

ARTICLE

Happy me: Asymmetric bidirectional links between the self and positivity underpin biased processing

Naomi A. Lee^{1,2}  | Douglas Martin¹  | Jie Sui¹
¹School of Psychology, King's College, University of Aberdeen, Aberdeen, UK

²Department of Psychology and Mental Health, University of Worcester, Henwick Grove, Worcester, UK

Correspondence

Naomi A. Lee, Department of Psychology, University of Worcester, Henwick Grove, Worcester WR2 6AJ, UK.

Email: naomi-anne.lee2@worc.ac.uk

Funding information

Leverhulme Trust, Grant/Award Number: RPG-2019-010

Abstract

Understanding how the self relates to emotional valence is fundamental to theories of self-concept and cognitive bias. While separate literatures have documented advantages for information associated with the self and information of positive valence, the mechanisms and directionality of their interaction remain poorly understood. Across three experiments combining evaluative priming and shape-label matching tasks, the current research examined whether self-positivity and self-negativity reflect bidirectional expectations between self and emotional information. Results revealed an asymmetric bidirectional relationship: positive emotion primes enhanced self-stranger bias magnitudes by inhibiting processing of stranger-related targets rather than facilitating self-associated targets (Experiment 1a), while self-primes facilitated positivity-bias magnitudes by enhancing processing of positive emotional targets (Experiment 2). Critically, no analogous effects emerged for negative valence (Experiment 1b). These findings suggest that the dominant mechanism linking self and positivity depends on the direction in which the association is activated: emotional context shapes self-prioritization through inhibition of distal-other representations, while self-context shapes emotional processing through direct facilitation of positive information. The theoretical implications for self-concept maintenance and predictive processing accounts of social cognition are discussed.

KEYWORDS

emotion, predictive processing, self, self-concept, self-positivity, social cognition

This is an open access article under the terms of the [Creative Commons Attribution](https://creativecommons.org/licenses/by/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

© 2026 The Author(s). *British Journal of Psychology* published by John Wiley & Sons Ltd on behalf of The British Psychological Society.

BACKGROUND

Humans exhibit a fundamental cognitive self-centredness, with self-related information accessing prioritized processing across cognitive domains (i.e., memory, attention, decision making; Alexopoulos et al., 2012; Cherry, 1953; Moray, 1959; Rogers et al., 1999; Sui et al., 2012). However, not all self-related information is equally prioritized. Importantly, self-prioritization is dynamically shaped by context. Self-biases are enhanced when the self is paired with preferred objects and weaken or disappear when paired with non-preferred ones (Golubickis et al., 2021), and cultural relevance shapes which stimuli receive preferential treatment (Dalmaso et al., 2026). These findings suggest that the self does not occupy a permanent tier in a static processing hierarchy. Rather, it operates on a hierarchy continuously calibrated by the evaluative and social properties of incoming information.

Among contextual factors, emotional valence holds a central theoretical position. Decades of research document a self-positivity bias. Positive attributes, outcomes, and associations are processed with special efficiency when linked to the self (Chen et al., 2014; Fields & Kuperberg, 2015; Watson et al., 2007; Xia et al., 2021; Zhang et al., 2023). This pattern aligns with self-enhancement theory (Sedikides & Gregg, 2008), which holds that individuals are motivated to maintain a positive self-concept, leading to preferential processing of positively valenced self-related information. The adaptive value of this bias is well documented. Self-positivity is reliably associated with psychological wellbeing, effective self-regulation, and prosocial functioning (Alicke & Sedikides, 2009; Hung et al., 2025; Lee et al., 2023; Markus & Kitayama, 1991; Sui et al., 2021).

Recent research has challenged accounts of exclusive self-positivity. When positive and negative emotional targets are each assessed against a neutral baseline, rather than in direct opposition, self-related primes facilitate responses to both, revealing implicit self-evaluation bivalence (Zayas et al., 2022). This finding has been replicated when controlling for lexicality and stimulus extremity (Shi et al., 2025) and has been documented across cultures (Ni et al., 2025). The robustness of bivalence suggests that self-representations carry negative as well as positive evaluative connotations. The self-discrepancy theory (Higgins, 1987), where mismatches between actual, ideal, and ought selves generate negative affect alongside positive, offers one account of why this might be the case. Predictive processing models of self offer another (Albarracín et al., 2024; Friston & Stephan, 2007). The self-concept probabilistic expectations are derived from both self-affirming and self-threatening experiences, and both types of expectation can shape downstream processing. These theoretical perspectives share an important gap. They describe the existence of self-valence associations without specifying the mechanisms through which those associations bias processing, or whether those associations operate in one direction only or in both. The present study was designed to address both questions directly.

Disentangling facilitation and inhibition in self-valence processing

Even granting that self-related information biases emotional processing, the mechanism remains unclear. Two processes could produce faster responses to self-congruent valenced information: facilitation of congruent information, inhibition of incongruent information, or a combination of both. Facilitation implies that priming with self-relevant (or positive) information directly activates associated representations, lowering the threshold for processing congruent information. Inhibition implies that the prime suppresses competing representations, showing responses to incongruent information without necessarily enhancing responses to congruent information. Standard relative contrasts between positive and negative stimuli cannot resolve this ambiguity. Any observed self-positive advantage could arise from facilitation, inhibition or some combination (Wentura & Rothermund, 2014). A neutral baseline condition, allowing the direction of each effect to be estimated independently, is required.

Evaluative priming paradigms (Fazio et al., 1986) with neutral baselines offer an experimental method to independently estimate the direction of facilitatory and inhibitory effects. When prime and target share evaluative properties, responses are facilitated (faster reaction time); when they are

mismatched, responses may be slowed. By comparing responses following each prime type against a neutral baseline, it becomes possible to localize observed biases to facilitation of congruent targets, inhibition of incongruent targets, or both (Klinger & Greenwald, 1995; Wentura & Rothermund, 2014). Zayas et al. (2022) used this logic to reveal self-bivalence in evaluative priming. The present study applies it within a perceptual matching context, where the dissociation between facilitation and inhibition has not previously been examined.

