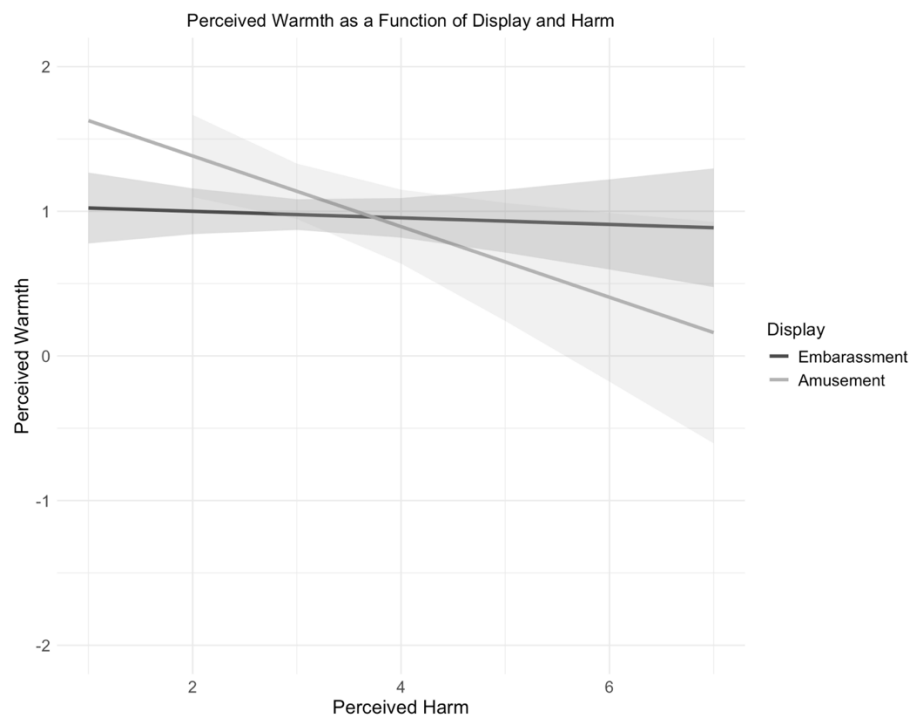


Figure S1. Perceived warmth as a function of Display (Amusement vs. Embarrassment) and Perceived Harm to third parties. Shaded regions represent 95% CIs.

Authenticity Sub-Scales. The scale of authenticity we used includes three sub-scales representing authentic living (i.e., how closely a person’s actions are perceived to align with their true self and core values), accepting external influence (i.e., the extent to which a person conforms to the expectations, norms, or pressures of others), and self-alienation (i.e., the degree to which someone is out of touch with their emotions or true self). We analyze each of

these sub-scales below. Note that in the present analysis scoring higher on accepting external influence and self-alienation represents an actor being seen as less authentic.

We first collapsed across the 12 scenarios presented to participants and ran a set of independent sample t-tests on each of the authenticity sub-scales. This revealed that the actor who displayed amusement was perceived to higher on authentic living ($M_{Amusement} = 0.88$, $SD = 0.71$ vs. $M_{Embarrassment} = 0.48$, $SD = 0.70$; $t(590) = 7.02$, $p < .001$, $d = .57$, 95 % $CI_d = [.40, .73]$), lower on accepting external influence ($M_{Amusement} = -0.55$, $SD = 0.86$ vs. $M_{Embarrassment} = 0.30$, $SD = 0.79$; $t(590) = -12.53$, $p < .001$, $d = 1.02$, 95 % $CI_d = [.86, 1.20]$), and lower on self-alienation ($M_{Amusement} = -1.04$, $SD = 0.80$ vs. $M_{Embarrassment} = 0.47$, $SD = 0.78$; $t(590) = -8.74$, $p < .001$, $d = .72$, 95 % $CI_d = [.55, .89]$) compared to the actor who displayed embarrassment.

Next, we ran a set of linear mixed-effects regression models on each of sub-scales, with Display condition as the fixed factor (Embarrassment Display = 0, Amusement Display = 1) and Scenario as a random factor. This confirmed that the effect of Display was robust across all sub-scales of authenticity when including random effects for Scenario ($B_{AuthenticLiving} = 0.40$, $SE = .06$, 95% $CI_B = [.28, .52]$, $t(588) = 6.92$, $p < .001$; $B_{ExternalInfluence} = -0.84$, $SE = .07$, 95% $CI_B = [-.98, -.70]$, $t(589) = -12.45$, $p < .001$; $B_{SelfAlienation} = -0.55$, $SE = .06$, 95% $CI_B = [-.67, -.43]$, $t(587) = -8.56$, $p < .001$). We then used the Stimulus package in R to run a heterogeneity analysis (Simonsohn, et al., 2025). This revealed that the level of variation in response to display across scenarios is larger than would be expected by chance for each sub-scale ($p_{AuthenticLiving} = .004$; $p_{ExternalInfluence} = .016$; $p_{Self-Alienation} < .001$).

To probe the source of this heterogeneity, we fit a linear mixed-effects model for each sub-scale that assessed whether each result was moderated by perceived harm to third parties. The model included fixed effects for Display (0 = Embarrassment Display; 1 = Amusement

Display), Perceived Harm (continuous), and their interaction, as well as random intercepts for Scenario.

Results show a significant Display x Perceived Harm interaction for the Authentic Living ($B = -.14$, $SE = .06$, 95% $CI_B = [-.26, -.02]$, $t(583) = 2.54$, $p = .011$) and Self-Alienation ($B = .23$, $SE = .06$, 95% $CI_B = [.11, .35]$, $t(581) = -3.62$, $p < .001$) sub-scales, but not for the External Influence sub-scale ($B = .04$, $SE = .07$, 95% $CI_B = [-.10, .18]$, $t(584) = -0.60$, $p = .055$). We explore the nature of the two significant interaction effects below.

For the Authentic Living sub-scale, neither the relationship between Amusement and perceived harm ($B = -0.11$, $SE = .06$, 95% $CI_B = [-.23, .01]$, $t(19) = -2.01$, $p = .06$) nor the relationship between Embarrassment and Perceived Harm was significant ($B = 0.03$, $SE = .05$, 95% $CI_B = [-.07, .13]$, $t(17) = 0.58$, $p = .57$), despite the interaction being significant. A floodlight analysis revealed the ranges of Perceived Harm for which the effect of display was statistically significant. Results revealed that displaying amusement (vs. embarrassment) significantly increased perceptions of Authentic Living when Perceived Harm was below 4.53. There was no range of Perceived Harm where embarrassment displays led to greater ratings of Authentic Living relative to amusement displays.

For the Self-Alienation sub-scale, those who displayed amusement were perceived to be more self-alienated the more their faux pas was perceived to harm others ($B = 0.18$, $SE = .07$, 95% $CI_B = [.04, .32]$, $t(17) = -2.75$, $p = .01$). In contrast the relationship between displaying embarrassment and perceived harm was not significant ($B = 0.04$, $SE = .07$, 95% $CI_B = [-.10, .18]$, $t(16) = 0.66$, $p = .52$). A floodlight analysis revealed the ranges of Perceived Harm for which the effect of display was statistically significant. This showed that displaying amusement (vs. embarrassment) significantly decreased perceptions of Self-Alienation when Perceived Harm was below 4.56. There was no range of Perceived Harm

where amusement displays led to increased ratings of Self-Alienation relative to embarrassment displays. Figure S2 displays these results.