A second foundational question is whether self-valence associations operate bidirectionally. The dominant research tradition asks whether self-relevant information biases subsequent emotional processing. Whether the reverse also holds, whether emotional information generates expectations of self-relevance that bias subsequent identity-related processing, has rarely been asked. If positive valence activates expectations of self-relevance, emotional context should modulate self-prioritization in ways that mirror the effect of self-context on emotional processing. However, this is not guaranteed. The relationship between self and positivity may be structurally asymmetric, with distinct mechanisms governing each direction. Whether bidirectionality exists, and whether any such bidirectionality is symmetric or asymmetric in mechanisms, are open questions with direct implications for theories of self-concept structure.

Self-valence influence on perceptual bias

Evaluative priming paradigms provide insight into the mechanisms linking self and valence, while shape-label matching tasks (Sui et al., 2012) offer a means of examining how these relationships influence biased perceptual processing. In the task, participants learn arbitrary personal shape-label pairings (e.g., You are a triangle, your friend is a circle, a stranger is a square). Following instruction, shape-label pairs are presented, and participants judge as quickly and accurately as possible whether each pair matches one of the learned associations (e.g., You-Triangle = match, Friend-Square = mismatch). Robust self-biases are observed in this task; participants are faster and more accurate in responding to matched self-related pairs in comparison to other-related pairs (Sui et al., 2012). In some literature, a friend-bias (advantageous responses to the friend in comparison to strangers) is also observed (Sui et al., 2012; Sui & Humphreys, 2017; Sun et al., 2016). When the task is adapted so that pairs are learned with emotional expressions rather than individuals, participants show a positivity bias—faster and more accurate responses to positive over negative and neutral emotions (Stolte et al., 2017). These emotion-biases appear to operate through similar mechanisms as self-biases, with both reflecting enhanced perceptual processing of prioritized stimulus categories (Stolte et al., 2017; Yankouskaya & Sui, 2021).

Self-biases in the matching task are reliably enhanced when the self is associated with positive stimulus (a happy face) and the stranger with negative stimulus (a sad face) (Constable et al., 2021). This effect has been replicated with more abstract valenced constructs such as light luminance and shape symmetry (Constable et al., 2021; Vicovaro et al., 2022). However, existing designs typically confound valence with identity, pairing the self with positive stimuli and others with negative ones, making it difficult to determine whether it is valence that modulates self-related processing, self-relevance that modulates emotional processing, or both. Moreover, these designs cannot reveal the mechanism whether observed enhancements arise from facilitation of self-congruent information, inhibition of self-incongruent information, or some mixture. By decoupling valence from identity through primes that are orthogonal to trial targets, the present design directly addresses these limitations.

The current research

The current research investigated the mechanisms and directionality of self-valence associations by combining evaluative priming and shape-label matching in a design that systematically varies the direction of priming. Experiment 1a examined whether positive emotion primes would modulate self-bias magnitudes, and whether any modulation would reflect facilitation of self-target processing, inhibition

of stranger target processing, or both. A larger self-stranger bias driven by faster responses to the self (regardless of stranger responses) would indicate self-positive facilitation, consistent with intrinsic self-positivity. A bias driven by slower responses to strangers (regardless of self-responses) would indicate stranger-positive inhibition, consistent with relative self-positivity, in which positivity is not intrinsic to the self but is withheld from distant others.

Experiment 1b applied the same logic to negative emotion primes. The self-positivity account predicts that negative primes should reduce self-bias magnitudes, driven by slower responses to self-targets and faster responses to other targets (Constable et al., 2021; Lee et al., 2023; Vicovaro et al., 2022). The bivalence account (Zayas et al., 2022) predicts that negative primes should facilitate responses to both self and others, resulting in no change in self-bias magnitude.

Experiment 2 reversed the priming relationship. Personal identity primes (self versus stranger) preceded emotional targets (happy, neutral, and sad), allowing examination of whether self-relevance modulates positivity bias magnitudes. If self-positive associations operate bidirectionally, self-primes should enhance the positivity bias, driven by facilitation of happy targets. If the relationship is asymmetric, this pattern should differ from the mechanism revealed in Experiment 1a.

EXPERIMENT 1: DO EMOTION PRIMES INFLUENCE PERSONAL BIASES?

Experiment 1a: Happy and neutral priming preceding personal shape-label matching

Introduction

Experiment 1a examined whether emotional context modulates self-bias magnitudes by presenting emotional expression primes (happy vs. neutral) before shape-label matching trials. Based on the intrinsic self-positivity account, we would expect the size of self-stranger biases to be enhanced following a positive prime, driven by faster responses to the self-target (self-positivity facilitation). However, if the self is positive only relative to others, enhancement would instead be driven by slower responses to stranger targets (stranger-positive inhibition). These two predictions are distinguished only by identifying the locus of the effect, a distinction that requires a neutral prime baseline.

Method

Participants

The target sample size was 42, in line with Zayas et al. (2022) and within the range (30–56) previously used in similar self-affective shape-label matching tasks (Constable et al., 2021; Vicovaro et al., 2022). In total, 50 participants were recruited from SONA and Prolific (www.prolific.co), and 7 were excluded due to failing standard shape-label matching task chance-level exclusion criteria (average accuracy $\leq 55\%$ and/or average reaction time < 200 ms). The average age of the 43 participants was 27.79 years ($SD = 10.78$, range = 18–64), consisting of 21 females and 22 males; 1 participant was ambidextrous, 3 were left-handed, and 39 were right-handed.

All participants had normal or corrected-to-normal vision and were fluent in English. Informed consent was acquired from all participants before the experiment following procedures approved by a local ethics committee.

Sensitivity analysis was conducted to assess the robustness of the results with the sample size and composition. Resampling-based sensitivity analyses appropriate for hierarchical data (leave-one-subject-out and subsampling) indicated that the observed effects are not dependent on individual participants or specific sample compositions, and the interaction effect was consistently recovered even

when the sample size was reduced to ~20 participants (see [Supporting Information](#) for sensitivity analysis for all experiments).

Procedure

Participants were told to associate three white shapes (triangle, diamond, and pentagon) with three people (self, a previously named best friend, and a stranger). The pairings were presented in a random order. Participants had to judge whether shape-label pairs matched one of the previously learnt associations ([Figure 1](#)). Each trial began with a fixation cross displayed at the centre of the screen for 500 ms. A happy or neutral emotional schematic face prime was presented above the fixation cross for 200 ms. Participants were instructed that the prime did not indicate which trial would follow. The emotional faces consisted of three black lines depicting expressions. Following prime presentation, it was removed from the screen, and just the fixation cross remained for 50 ms. Subsequently, a shape-label pair was presented for 150 ms; the shape was presented above the fixation cross and the label below. The pairing could match a learnt association or could be a recombination (shape and label from separate pairings). A blank screen was presented until participants judged whether the shape-label pair matched a learnt association or until 1500 ms elapsed. Participants were instructed to respond as quickly and accurately as possible. The allocation of response keys ('b', 'v') was counterbalanced across participants. Following the response, feedback (correct, incorrect, too slow) was provided for 500 ms. Average accuracy and RT performance were reported at the end of each block.

There were six demonstration trials, three match and three mismatch. In these trials, no prime was presented. The shape-label pair remained on screen with instruction until a response was made. The instructions indicated the condition and response key needed (i.e., 'MATCH! So, you'd press "v" with the index finger of your left hand'. Or 'NOT A MATCH! So, you'd press "b" with the index finger of your right hand. Remember this shape matched with "v"'). Six self-paced practise trials were then completed, which were identical to experimental trials except the stimuli pairing remained onscreen until a response was made. No instructions were provided. Finally, six practice trials identical to the experimental trials were then completed.

For the experiment, each participant performed five blocks of 72 trials totalling 360 trials, resulting in 30 trials in each condition (e.g., happy primed self-match). At the end of the experiment, participants completed a series of questionnaires, but these data are not presented here. The experiment was conducted online through Inquisit 6 [Computer software] (2021). Retrieved from <https://www.millisecond.com>.

Data analysis

All analyses were performed in R (R Core Team, 2023). In line with standard practice, responses which are unlikely to reflect conscious decision-making processes (RT < 200 ms) were excluded from all analyses (Sui et al., 2012). RT analysis was conducted on correct responses. Timeout trials (no response in 1500 ms) were counted as incorrect trials.

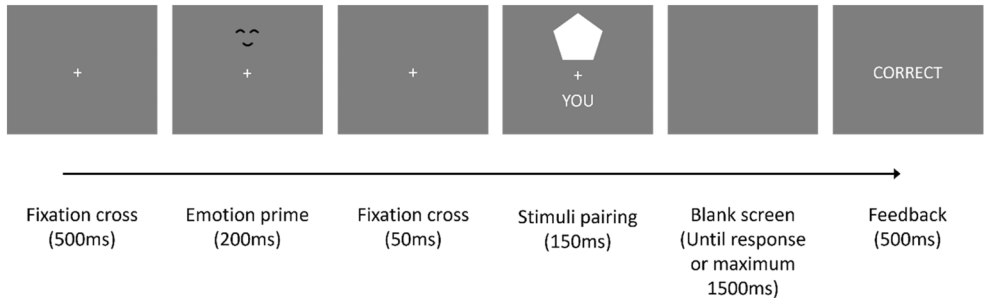


FIGURE 1 Example trial from Experiment 1. In Experiment 1a, the emotion prime could be a happy (shown) or neutral expression. In Experiment 1b, the emotion prime was a sad or neutral expression.

Generalized linear mixed effects models (GLMM) were conducted using lme4 (Bates et al., 2025) on RT data in match trials in line with standard shape-label matching paradigms (Sui et al., 2012). The fixed effect structure (i.e., prime, target, and prime*target interaction) was selected based on experimental design and hypotheses; therefore, no model comparison against simpler fixed effect model structures was conducted. The inverse Gaussian family with the identity link function was used to account for the long tail in the RT distribution (Lo & Andrews, 2015). The bound optimization by quadratic approximation (bodyqa) with a set maximum of 200,000 iterations was used. Maximal random effect structures were tested (i.e., each participant's intercepts and slopes varied as a function of targets, primes, and target*prime interactions) (Barr et al., 2013). If convergence or singularity issues were encountered, the random effect structure was simplified. The following equation was used:

$$\text{glmer}(\text{RT} \sim \text{personal target} * \text{emotional prime} + (1 + \text{personal target} + \text{emotional prime} | \text{participant}))$$

This includes fixed effects of personal target, emotional prime, and personal target*emotion prime interaction. Within random effects, each participant's intercept and slopes can be influenced by the target and prime.

The model was run through the joint_test function from emmeans (Lenth et al., 2025) to produce a Type-III-ANOVA-like omnibus test to assess the overall significance of main and interaction effects. Any significant results were then explored using pairwise comparisons and custom contrasts. Custom contrasts with Sidak adjustment were used to test whether self and friend bias sizes vary between emotional primes. Pairwise comparisons within emmeans were then used with Tukey adjustment to assess simple effects of prime within each target (i.e., positivity prime biases for each personal target).

Transparency and openness

All data, analysis code, and research materials are available at the Open Science Framework and can be accessed at (https://osf.io/enfma/overview?view_only=a357400ecae541949e5aba49d4a74842). Data were analyzed using R, version 4.5.0 (R Core Team, 2023) and the package lme4, version 1.1–37, and emmeans, version 2.0.0. This study was not preregistered.

Results

There was a significant main effect of personal target $F(2, \text{inf}) = 40.28, p < .0001$. Replicating previous findings of self-biases, responses were faster to the self (EMM = 716, SE = 9.46 [95% CI: 698–735]) in comparison to the friend (EMM = 797, SE = 16.10, 95% CI = [766–829]), $\beta = -81.3$ [95% CI: -112.9 to -49.6], $p < .0001$, and stranger (EMM = 811, SE = 14.30 [95% CI: 783–839]), $\beta = -94.4$ [95% CI: -122.4 to -66.5], $p < .0001$. However, no friend-bias was observed, with no difference between friend and stranger response times, $\beta = -13.1$ [95% CI: -49.9 to 23.7], $p = .68$. There was no significant main effect of prime $F(1, \text{inf}) = 0.64, p = .42$.

A significant interaction between personal targets and prime was observed $F(2, \text{inf}) = 5.02, p = .007$ (see Figure 2). As predicted, the self-stranger bias was larger following the happy prime relative to the neutral prime $\beta = -25.64$ [95% CI: -45.1 to -6.18], $p = .005$. However, relative to neutral primes, there was no evidence that happy primes modulated the size of the self-friend bias, $\beta = -6.05$ [95% CI: -25.0 to 12.89], $p = .83$ or friend-stranger bias $\beta = -19.60$ [95% CI: -41.4 to 2.17], $p = .09$.