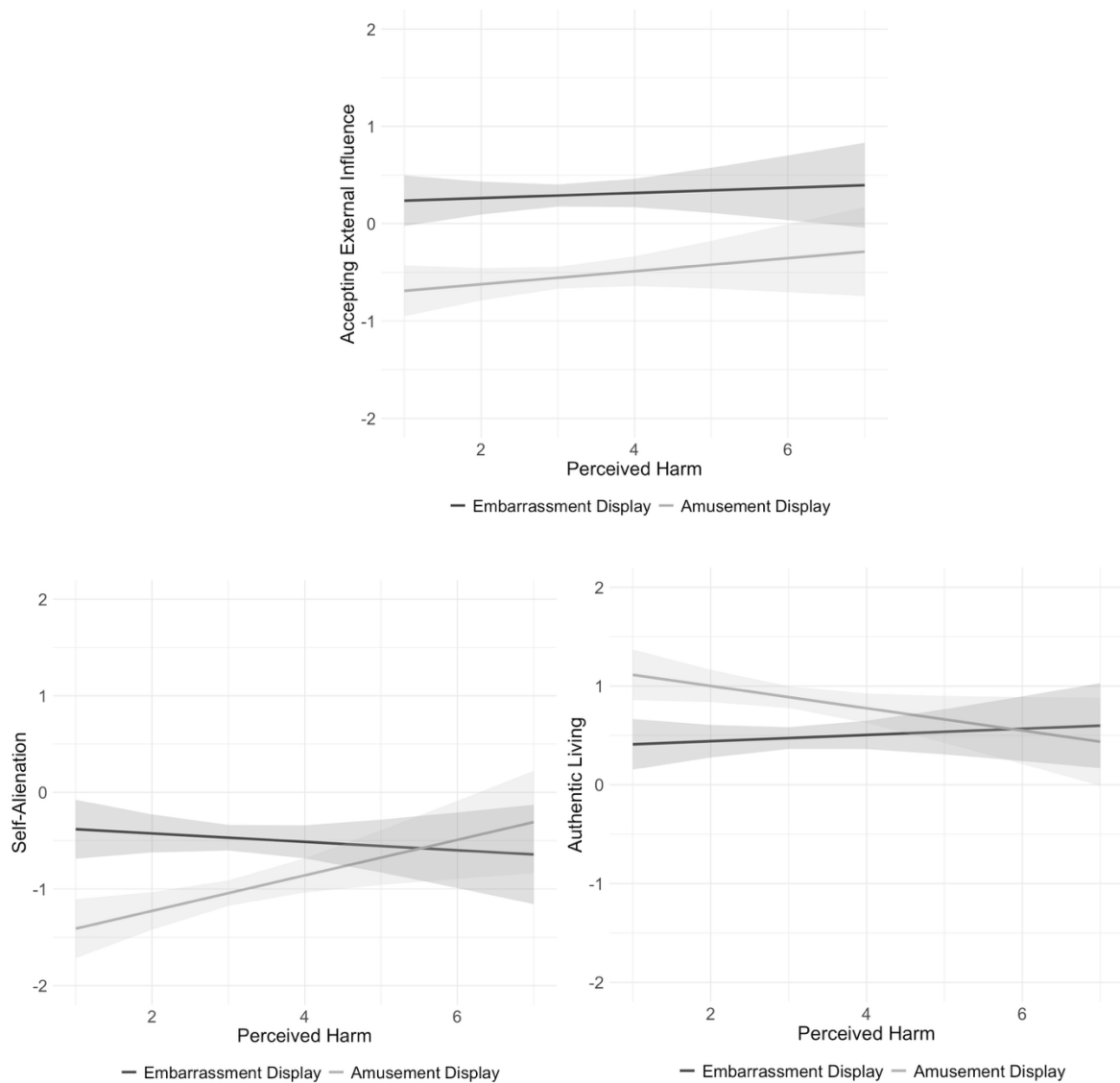


Figure S2. Ratings of each authenticity sub-scales as a function of Display (Amusement vs. Embarrassment) and Perceived Harm to third parties. Shaded regions represent 95% CIs. Note that higher levels of Self-Alienation and Accepting External Influence correspond to being less authentic.

Resilience. An independent sample t-test revealed that the partner who displayed amusement was perceived to be more resilient ($M = 1.02$, $SD = .69$) compared to the partner who displayed embarrassment ($M = .11$, $SD = .78$; $t(590) = 14.87$, $p < .001$, $d = 1.24$, 95% $CI_d = [1.06, 1.41]$) upon sharing a faux pas.

Next, we ran a linear mixed-effects regression model on resilience, with Display condition as the fixed factor (Embarrassment Display = 0, Amusement Display = 1) and Scenario as a random factor. This confirmed that the effect of Display was robust on resilience when including random effects for Scenario ($B = 0.90$, $SE = .06$, 95% $CI_B = [.78, 1.02]$, $t(588) = 14.93$, $p < .001$). We then used the Stimulus package in R to run a heterogeneity analysis (Simonsohn, et al., 2025), which revealed significant heterogeneity across scenario in response to Display ($p < .001$).

To probe this heterogeneity, we thus fit a linear mixed-effects model which included fixed effects for Display (0 = Embarrassment; 1 = Amusement), Perceived Harm (continuous), and their interaction as well as random intercepts for Scenario. This revealed as significant Display x Perceived Harm interaction ($B = -.13$, $SE = .06$, 95% $CI_B = [-.25, -.01]$, $t(583) = -2.25$, $p = .024$). When an actor displayed amusement, they were perceived as less resilient, the more harm was caused to others ($B = -.12$, $SE = .06$, 95% $CI_B = [-.24, -.002]$, $t(19) = -2.06$, $p = .05$). However, when an actor displayed embarrassment, the level of harm caused to others did not significantly relate to perceived resilience ($B = .02$, $SE = .06$, $t(18) = 0.29$, 95% $CI_B = [-.10, .13]$, $p = .77$). A Johnson-Neyman analysis revealed that those who displayed amusement were considered more resilient relative to those who displayed embarrassment for all values of perceived harm below 6.7, while there were no values of perceived harm for which those who displayed embarrassment were perceived significantly more resilient than those who displayed amusement. Figure S3 presents these results.

Social Competence. An independent sample t-test revealed that the partner who displayed amusement was perceived to be more socially competent ($M = .99$, $SD = .65$) compared to the partner who displayed embarrassment ($M = .43$, $SD = .67$; $t(590) = 10.27$, $p < .001$, $d = .85$, 95% $CI_d = [.68, 1.02]$,) upon sharing a faux pas.

Next, we ran a linear mixed-effects regression model on social competence, with Display condition as the fixed factor (Embarrassment Display = 0, Amusement Display = 1) and Scenario as a random factor. This confirmed that the effect of Display was robust on social competence when including random effects for Scenario ($B = .54$, $SE = .05$, $t(586) = 10.13$, 95% $CI_B = [.44, .64]$, $p = .77$). We then used the Stimulus package in R to run a heterogeneity analysis (Simonsohn, et al., 2025), which revealed significant heterogeneity across scenario in response to Display ($p < .001$).

To probe this heterogeneity, we thus fit a linear mixed-effects model which included fixed effects for Display (0 = Embarrassment Display; 1 = Amusement Display), Perceived Harm (continuous), and their interaction as well as random intercepts for Scenario. This revealed as significant Display x Perceived Harm interaction ($B = -.22$, $SE = .05$, 95% $CI_B = [-.32, -.12]$, $t(581) = -4.23$, $p < .001$). When an actor displayed amusement, they were perceived as less socially competent, the more harm was caused to others ($B = -.18$, $SE = .06$, 95% $CI_B = [-.30, -.06]$, $t(16) = -3.00$, $p = .008$). However, when an actor displayed embarrassment, the level of harm caused to others did not significantly relate to perceived social competence ($B = .04$, $SE = .06$, $t(15) = 0.78$, $p = .45$). A Johnson-Neyman analysis revealed that those who displayed amusement were considered more socially competent relative to those who displayed embarrassment for all values of perceived harm below 4.7, while there were no values of perceived harm for which those who displayed embarrassment were perceived significantly more socially competent than those who displayed amusement. Figure S3 presents these results.

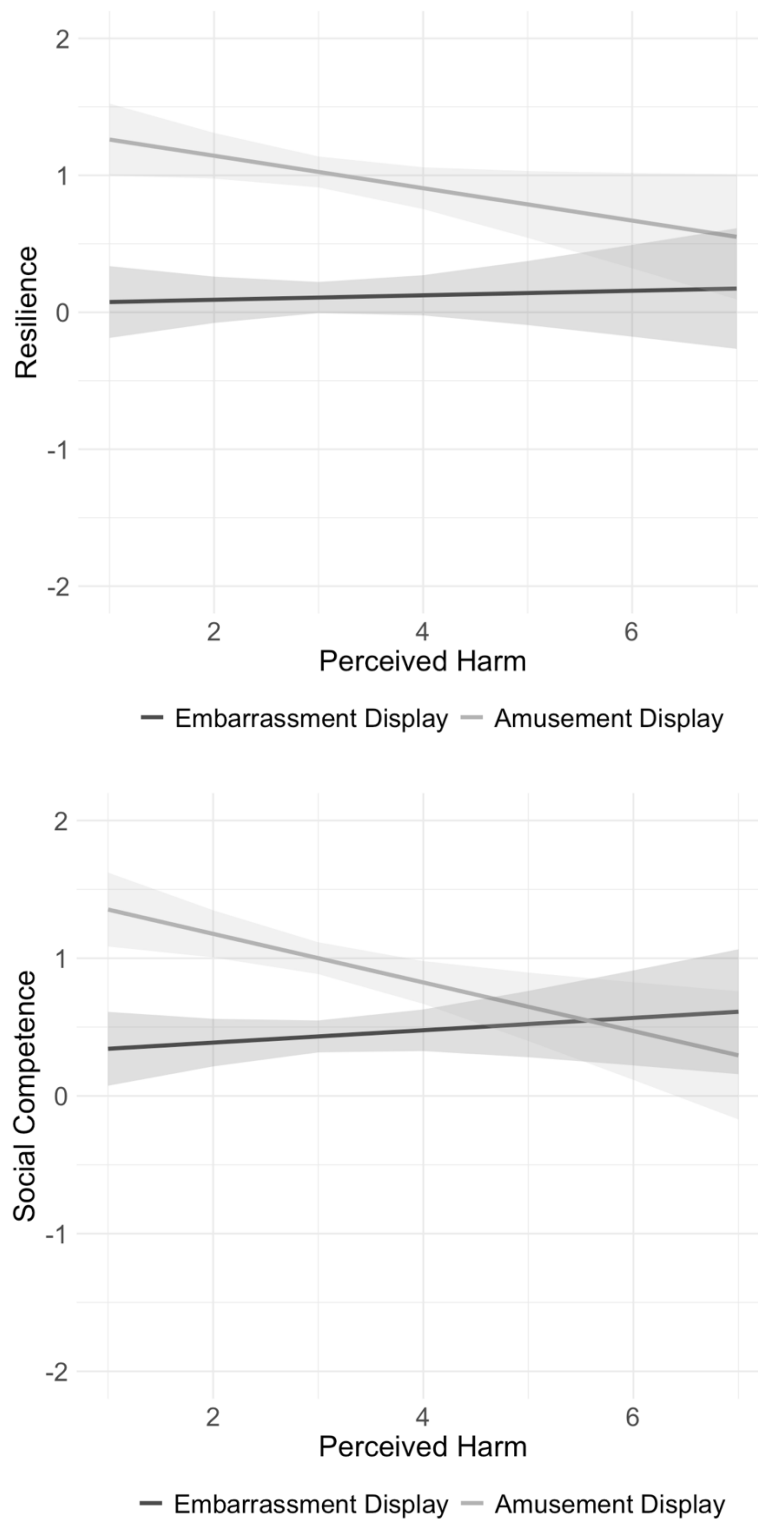


Figure S3. Perceived resilience and perceived social competence as a function of emotional display (Amusement vs. Embarrassment) and the extent to which the action harmed a third party. Dashed lines represent the 95% CIs. Note that since our stimuli only includes perceived harm ratings between 1.50 and 4.87, results outside this range should be treated with caution.

STUDY 2

Manipulation Check. In order to conduct a more detailed analysis of the manipulation check, we recoded the data as follows: First, we created two variables called Match and Mismatch which consist of participants' responses to the two manipulation check questions based on their condition. Specifically, for the Embarrassment Display condition, perceived embarrassment responses were coded as the Match ratings, and perceived amusement responses were coded as the Mismatch ratings. Conversely, for the Amusement Display condition, perceived amusement responses were coded as the Match ratings, and perceived embarrassment responses were coded as the Mismatch ratings.

Next, we ran a 2 (Display) x 2(Match) mixed ANOVA, where Match was a within-subjects factor and Display was a between-subjects factor. This revealed a significant main effect of Match ($F(1, 387) = 392.51, p < .001, \eta^2 = .50$), such that emotions that matched the manipulation ($M = 5.93, SD = .06$) were higher than those that mis-matched ($M = 3.32, SD = .09$). The main effect of Display ($F(1, 387) = .07, p = .79, \eta^2 < .001$) and the Display x Match interaction ($F(1, 387) = .89, p = .35, \eta^2 = .002$) were not significant. The non-significant interaction suggests that participants evaluated the difference between the focal and non-focal emotions to be similar in size across the two display conditions.

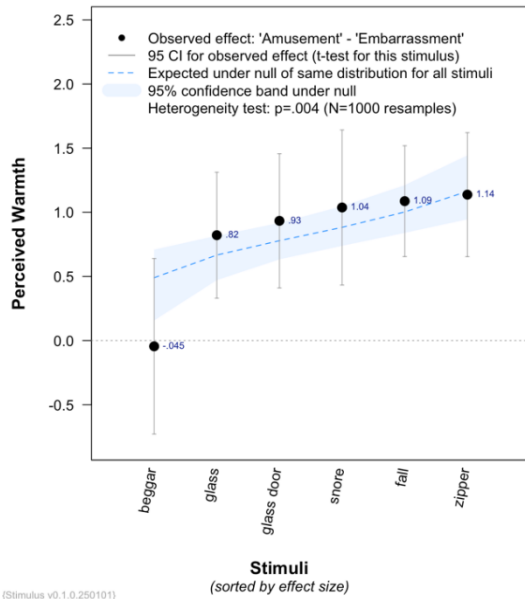
Authenticity Sub-Scales. We first collapsed across the 4 scenarios presented to participants and ran a set of independent sample t-tests on each of the authenticity sub-scales. This revealed that the actor who displayed amusement was perceived to higher on authentic living ($M_{Amusement} = 0.99, SD = 0.62, M_{Embarrassment} = 0.59, SD = 0.70; t(387) = 6.08, p < .001, d = .61, 95\% CI_d = [.40, .81]$), lower on accepting external influence ($M_{Amusement} = -0.58, SD = 0.84, M_{Embarrassment} = 0.12, SD = 0.92; t(387) = 7.84, p < .001, d = .79, 95\% CI_d = [.59, 1.00]$), and lower on self-alienation ($M_{Amusement} = -1.08, SD = 0.72, M_{Embarrassment} = -0.63, SD =$

0.81; $t(387) = 5.72, p < .001, d = .59, 95\% CI_d = [.39, .79]$) compared to the actor who displayed embarrassment.