In contrast to the prediction based on implicit self-positivity, modulation of the self-stranger bias was not a consequence of facilitation of responses to self-targets following happy primes, which did not differ relative to neutral primes, $\beta = -6.84$ [95% CI: -18.03 to 4.36], $p = .23$. Instead, in support of relative self-positivity (where positivity is exclusive to the self, relative to others), it seems the larger self-stranger bias that was observed following happy primes was driven by the significant inhibitory positivity affect at the stranger target, with responses significantly slower following the happy prime in comparison to the neutral prime, $\beta = 18.81$ [95% CI: 4.28–33.33], $p = .01$. No significant differences

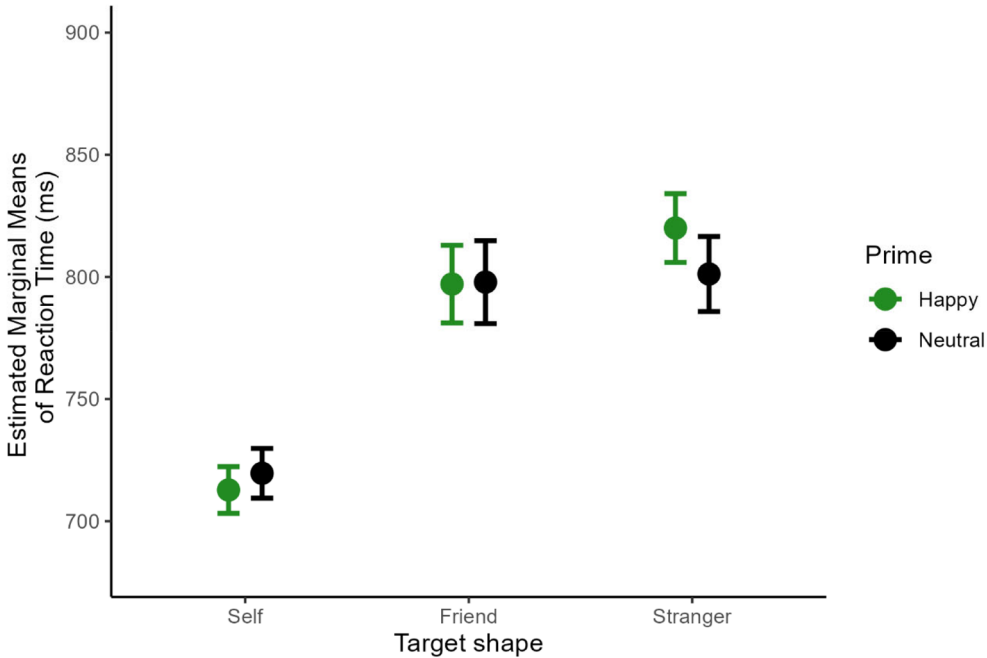


FIGURE 2 EMMs of RT as a function of Prime (Happy vs. Neutral) and Target (Self, Friend, or Stranger) in Experiment 1a. Error bars show standard error.

in response time following the different primes were observed for friend targets, $\beta = -0.79$ [95% CI: -14.46 to 12.88], $p = .91$.

Discussion

Experiment 1a revealed that positive emotional primes enhance the self-stranger bias, replicating and extending previous findings (Constable et al., 2021; Vicovaro et al., 2022). The enhancement was driven entirely by inhibitory stranger-associated processing, with no evidence of facilitation of self-associated processing. This finding suggests that positive emotion may specifically interfere with the processing of socially distant others, therefore indirectly increasing the relative prioritization for self-related information. This may explain why we found no evidence of differences in self-friend bias magnitudes between positive and neutral priming. The friend is sometimes considered part of the extended self, so this may have resulted in the positive prime leading to predictions of either the self or the friend (i.e., the extended self). Together these findings suggest that positive emotion context shapes self-prioritization not by making the self more accessible, but by making positivity incompatible with representing a stranger.

Experiment 1b: Sad and neutral priming preceding personal shape-label matching

Introduction

Experiment 1b tested whether negative emotional context modulates self-bias sizes. Two competing predictions were tested. The self-positivity account holds that positive information is selectively associated with the self and negative information with others; accordingly, sad primes should reduce

self-bias magnitudes by slowing responses to self-targets and speeding responses to other targets. The self-bivalence account (Zayas et al., 2022) holds that both positive and negative information are associated with the self. Therefore, negative primes should facilitate responses to the self and other associated targets. Consequently, sad primes should facilitate responses to both self and other targets, leaving self-bias magnitudes unchanged.

Method

Participants

Target participant sample and participant exclusion criteria were identical to Experiment 1a. In total, 53 participants were recruited from SONA and Prolific (www.prolific.co), 11 were excluded due to performance at chance level. The average age of the 42 participants was 25.00 years ($SD=8.48$, range=18–55), consisting of 24 females and 18 males; four participants were left-handed, and 38 were right-handed.

Procedure

The procedure was identical to Experiment 1a except that the primes were sad and neutral emotion faces.

Data analysis

The analysis was identical to Experiment 1a, the only difference being in the prime valences—sad and neutral primes were used here. The equation of the GLMM used was:

$$\text{glmer}(\text{RT} \sim \text{personal target} * \text{emotional prime} + (1 + \text{personal target} | \text{participant}))$$

This includes fixed effects of personal target, emotional prime, and personal target*emotion prime interaction. Within random effects, each participant's intercept and slopes can be influenced by the target.

Results

Mirroring Expt. 1a, there was a significant main effect of personal target $F(2, \text{inf})=35.13$, $p<.0001$. Responses were faster to the self ($EMM=740$, $SE=11.6$ [95% CI: 717–763]) in comparison to the friend ($EMM=805$, $SE=16.5$ [95% CI: 772–837]), $\beta=-64.8$ [95% CI: -95.70 to -33.86], $p<.0001$, and stranger ($EMM=844$, $SE=16.7$ [95% CI: 811–877]), $\beta=-104.1$ [95% CI: -135.4 to -72.8], $p<.0001$, replicating previous findings of self-biases. A friend-bias was also observed; responses were faster to the friend than to the stranger, $\beta=-34.9$ [95% CI: -77.6 to -1.15], $p=.04$. There was no significant main effect of prime $F(1, \text{inf})=0.60$, $p=.44$.

In contrast to Expt. 1a, no interaction between personal targets and prime was observed, $F(2, \text{inf})=0.54$, $p=.59$ (see Figure 3). Despite the non-significant interaction, the planned contrasts were still run. No prime biases were observed at targets ($ps>.31$), indicating no intrinsic personal-negative valences.