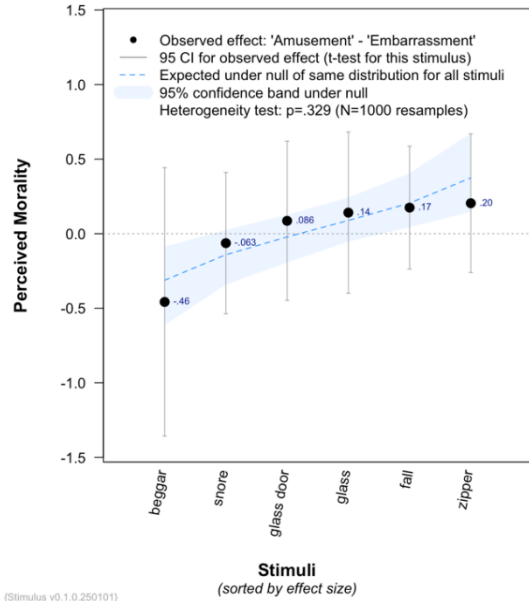
Next, we ran a set of linear mixed-effects regression models on each of sub-scale, with Display condition as the fixed factor (Embarrassment Display = 0, Amusement Display = 1) and Scenario as a random factor. This confirmed that the effect of Display was robust across all sub-scales of authenticity when including random effects for Scenario (Authentic Living: $B = -0.43, SE = .07, 95\% CI_B = [-.57, -.29], t(387) = -6.36, p < .001$; External Influence: $B = 0.73, SE = .09, 95\% CI_B = [.55, .91], t(387) = 8.18, p < .001$; Self-Alienation: $B = 0.45, SE = .08, 95\% CI_B = [.29, .61], t(387) = 5.84, p < .001$). We then ran a heterogeneity analysis, which revealed that the variation in the effect of Display was not significant across the scenarios for either of sub-scales (Authentic Living: $p = .44$; External Influence: $p = .48$; Self-Alienation: $p = .65$).

STUDY 4

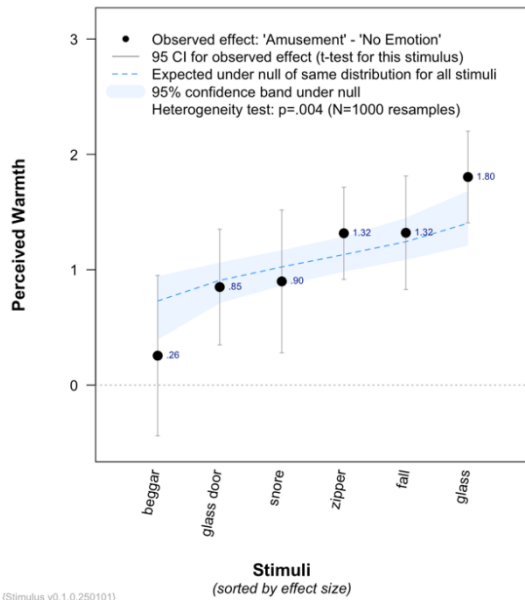
Stimulus Plots. Below we display the stimulus plots for perceived warmth and morality that compares the Amusement- and Embarrassment-Display conditions (top row), followed by comparisons of Amusement and Neutral-Display conditions (second row), and the Embarrassment and No Emotion conditions (last row).



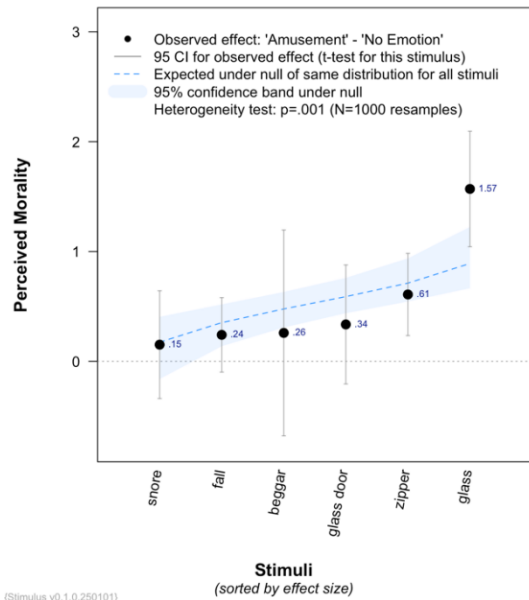
{Stimulus v0.1.0.250101}



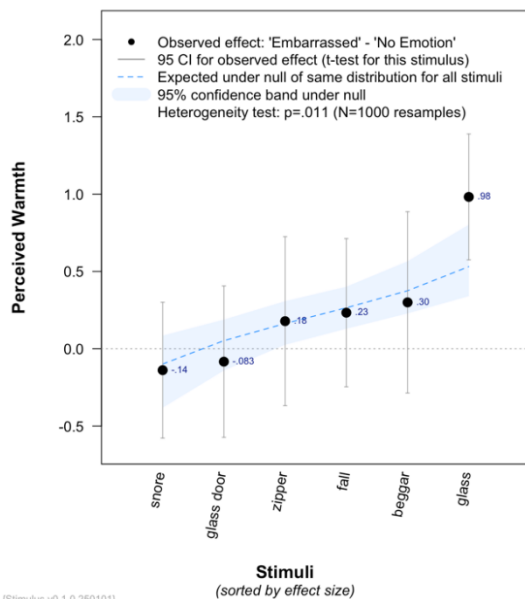
{Stimulus v0.1.0.250101}



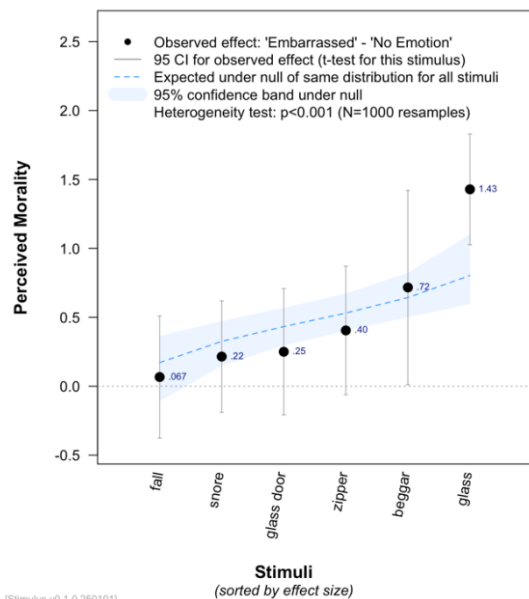
{Stimulus v0.1.0.250101}



{Stimulus v0.1.0.250101}



{Stimulus v0.1.0.250101}



{Stimulus v0.1.0.250101}

Authenticity sub-scales. We first collapsed across the 6 scenarios presented to participants and ran a set of one-way ANOVAs on each of the authenticity sub-scales using Display as the independent variable. These revealed a significant main effect of Display on authentic living ($F(2, 286) = 9.58, p < .001, \eta^2 = .06$), accepting external influence ($F(2, 286) = 33.69, p < .001, \eta^2 = .19$), and self-alienation ($F(2, 286) = 18.69, p < .001, \eta^2 = .12$).

A planned contrasts analysis showed that participants evaluated those who displayed amusement higher on authentic living ($M = 0.78, SD = 0.70$), lower on accepting external influence ($M = -0.52, SD = 0.80$), and lower on self-alienation ($M = -0.90, SD = 0.81$), compared to those who displayed embarrassment ($M_{AuthenticLiving} = 0.44, SD = 0.63, t(286) = 3.67, p < .001, d = .51, 95\% CI_d = [.22, .79]$; $M_{ExternalInfluence} = 0.29, SD = 0.66, t(286) = 7.29, p < .001, d = 1.10, 95\% CI_d = [.81, 1.41]$; $M_{SelfAlienation} = -0.39, SD = 0.65, t(286) = 4.92, p < .001, d = -.69, 95\% CI_d = [-.40, -.98]$). Moreover, those who displayed amusement were also evaluated higher than those who displayed a neutral reaction on authentic living ($M = 0.42, SD = 0.59, t(286) = 3.89, p < .001, d = .55, 95\% CI_d = [.26, .84]$) and self-alienation ($M = -0.32, SD = 0.68, t(286) = 5.60, p < .001, d = .77, 95\% CI_d = [.48, 1.07]$), but not on accepting external influence ($M = -0.48, SD = 0.86, t(286) = .35, p = .73, d = .05, 95\% CI_d = [-.23, .33]$). Finally, those who displayed embarrassment were not rated significantly differently than those who displayed a neutral reaction on authentic living ($t(286) = .24, p = .81, d = .06, 95\% CI_d = [.02, .12]$) and self-alienation ($t(286) = .71, p = .48, d = .10, 95\% CI_d = [-.18, .39]$), while they were rated higher on accepting external influence ($t(286) = 6.90, p < .001, d = 1.01, 95\% CI_d = [.70, 1.31]$).

Next, we ran a set of linear mixed-effects regression models, with Display condition as the fixed factor (dummy coded in turn so that each condition could serve as the reference category) and Scenario as a random factor on each dependent variable. We qualitatively replicated the above analysis on all authenticity sub-scales for all contrasts: Amusement

Display (referent) versus Embarrassment Display (Authentic Living: $B = -0.35$, $SE = .09$, 95% $CI_B = [-.52, -.17]$, $t(284) = -3.84$, $p < .001$; External Influence: $B = 0.80$, $SE = .11$, 95% $CI_B = [.58, 1.02]$, $t(283) = 7.29$, $p < .001$; Self-Alienation: $B = 0.51$, $SE = .10$, 95% $CI_B = [.31, .71]$, $t(284) = 4.93$, $p < .001$); Amusement Display (referent) versus Neutral Display (Authentic Living: $B = -0.37$, $SE = .09$, 95% $CI_B = [-.54, -.19]$, $t(284) = -4.01$, $p < .001$; External Influence: $B = 0.13$, $SE = .11$, 95% $CI_B = [-.09, .35]$, $t(284) = 0.12$, $p = .09$; Self-Alienation: $B = 0.58$, $SE = .10$, 95% $CI_B = [.38, .78]$, $t(286) = 5.60$, $p < .001$); and Embarrassment Display (referent) versus Neutral Display (Authentic Living $B = -0.02$, $SE = .09$, 95% $CI_B = [-.16, .20]$, $t(281) = -0.20$, $p = .84$; External Influence: $B = -0.78$, $SE = .10$, 95% $CI_B = [-.58, .98]$, $t(281) = 7.17$, $p < .001$; Self-Alienation: $B = 0.07$, $SE = .10$, 95% $CI_B = [-.12, .27]$, $t(282) = 0.69$, $p = .49$).

Finally, we ran a heterogeneity analysis across each possible pair of display conditions, which did not reveal any significant heterogeneity across scenario in response to Display for either of sub-scales (Authentic Living: $p = .63$; External Influence: $p = .91$; Self-Alienation: $p = .73$).

Resilience. A one-way ANOVA revealed a main effect of Display on perceived resilience ($F(2, 286) = 42.27$, $p < .001$, $\eta^2 = .23$). A planned contrasts analysis showed that participants evaluated those who displayed amusement as more resilient ($M = 1.00$, $SD = .68$) than those who displayed embarrassment ($M = 0.07$, $SD = .71$; $t(286) = 9.19$, $p < .001$, $d = 1.35$, 95% $CI_d = [1.04, 1.66]$) and those who displayed a neutral reaction ($M = 0.50$, $SD = .74$; $t(286) = 4.88$, $p < .001$, $d = .70$, 95% $CI_d = [.41, .99]$). Moreover, those who displayed embarrassment were rated as more resilient than those who displayed a neutral reaction ($t(286) = 4.26$, $p < .001$, $d = .61$, 95% $CI_d = [.32, .90]$).

Next, we ran a linear mixed-effects regression model, with Display condition as the fixed factor (dummy coded in turn so that each condition could serve as the reference

category) and Scenario as a random factor on each dependent variable. We qualitatively replicated the above analysis on resilience for all contrasts: Amusement Display (referent) versus Embarrassment Display ($B = -.93$, $SE = .10$, 95% $CI_B = [-1.13, -.73]$, $t(284) = -9.31$, $p < .001$); Amusement Display (referent) versus Neutral Display ($B = -.49$, $SE = .10$, 95% $CI_B = [-.69, -.29]$, $t(284) = -4.81$, $p < .001$); and Embarrassment Display (referent) versus Neutral Display ($B = .45$, $SE = .10$, 95% $CI_B = [.25, .65]$, $t(282) = 4.48$, $p < .001$).

Finally, we ran a heterogeneity analysis across each possible pair of display conditions. This did not reveal significant heterogeneity in responses to the display manipulation across Scenario ($ps > .43$).

Social Competence. A one-way ANOVA revealed a main effect of Display on perceived social competence ($F(2, 286) = 44.61$, $p < .001$, $\eta^2 = .24$). A planned contrasts analyses showed that participants evaluated those who displayed amusement as more socially competent ($M = 0.83$, $SD = .68$) than those who displayed embarrassment ($M = 0.17$, $SD = .64$; $t(286) = 7.00$, $p < .001$, $d = 1.00$, 95% $CI_d = [.70, 1.30]$) and those who displayed a neutral reaction ($M = -0.02$, $SD = .65$; $t(286) = 8.98$, $p < .001$, $d = 1.28$, 95% $CI_d = [.97, 1.59]$). Moreover, those who displayed embarrassment were rated as more socially competent than those who displayed a neutral reaction ($t(286) = 2.02$, $p = .044$, $d = .29$, 95% $CI_d = [.01, .58]$).

Next, we ran a linear mixed-effects regression model, with Display condition as the fixed factor (dummy coded in turn so that each condition could serve as the reference category) and Scenario as a random factor on each dependent variable. We qualitatively replicated the above analysis on resilience for all contrasts: Amusement Display (referent) versus Embarrassment Display ($B = -.69$, $SE = .09$, 95% $CI_B = [-.87, -.51]$, $t(283) = -7.47$, $p < .001$); Amusement Display (referent) versus Neutral Display ($B = -.87$, $SE = .09$, 95% $CI_B =$

[-1.05, -.69], $t(283) = -9.38, p < .001$; and Embarrassment Display (referent) versus Neutral Display ($B = -0.18, SE = .09, 95\% CI_B = [-.35, -.004], t(281) = -1.98, p = .048$).