Discussion

In line with the prediction based on self-bivalence, negative primes did not alter self-bias sizes, compared to neutral primes. However, in contrast to the self-bivalence account, there were also no differences in that sad primes facilitated responses to self or other targets. The absence of any prime effect in

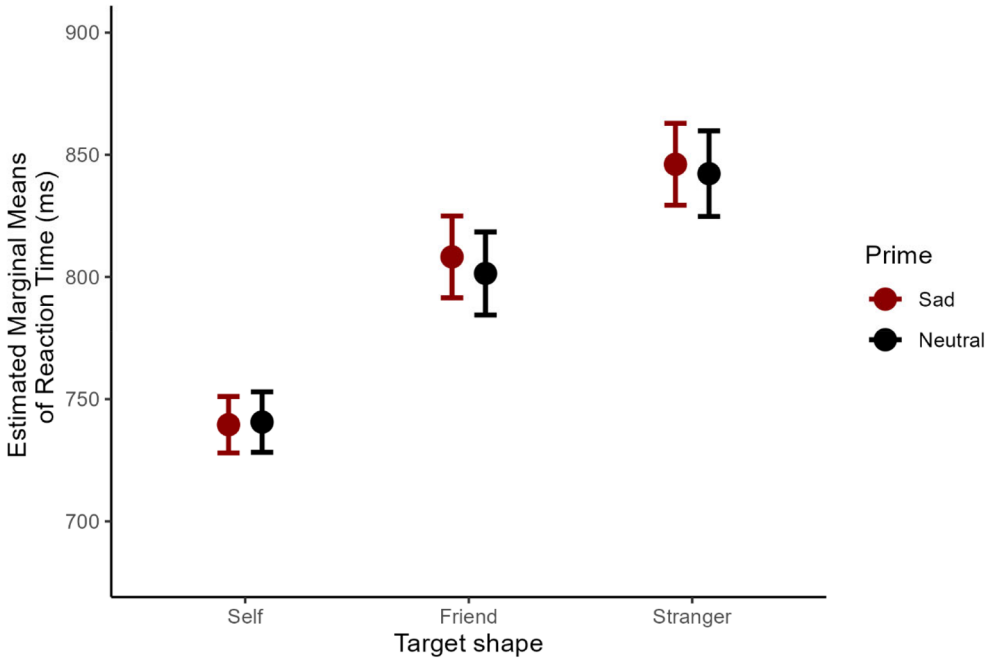


FIGURE 3 EMMs of RT as a function of Prime (Sad vs. Neutral) and Target (Self, Friend, or Stranger) in Experiment 1b. Error bars show standard error.

Experiment 1b, neither the reduced self-bias predicted by the self-positivity account nor the facilitation predicted by bivalence, suggests that while positivity might inhibit responses to the stranger, indicating that positivity was restricted solely to the self and extended self (i.e., friend), the stranger valence was not negative. Experiment 2 examined the complementary direction. Whether personal-identity primes selectively modulate emotional bias magnitudes.

EXPERIMENT 2: DO PERSON-BASED PRIMES INFLUENCE EMOTIONAL BIASES?

Introduction

Experiment 2 tested the directionality of self-valence interactions by reversing the priming logic of Experiment 1. In Experiment 1a and 1b, emotional primes preceded personal identity targets. In Experiment 2, personal-identity primes (self, stranger) preceded emotional targets (happy, neutral, sad). This design allowed a direct test of whether self-relevant information selectively modulates positivity-bias magnitudes, and whether any such modulation is driven by facilitation of happy target processing, inhibition of sad target processing, or both.

A key design difference between experiments reflects their distinct theoretical goals. In Experiment 1a and 1b, positive and negative primes were tested in separate experiments, each compared to a neutral baseline. This approach allowed the effects of positive and negative valence to be examined in isolation, minimizing potential carryover or contrast effects that may arise when multiple emotional primes are included within the same design. In contrast, Experiment 2 required a within-subjects manipulation of emotional targets (happy, neutral, sad). This allowed direct comparison of positivity-bias magnitudes across emotional conditions within the same participants, which is critical for assessing whether self-related primes selectively enhance positive processing. Accordingly, three emotional targets (happy,

neutral, and sad) were included, mirroring the categories used across Experiments 1a and 1b while enabling within-task comparison of valence effects.

Based on the asymmetric bidirectionality hypothesis, we predicted that, relative to stranger primes, self-primes would increase positivity-bias magnitudes, driven by faster responses to positive (happy) targets following self-priming.

Method

Participants

Target participant sample and participant exclusion criteria were identical to Experiment 1a. In total, 50 participants were recruited from SONA and Prolific (www.prolific.co); 8 were excluded due to performance at chance level. The average age of the 42 participants was 23.29 years ($SD=7.67$, range=18–61), consisting of 20 females, 20 males, and 2 participants who preferred not to answer; 1 participant was ambidextrous, 3 participants were left-handed, and 38 participants were right-handed.

Procedure

The procedure was identical to Experiment 1a and 1b, except the targets were three emotions (happy, neutral, and sad), and personal primes were used (Figure 4). The emotional stimuli were identical to those used as primes in Experiment 1a and 1b with a surrounding circle outline, and all lines were white. The primes used were the participant's first name and the participant's best friend's first name presented in black text.

Data analysis

The analysis was identical to Experiment 2 except the primes represented the self and a stranger, so the self-stranger prime bias was considered. The equation for the GLMM was:

$$\text{glmer}(\text{RT} \sim \text{emotional target} * \text{personal prime} + (1 + \text{emotional target} + \text{personal prime} | \text{participant}))$$

This includes fixed effects of emotional target, personal prime, and emotional target* personal prime interaction. Within random effects, each participant's intercept and slopes can be influenced by the targets and primes.

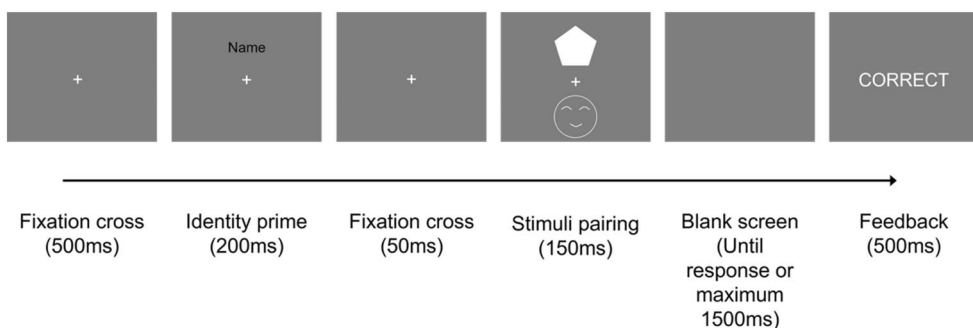


FIGURE 4 Example trial from Experiment 2.