We then ran a heterogeneity analysis across each possible pair of display conditions. This revealed significant heterogeneity in response to display condition in two of the contrasts (Amusement vs. Embarrassment: $p = .34$; Amusement vs. Neutral: $p = .009$; Embarrassment vs. Neutral: $p = .036$). To probe this heterogeneity, we fit a linear mixed-effects model that included fixed effects for Display (dummy coded in turn so that each condition could serve as the reference category), Perceived Harm (continuous), and their interaction, as well as random intercepts for Scenario. However, the Display x Perceived Harm interactions were not significant (Amusement vs. Embarrassment: $B = .13, SE = .11, t(280) = 1.15, p = .24$; Amusement vs. Neutral: $B = .12, SE = .11, t(280) = 1.09, p = .27$; Embarrassment vs Neutral: $B = -.01, SE = .11, t(279) = -0.07, p = .95$).

STUDY 5

Embarrassment Gap – Additional Analysis. A follow-up analysis on the embarrassment gap (i.e., the difference between perceived and appropriate embarrassment) examined which one of the components drive the differences. To assess this, we conducted two separate one-way ANOVAs on the perceived embarrassment and appropriate embarrassment. This revealed that the effect size of Display on perceived embarrassment ($\eta^2 = .12$) was 5.5 times larger than it was on appropriate embarrassment ($\eta^2 = .02$). This suggests that the effect of Display on the embarrassment gap was primarily driven by differences in perceived embarrassment rather than appropriate embarrassment.

Robustness Checks on Appeasement and Emotional Calibration. As in the previous studies, to test robustness, we then ran a set of non-preregistered analyses. First, we ran linear mixed-effects regression models for each dependent variable, including Display as a fixed factor (dummy coded in turn so that each condition could serve as the reference

category) and Scenario as a random factor. We qualitatively replicated the above analysis on appeasement and perceived emotional miscalibration for all contrasts: Amusement (referent) versus Embarrassment-Display conditions (Appeasement: $B = .16$, $SE = .13$, 95% $CI_B = [-.09, .41]$, $t(582) = 1.22$, $p = .22$; Emotional Miscalibration: $B = .32$, $SE = .15$, 95% $CI_B = [.03, .61]$, $t(582) = 3.12$, $p = .03$); Amusement (referent) versus Neutral-Display conditions (Appeasement: $B = -1.19$, $SE = .13$, 95% $CI_B = [-1.44, -.94]$, $t(581) = -9.12$, $p < .001$; Emotional Miscalibration: $B = .59$, $SE = .15$, 95% $CI_B = [.30, .88]$, $t(581) = 3.95$, $p < .001$); and Embarrassment (referent) vs. Neutral-Display conditions (Appeasement: $B = -1.35$, $SE = .13$, 95% $CI_B = [-1.60, -1.09]$, $t(582) = -10.36$, $p < .001$; Emotional Miscalibration: $B = .28$, $SE = .15$, 95% $CI_B = [-.01, .57]$, $t(582) = 1.84$, $p = .07$).

Next, we ran a test of heterogeneity in response to display across all possible pairs of conditions. As in Studies 2 and 3, this did not reveal significant heterogeneity in responses to the display manipulation across scenarios for either the appeasement ($ps > .096$) or emotional calibration dependent variables ($ps > .051$).

Finally, we also replicate our results for the analyses comparing perceived and appropriate embarrassment including Scenario as a random factor. To test this, we ran linear mixed-effects regressions that included Display as a fixed factor (dummy coded in turn so that each condition could serve as the reference category), Embarrassment Type (Perceived, Appropriate) and their interaction as fixed effects, and Scenario and Participant ID as random factors. The results show that participants felt the actor was perceived to be more embarrassed than they ought to be across each display condition (Embarrassment: $B = 1.23$, $SE = .14$, 95% $CI_B = [.96, 1.50]$, $t(583) = 8.73$, $p < .001$; Amusement: $B = 0.40$, $SE = .14$, 95% $CI_B = [-.23, .31]$, $t(584) = 2.80$, $p = .005$; Neutral: $B = 1.29$, $SE = .14$, 95% $CI_B = [1.02, 1.56]$, $t(584) = 9.13$, $p < .001$). Moreover, this difference was significantly smaller in the Amusement-Display condition compared to both the Embarrassment-Display ($B = 0.83$, $SE =$

.20, 95% $CI_B = [.44, 1.22]$, $t(584) = 4.16$, $p < .001$) and Neutral-Display ($B = 0.89$, $SE = .20$, 95% $CI_B = [.50, 1.28]$, $t(584) = 4.34$, $p < .001$) conditions. There was no significant difference in the size of this difference between the Embarrassment- and Neutral Display conditions ($B = 0.05$, $SE = .20$, 95% $CI_B = [-.34, .44]$, $t(584) = 0.28$, $p = .78$).

STUDY 6

Embarrassment Gap – Additional Analysis. First, a follow-up analysis to the reported analysis in the manuscript further examined what drove changes to the embarrassment gap (i.e., the difference between perceived and appropriate embarrassment). Our theorizing supposes that change to the embarrassment gap should be driven primarily by two forces. First, changes in emotional display should influence perceptions of experienced embarrassment more than perceptions of whether embarrassment is appropriate. Second, harming others (rather than the self) should increase the belief that embarrassment is appropriate, more than changing how embarrassed actors are perceived to be. To assess these, we ran two separate two-way ANOVAs on perceived embarrassment and appropriate embarrassment using Display and Harm as independent variables. Consistent with our first proposition, the effect sizes of Display on perceived embarrassment were 7.6 and 3.2 times those of appropriate embarrassment for the Self Harm ($d_{\text{Perceived}} = 1.68$ vs $d_{\text{Appropriate}} = .22$) and Other Harm ($d_{\text{Perceived}} = 1.88$ vs $d_{\text{Appropriate}} = .59$) conditions, respectively. Consistent with our second proposition, the effect sizes of Harm on appropriate embarrassment were 2.7 and 23.2 times larger than on perceived embarrassment, for the Embarrassment Display ($d_{\text{Perceived}} = 0.71$ vs $d_{\text{Appropriate}} = 0.26$) and Amusement Display ($d_{\text{Perceived}} = 1.63$ vs $d_{\text{Appropriate}} = 0.07$) conditions, respectively.

Embarrassment Gap – Preregistered Analysis. In addition to the reported analysis in the manuscript, we preregistered that we would conduct a two-way ANOVAs on the embarrassment gap (which we referred to as the embarrassment gap in the preregistration) as

the dependent variables using Display and Harm as independent variables. To do so, we subtracted the appropriate embarrassment from perceived embarrassment to create embarrassment gap as the DV. The two-way ANOVA on this measure revealed a main effect of Display ($F(1, 787) = 339.05, p < .001, \eta^2 = .30$), a main effect of Harm ($F(1, 787) = 154.23, p < .001, \eta^2 = .16$), and a Display x Harm interaction ($F(1, 787) = 33.75, p < .001, \eta^2 = .04$).

A planned contrast analysis revealed that, in the Self-Harm condition, embarrassment gap was smaller when individuals displayed amusement ($M = .92, SD = 1.81$) compared to when they displayed embarrassment ($M = 2.83, SD = 2.09, t(787) = 8.96, p < .001, d = .98, 95\% CI_d = [.77, 1.18]$). In contrast, in the Other-Harm condition, embarrassment gap was smaller when individuals displayed amusement ($M = -1.85, SD = 2.56$) compared to when they displayed embarrassment ($M = 1.83, SD = 2.03, t(787) = 17.04, p < .001, d = 1.59, 95\% CI_d = [1.36, 1.82]$).

Emotional Miscalibration – Preregistered Analysis. Moreover, we preregistered that we would run a t-test on emotional miscalibration (which we referred to as the absolute embarrassment gap in the preregistration) comparing the magnitude of the gap in the Embarrassment-Display and Amusement-Display conditions, focusing on the Self-Harm condition. This revealed that the absolute emotional miscalibration (i.e., interpersonal embarrassment) gap was smaller when individuals displayed amusement ($M = 1.54, SD = 1.32$) compared to when they displayed embarrassment ($M = 2.96, SD = 1.90, t(398) = 8.69, p < .001, d = .87, 95\% CI_d = [.66, 1.08]$) in the Self-Harm condition.

Authenticity sub-scales. We ran three separate two-way ANOVAs on each of sub-scale of authenticity using Display and Harm as independent variables. These revealed a significant Display x Harm interaction on authentic living ($F(1, 787) = 56.22, p < .001, \eta^2 =$

.07) and self-alienation ($F(1, 787) = 68.27, p < .001, \eta^2 = .08$), but not on acceptance of external influence ($F(1, 787) = .37, p = .55, \eta^2 < .001$).

A planned contrast analysis revealed that actors who harmed themselves (Self-Harm condition) were perceived as having a more authentic life ($t(787) = 9.72, p < .001, d = 1.01, 95\% CI_d = [.80, 1.22]$) and less self-alienated ($t(787) = 10.23, p < .001, d = 1.16, 95\% CI_d = [.95, 1.38]$) when they displayed amusement compared to embarrassment. However, when the faux pas inflicted harm on a third party (Other-Harm condition), those who displayed amusement were not perceived as having a more or less authentic life ($t(787) = .93, p = .35, d = .09, 95\% CI_d = [-.11, .29]$) or being more or less self-alienated ($t(787) = 1.50, p = .13, d = .15, 95\% CI_d = [-.04, .35]$) than those who displayed embarrassment. Moreover, those displaying amusement were perceived as being less open to external influence than those who display embarrassment across both Self-Harm ($t(787) = 11.35, p < .001, d = 1.14, 95\% CI_d = [.93, 1.35]$) and Other-Harm ($t(787) = 10.36, p < .001, d = 1.05, 95\% CI_d = [.84, 1.26]$) conditions. Table S1 displays these results.

Table S1
Study 6: Means and Standard Deviations

DV	Self-Harm Condition		Other-Harm Condition	
	Embarrassment Display ($N = 189$)	Amusement Display ($N = 211$)	Embarrassment Display ($N = 204$)	Amusement Display ($N = 187$)
Authentic Living	0.22 (0.72)	0.93 (0.69)	0.49 (0.75)	0.42 (0.78)
External Influence	0.50 (0.70)	-0.42 (0.89)	0.37 (0.81)	-0.48 (0.81)
Self-Alienation	-0.21 (0.78)	-1.02 (0.61)	-0.40 (0.78)	-0.28 (0.79)

Note: Standard deviations are presented in parentheses.

Moderated Mediation. We ran two separate sets of moderated mediation analyses, which include either the perceived emotional calibration or the perceived embarrassment as the only mediator in the model, for robustness checks.

Model 1 included each character judgment (warmth, morality, competence, and authenticity) as the dependent variable, Display (1 = Amusement, 0 = Embarrassment) as the independent variable, the perceived emotional miscalibration as the (continuous) mediator, and Harm (1 = Other, 0 = Self) as the moderator. Consistent with our moderated mediation hypothesis, we found significant index of moderated mediation on warmth, competence, and morality (Warmth = $-.14$, $SE = .04$, 95% $CI = [-.23, -.07]$; Competence = $-.09$, $SE = .04$, 95% $CI = [-.18, -.02]$; Morality = $-.12$, $SE = .05$, 95% $CI = [-.22, -.04]$), while there was no significant moderation for authenticity (Authenticity = $-.02$, $SE = .03$, 95% $CI = [-.07, .03]$). Table S2 presents all paths for each model, replicating results presented in the paper. Finally, the direct effect of Display was significant for warmth ($B = 0.51$, $SE = 0.05$, 95% $CI = [.41, .62]$), morality ($B = -0.27$, $SE = 0.06$, 95% $CI = [-.40, -.15]$), and authenticity ($B = 0.52$, $SE = 0.05$, 95% $CI = [.43, .61]$), but not for competence ($B = -0.05$, $SE = 0.06$, 95% $CI = [-.16, .07]$).

Model 2 included each character judgment (warmth, morality, competence, and authenticity) as the dependent variable, Display (1 = Amusement, 0 = Embarrassment) as the independent variable, the perceived embarrassment as the (continuous) mediator, and Harm (1 = Other, 0 = Self) as the moderator. Consistent with our moderated mediation hypothesis, we found significant index of moderated mediation on warmth, competence, and morality (Warmth = $-.03$, $SE = .02$, 95% $CI = [-.07, .001]$; Competence = $-.03$, $SE = .02$, 95% $CI = [-.07, .001]$; Morality = $-.05$, $SE = .03$, 95% $CI = [-.12, -.001]$), while there was no significant moderation for authenticity (Authenticity = $.02$, $SE = .01$, 95% $CI = [-.001, .05]$). Table S2 presents all paths for each model, replicating results presented in the paper. Finally, the direct

effect of Display was significant for warmth ($B = 0.71$, $SE = 0.07$, 95% $CI = [.57, .86]$) and authenticity ($B = 0.40$, $SE = 0.06$, 95% $CI = [.28, .52]$), but not for competence ($B = 0.14$, $SE = 0.08$, 95% $CI = [-.02, .29]$) and morality ($B = 0.06$, $SE = 0.08$, 95% $CI = [-.10, .23]$).

Table S2

Study 6: Moderated Mediation Results

	<i>Indirect Effect</i>	<i>SE</i>	<i>95% CI</i>	a pathway			b pathway		
				<i>b</i>	<i>SE</i>	<i>p</i>	<i>b</i>	<i>SE</i>	<i>p</i>
<u>Model 1 (Emotional Miscalibration is the mediator)</u>									
DV = Warmth									
Self Harm	0.11	0.03	[.06, .16]	-1.41	0.17	<.001	-0.07	0.02	<.001
Other Harm	-0.04	0.02	[-.08, -.01]	0.48	0.17	.005			
DV = Competence									
Self Harm	0.07	0.03	[.02, .13]	-1.41	0.17	<.001	-0.05	0.02	.005
Other Harm	-0.02	0.02	[-.06, -.002]	0.48	0.17	.005			
DV = Morality									
Self Harm	0.09	0.03	[.03, .15]	-1.41	0.17	<.001	-0.06	0.02	<.001
Other Harm	-0.03	0.02	[-.07, -.004]	0.48	0.17	.005			
DV = Authenticity									
Self Harm	0.01	0.02	[-.07, .03]	-1.41	0.17	<.001	-0.01	0.01	.44
Other Harm	-0.005	0.007	[-.03, .01]	0.48	0.17	.005			
<u>Model 2 (Perceived Embarrassment is the mediator)</u>									
DV = Warmth									
Self Harm	-0.15	0.05	[-.24, -.05]	-2.29	0.14	<.001	0.06	0.02	<.001
Other Harm	-0.17	0.06	[-.30, -.06]	-2.69	0.14	<.001			
DV = Competence									
Self Harm	-0.15	0.05	[-.26, -.04]	-2.29	0.14	<.001	0.06	0.02	.002
Other Harm	-0.17	0.07	[-.32, -.05]	-2.69	0.14	<.001			
DV = Morality									
Self Harm	-0.28	0.06	[-.39, -.16]	-2.29	0.14	<.001	0.12	0.02	<.001
Other Harm	-0.33	0.08	[-.49, -.18]	-2.69	0.14	<.001			
DV = Authenticity									
Self Harm	0.11	0.05	[.03, .21]	-2.29	0.14	<.001	-0.05	0.02	.002
Other Harm	0.13	0.05	[.03, .23])	-2.69	0.14	<.001			