Results

There was a significant main effect of emotional target $F(2, \text{inf})=39.86, p<.0001$. Replicating previous findings of positivity-biases, responses were faster to the happy emotion (EMM=754, SE=18.0 [95% CI: 719–789]) in comparison to the neutral emotion (EMM=804, SE=20.5 [95% CI: 764–844]), $\beta=-50.1$ [95% CI: -78.9 to -21.3], $p=.0001$, and sad emotion (EMM=856, SE=20.2 [95% CI: 816–895]), $\beta=-101.7$ [95% CI: -129.8 to -73.6], $p<.0001$. Responses were also faster to the neutral emotion in comparison to the sad emotion, $\beta=-51.6$ [95% CI: -88.6 to -14.6], $p=.003$. There was a non-significant main effect of prime $F(1, \text{inf})=3.21, p=.07$.

There was a significant interaction between emotional targets and prime $F(2, \text{inf})=7.40, p=.0006$ (see Figure 5). There were no differences in neutral-sad bias magnitudes, $\beta=-14.2$ [95% CI: -38.3 to 10.02], $p=.41$, or happy-neutral bias magnitudes, $\beta=-17.9$ [95% CI: -38.1 to 2.27], $p=.10$, between primes. However, the happy-sad bias magnitude was significantly larger following the self-prime in comparison to the stranger prime $\beta=-32.1$ [95% CI: -52.8 to -11.37], $p=.0006$. This is likely driven by a facilitatory self-advantage effect at the happy target; responses were significantly faster following the self-prime in comparison to the stranger prime, $\beta=-29.3$ [95% CI: -45.2 to -13.41], $p=.0003$. While there were no differences in response time following the different primes at the neutral, $\beta=-11.4$ [95% CI: -28.7 to 5.95], $p=.20$ or sad targets, $\beta=2.8$ [95% CI: -15.6 to 21.22], $p=.77$.

Discussion

Experiment 2 revealed a clear facilitatory effect of self-relevant primes on positive emotional processing. Self-priming enhanced the positivity bias, driven by faster responses to happy targets; neutral and sad targets were unaffected. This pattern directly contrasts with the mechanism revealed in Experiment

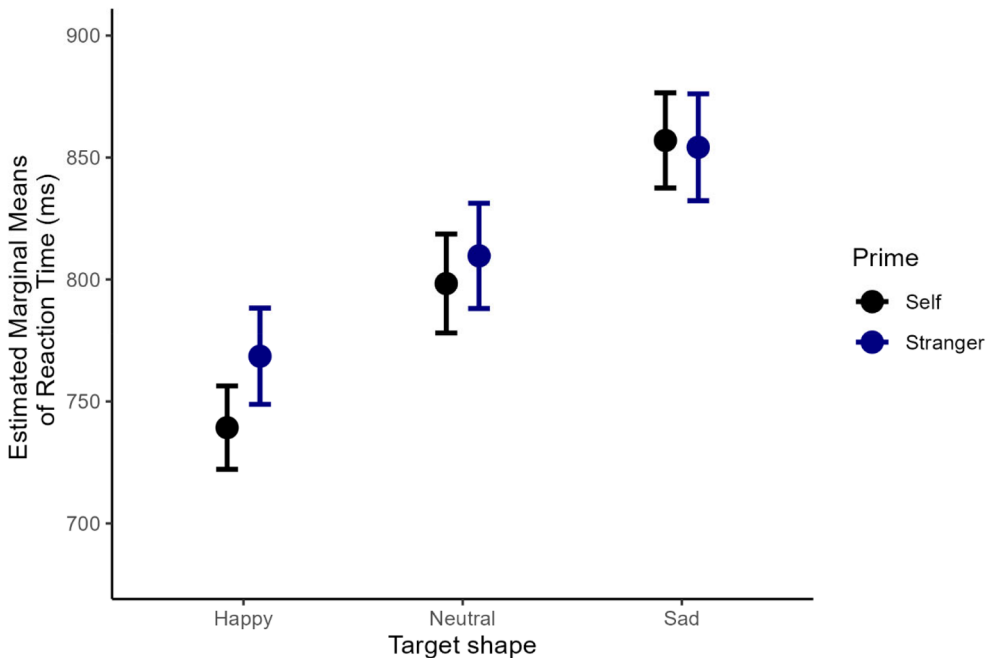


FIGURE 5 EMMs of RT as a function of Prime (Self vs. Friend) and Target (Happy, Neutral, or Sad) in Experiment 2. Error bars show standard error.

1a, where the self-stranger bias was enhanced not by faster self-target responses but by slower stranger-target responses. The complementarity of these findings is informative. Both experiments revealed enhanced self-positive processing, but through mechanistically distinct routes depending on the direction of priming. This double dissociation establishes that the mechanism linking self and positivity is a dynamic property of their association, depending critically on the direction from which the association is probed.

GENERAL DISCUSSION

Across three experiments, we examined the mechanisms and directionality of self-valence interactions by combining evaluative priming with a shape-label matching task. Two findings were observed. When emotional information preceded personal identity judgements, positive primes enhanced self-stranger bias through inhibition of stranger target processing, with no evidence of facilitation of self-target processing (Experiment 1a). When personal information preceded emotional judgements, self-primes enhanced the positivity-bias through facilitation of happy target processing, with no effect on negative or neutral targets (Experiment 2). Negativity valence produced no analogous effects in either direction (Experiment 1b). This pattern revealed an asymmetric bidirectional relationship between self and positivity. The association exists in both directions, but the dominant mechanism, inhibition or facilitation, depends on which member of the pair serves as prime.

The asymmetry in mechanisms carries direct implications for self-concept maintenance. This pattern aligns with self-enhancement theory (Sedikides & Gregg, 2008), suggesting that self-biases are driven by self-enhancement such that self-positive information receives prioritized processing to maintain intrinsic self-positivity and healthy mental wellbeing (Lee et al., 2023). However, self-enhancement theory does not obviously predict that the mechanism should differ by direction: why should the self to positivity link operate through facilitation while the positivity to self-link operates through inhibition? Understanding this requires a more mechanistic theoretical framework.