F: GENDER INTERACTIONS

STUDY 4

Warmth. We ran a 2 Display (Embarrassment, Amusement, Neutral) x 2 Actor Gender (Male vs. Female) ANOVA on warmth. This revealed a significant Display x Actor Gender interaction ($F(2, 283) = 5.29, p = .006, \eta^2 = .04$; See Figure S4). A planned contrasts analysis showed that benefits of Amusement versus Embarrassment Displays holds for both genders, while this effect is larger for male actors ($t(286) = 7.30, p < .001$) compared to female actors ($t(283) = 3.61, p < .001$). The benefits of Amusement versus Neutral Displays also holds for both genders, but it was larger for males ($t(286) = 9.23, p < .001$) than females ($t(286) = 4.62, p < .001$). Finally, the Embarrassment versus Neutral Displays contrast was not significant for either gender (Female: $t(286) = 1.11, p = .27$; Male: $t(286) = 1.83, p = .07$).

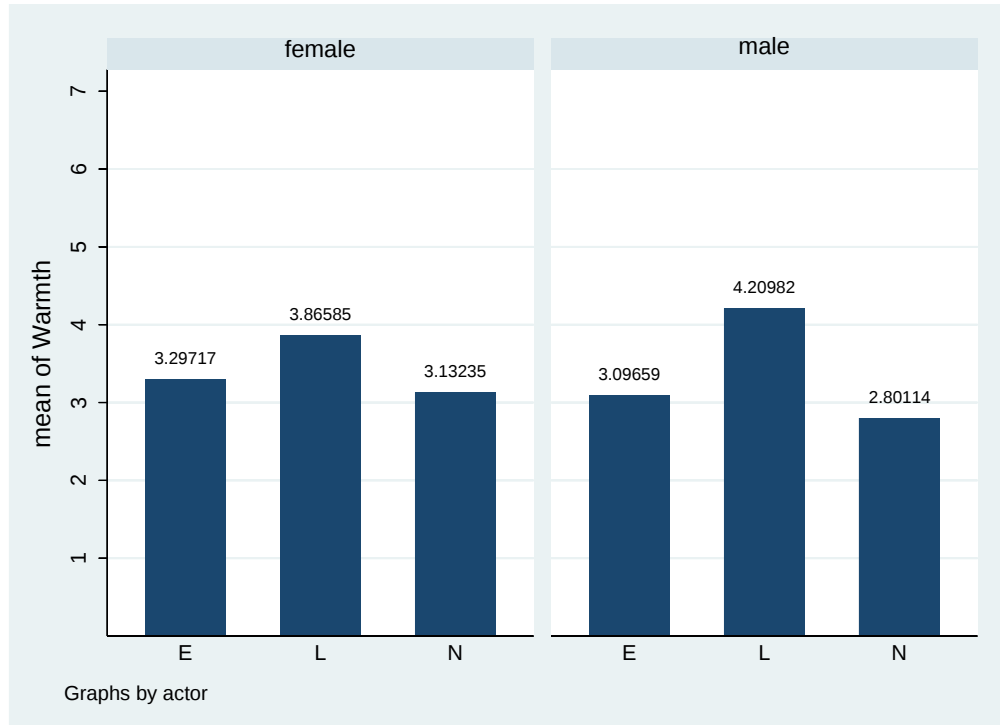


Figure S4. Means of perceived warmth across Display and Actor Gender conditions in Study 4. A = Amusement, E = Embarrassment, N = Neutral.

G: SUPPLEMENTARY STUDY S1

We aimed to recruit 300 participants via Amazon Mechanical Turk in exchange for payment and ended up with 299 participants. Per our preregistration plan, we excluded 12 participants who failed the attention check, leaving us with 287 participants (48.4% male, 49.5% female, 1.4% non-binary, .4% self-described, and .4% did not disclose; $M_{age} = 40$ years old, $SD = 13.9$).

This study employed the same 3 Display (Embarrassment, Amusement, Neutral) x 6 Scenario between-subjects design as in Study 4. Participants were presented with the same scenarios and the same manipulation describing the actor's emotional displays. They then completed two different dependent measures in a random order. First, they rated how generous, cooperative, morally upright, and trustworthy this person was as a four-item measure of perceived prosociality (adapted from Feinberg et al., 2012; $\alpha = .93$). Items were presented in a random order and rated on 7-point scales (1 = not at all, 7 = extremely). Second, they rated the extent to which they thought this agent acknowledged that "something socially awkward happened" and "a minor unfortunate event happened" on 7-point scales (1 = not at all, 7 = to a great extent). We collapsed across these two items to construct a measure of norm violation recognition ($r = .90$, $p < .001$). Last, participants reported how embarrassed they thought their interaction partner was on a 7-point scale (1 = "not at all," 7 = "extremely").

First, a one-way ANOVA revealed a significant main effect of Display on perceived prosociality, $F(2, 284) = 21.41$, $p < .001$, $\eta^2 = .13$. A planned contrast analysis showed that those who displayed amusement ($M = 4.99$, $SD = 1.03$) were perceived as equally prosocial as those who displayed embarrassment ($M = 4.77$, $SD = 1.07$; $t(284) = 1.35$, $p = .18$, $d = .21$, 95% $CI_d = [-.07, .50]$). In addition, both displaying amusement ($t(284) = 6.23$, $p < .001$, $d =$

.88, 95% $CI_d = [.58, 1.17]$) and embarrassment ($t(284) = 4.79, p < .001, d = .67, 95\% CI_d = [.38, .97]$) were perceived as more prosocial than neutral displays ($M = 3.99, SD = 1.24$).

Second, a one-way ANOVA revealed a significant main effect of Display on perceptions of norm violation recognition, $F(2, 284) = 154.28, p < .001, \eta^2 = .52$. A planned contrast analysis showed that displaying amusement ($M = 6.12, SD = .84$) signaled greater recognition that a norm was violated compared to displaying embarrassment ($M = 5.59, SD = 1.43; t(284) = 2.56, p = .011, d = .46, 95\% CI_d = [.17, .75]$) or not displaying a neutral reaction ($M = 2.75, SD = 1.85; t(284) = 16.37, p < .001, d = 2.35, 95\% CI_d = [1.98, 2.71]$). Embarrassment displays also signaled greater recognition that a norm was violated compared to neutral displays ($t(284) = 13.58, p < .001, d = 1.72, 95\% CI_d = [1.38, 2.05]$).

As a secondary analysis, a one-way ANOVA revealed a significant main effect of Display on the perceived embarrassment of the individual ($F(2, 284) = 73.36, p < .001, \eta^2 = .34$). This person was perceived as feeling less embarrassed when displaying amusement ($M = 4.20, SD = 1.60$) compared to displaying embarrassment ($M = 6.25, SD = .88; t(284) = 9.29, p < .001, d = 1.59, 95\% CI_d = [1.26, 1.91]$), but more embarrassed compared to a neutral display ($M = 3.72, SD = 1.87; t(284) = 2.17, p = .03, d = .27, 95\% CI_d = [-.01, .55]$). In addition, those who displayed embarrassment were also perceived as feeling greater embarrassment than those who displayed a neutral reaction ($t(284) = 11.45, p < .001, d = 1.72, 95\% CI_d = [1.39, 2.06]$).