A predictive processing account offers a candidate (Albarracín et al., 2024; Friston & Stephan, 2007). If self-representations function as strong priors generating probabilistic expectations about upcoming information, then self-relevant primes activate a prediction of positive content. When this prediction is confirmed by a happy target, prediction error is minimized, facilitating response. This account explains the facilitation observed in Experiment 2. The converse case is different. A positive prime does not generate a specific prediction of self-relevance. Rather, it generates a broader prediction of self or social proximity that places strangers outside the expected range. The result is not facilitation of self or friend targets, whose responses are in the expected direction but do not reach significance, but inhibition of stranger-target processing, because stranger-positivity represents a violated prediction. On this view, the asymmetry in mechanism reflects an asymmetry in prior probability. The self predicts positivity strongly and specifically, while positivity predicts against stranger-relevance rather than for self-relevance.

An alternative account involves reward-based encoding. Self-associated stimuli may automatically engage reward systems (de Greck et al., 2010; Enzi et al., 2009; Northoff & Hayes, 2011; Sui et al., 2015; Sui & Humphreys, 2015; Yankouskaya et al., 2017), with positive valence similarly activating reward networks. The convergence of these reward signals in self-positive contexts could enhance processing efficiency through increased motivational salience. This reward-based account is consistent with neuroimaging evidence showing overlapping activation in the ventral striatum and medial prefrontal cortex for both self-related and reward-related stimuli, suggesting a shared representational substrate (Enzi et al., 2009; Fan et al., 2016; Sui & Humphreys, 2015; Yankouskaya et al., 2017). If positive valence and self-relevance both activate reward-related circuitry, their co-occurrence in self-positive combinations may produce enhanced processing through additive reward signals, explaining the self to positive facilitation in Experiment 2. In the reverse direction, the failure of stranger-positive combinations to resolve either self-relevance or positive valence together may manifest as competitive suppression,

explaining the stranger-positive inhibition in Experiment 1a. These two accounts, predictive processing and reward coding, make overlapping predictions and may describe complementary aspects of the same mechanism. Future work examining both priming directions simultaneously would be informative.

The distinction between facilitatory self-positive associations (self to emotion direction) and inhibitory stranger-positive associations (emotion to self-direction) may have practical implications beyond theory. If self-positive biases serve an adaptive function in maintaining wellbeing (Lee et al., 2023; Sui et al., 2021; Wang et al., 2023), then clinical conditions characterized by reduced self-positivity may involve selective disruption of one or both mechanisms. Training-based interventions designed to strengthen self-positive associations, such as self-positivity bias training, may benefit from knowing which component of the association is disrupted. The current design provides a methodological template for characterizing the specific directional and mechanistic profile of self-valence processing in clinical populations.

The selective inhibition of stranger-target processing following happy primes, with no reliable effect on friend-target processing, suggests that the relevant boundary is social distance rather than a strict self-other distinction. The friend occupies an intermediate position between self and stranger and is sometimes conceptualized as part of an extended self (Jiang & Sui, 2025; Sun et al., 2016). If positive valence generates expectations of self and social proximity, friend targets may be partially spared from the inhibitory consequences directed at strangers. The current data are consistent with this interpretation; friend targets show a non-significant trend toward facilitation following happy primes, but the effect does not reach significance and should not be over-read. Parametrically varying social distance (e.g., family, best friend, acquaintance, stranger) in future work would characterize the full gradient of this effect and clarify whether inhibition of positive-prime responses scales continuously with social proximity.

The absence of self-negativity effects across Experiment 1b and 2 contrasts with recent findings of implicit self-evaluation bivalence (Ni et al., 2025; Shi et al., 2025; Zayas et al., 2022), which demonstrate that self-primes facilitate responses to both positive and negative emotional targets. This discrepancy may reflect differences in task demands and the nature of the relationships being examined. Previous bivalence studies used evaluative priming tasks where participants explicitly judged whether targets were positive or negative, whereas the current experiments required perceptual discrimination of shape-label identity or emotional category, tasks with minimal evaluative demands. If self-negativity associations exist at the level of explicit evaluation but do not modulate low-level perceptual matching processes, then bivalence and null negativity effects are reconcilable with a dual-process framework. Implicit associations may require evaluative elaboration to influence deliberate responding, yet remain below threshold for influencing perceptual matching. Alternatively, self-negativity exists but is weaker than self-positivity associations in healthy participants, requiring different study contexts or measurement approaches to be revealed in perceptual tasks. Future research comparing evaluative and perceptual tasks within the same individuals would directly test these accounts.

Constraints on generality

The present findings are based on online recruited samples of predominantly young adults that are fluent in English. Therefore, the results may not generalize to older populations that have previously been shown to have altered self-biases (Sui & Humphreys, 2017) or individuals with clinical conditions such as depression who have reduced self-positivity biases (Wang et al., 2023). Due to the research focusing on low-level perceptual processes, the tasks involved abstract shape-label associations with schematic emotional faces which differ from naturalistic social interactions, limiting the ecological validity of the paradigm. This is in line with previous work (Constable et al., 2021), which allowed for directionality effects explored here to be considered in relation to previous findings. It would be interesting for future research to explore whether effects may vary with richer stimuli. Additionally, only happy, neutral, and sad emotions were examined, so results should not be assumed to extend to other affective states.

Finally, the observed effects reflect perceptual matching processes under controlled conditions and may not apply to explicit evaluations or real-world contexts.

CONCLUSION

Self and positivity are linked not by a single mechanism but by two that depend on direction. When emotional information precedes self-other judgements, positive primes enhance self-bias through inhibition of distal social targets. When self-relevant information precedes emotional judgements, self-primes enhance positivity bias through facilitation of congruent positive targets. Negative valence does not reciprocate this relationship at the level of perceptual processing. These findings extend self-enhancement theory and enrich predictive processing accounts of social cognition by demonstrating that self-positivity is mechanistically more complex than a simple association. It is a directionally asymmetric relationship in which distinct computational processes converge to maintain the privileged link between self and the good.

AUTHOR CONTRIBUTIONS

Naomi A. Lee: Conceptualization; methodology; formal analysis; investigation; data curation; writing – original draft; writing – review and editing; project administration. **Douglas Martin:** Conceptualization; writing – review and editing; supervision. **Jie Sui:** Conceptualization; methodology; writing – review and editing; supervision; funding acquisition.

FUNDING INFORMATION

This work was supported by the Leverhulme Trust (RPG-2019-010).

CONFLICT OF INTEREST STATEMENT

The authors declare no conflicts of interest.

DATA AVAILABILITY STATEMENT

The data that support the findings of this study are openly available in Open Science Framework at https://osf.io/enfma/overview?view_only=a2778900c8ce490085c571c3dbe1af76.

ETHICS STATEMENT

The research protocol was approved by the University of Aberdeen Research Ethics Committee. All participants provided informed consent prior to participation.

DISCLOSURE STATEMENT

All studies, measures, manipulations, and data/participant exclusions are reported in the manuscript.

ORCID

Naomi A. Lee  <https://orcid.org/0000-0002-0973-6394>

Douglas Martin  <https://orcid.org/0000-0002-9718-7662>

REFERENCES

- Albarracín, M., Bouchard-Joly, G., Sheikhbahe, Z., Miller, M., Pitliya, R. J., & Poirier, P. (2024). Feeling our place in the world: An active inference account of self-esteem. *Neuroscience of Consciousness*, 2024(1), niae007. <https://doi.org/10.1093/nc/niae007>
- Alexopoulos, T., Muller, D., Ric, F., & Marendaz, C. (2012). I, me, mine: Automatic attentional capture by self-related stimuli. *European Journal of Social Psychology*, 42(6), 770–779. <https://doi.org/10.1002/ejsp.1882>
- Alicke, M. D., & Sedikides, C. (2009). Self-enhancement and self-protection: What they are and what they do. *European Review of Social Psychology*, 20(1), 1–48. <https://doi.org/10.1080/10463280802613866>

- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Bates, D., Maechler, M., Bolker, B., Walker, S., Christensen, R. H. B., Singmann, H., Dai, B., Scheipl, F., Grothendieck, G., Green, P., Fox, J., Bauer, A., Krivitsky, P. N., Tanaka, E., Jagan, M., & Boylan, R. D. (2025). *lme4: Linear mixed-effects models using 'Eigen' and S4* (Version 1.1–37) [Computer software]. <https://cran.r-project.org/web/packages/lme4/index.html>
- Chen, Y., Zhong, Y., Zhou, H., Zhang, S., Tan, Q., & Fan, W. (2014). Evidence for implicit self-positivity bias: An event-related brain potential study. *Experimental Brain Research*, 232(3), 985–994. <https://doi.org/10.1007/s00221-013-3810-z>
- Cherry, E. C. (1953). Some experiments on the recognition of speech, with one and with two ears. *The Journal of the Acoustical Society of America*, 25(5), 975–979. <https://doi.org/10.1121/1.1907229>
- Constable, M. D., Becker, M. L., Oh, Y.-I., & Knoblich, G. (2021). Affective compatibility with the self modulates the self-prioritisation effect. *Cognition and Emotion*, 35(2), 291–304. <https://doi.org/10.1080/02699931.2020.1839383>
- Dalmaso, M., Vicovaro, M., Saito, T., & Watanabe, K. (2026). We are what we eat: Cross-cultural self-prioritization effects for food stimuli. *British Journal of Psychology*, 117(1), 177–191. <https://doi.org/10.1111/bjop.70018>
- de Greck, M., Enzi, B., Prösch, U., Gantman, A., Tempelmann, C., & Northoff, G. (2010). Decreased neuronal activity in reward circuitry of pathological gamblers during processing of personal relevant stimuli. *Human Brain Mapping*, 31(11), 1802–1812. <https://doi.org/10.1002/hbm.20981>
- Enzi, B., de Greck, M., Prösch, U., Tempelmann, C., & Northoff, G. (2009). Is our self nothing but reward? Neuronal overlap and distinction between reward and personal relevance and its relation to human personality. *PLoS One*, 4(12), e8429. <https://doi.org/10.1371/journal.pone.0008429>
- Fan, W., Zhong, Y., Li, J., Yang, Z., Zhan, Y., Cai, R., & Fu, X. (2016). Negative emotion weakens the degree of self-reference effect: evidence from ERPs. *Frontiers in Psychology*, 7. <https://doi.org/10.3389/fpsyg.2016.01408>
- Fazio, R. H., Sanbonmatsu, D. M., Powell, M. C., & Kardes, F. R. (1986). On the automatic activation of attitudes. *Journal of Personality and Social Psychology*, 50(2), 229–238. <https://doi.org/10.1037/0022-3514.50.2.229>
- Fields, E. C., & Kuperberg, G. R. (2015). Loving yourself more than your neighbor: ERPs reveal online effects of a self-positivity bias. *Social Cognitive and Affective Neuroscience*, 10(9), 1202–1209. <https://doi.org/10.1093/scan/nsv004>
- Friston, K. J., & Stephan, K. E. (2007). Free-energy and the brain. *Synthese*, 159(3), 417–458. <https://doi.org/10.1007/s1122-9-007-9237-y>
- Golubickis, M., Ho, N. S. P., Falbén, J. K., Schwertel, C. L., Maiuri, A., Dublas, D., Cunningham, W. A., & Macrae, C. N. (2021). Valence and ownership: Object desirability influences self-prioritization. *Psychological Research*, 85(1), 91–100. <https://doi.org/10.1007/s00426-019-01235-w>
- Higgins, E. T. (1987). Self-discrepancy: A theory relating self and affect. *Psychological Review*, 94(3), 319–340. (1987-34444-001). <https://doi.org/10.1037/0033-295X.94.3.319>
- Hung, K., Feldborg, M., Moseley, R., Peng, K., & Sui, J. (2025). Digital selfie editing shows sex specific associations between processing biases and life satisfaction. *Scientific Reports*, 15(1), 14235. <https://doi.org/10.1038/s41598-025-99056-y>
- Jiang, M., & Sui, J. (2025). The social self: Categorisation of family members examined through the self-bias effect in new mothers. *Quarterly Journal of Experimental Psychology*, 78(12), 2816–2828. <https://doi.org/10.1177/17470218251332905>
- Klinger, M. R., & Greenwald, A. G. (1995). Unconscious priming of association judgments. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(3), 569–581. <https://doi.org/10.1037/0278-7393.21.3.569>
- Lee, N. A., Martin, D., & Sui, J. (2023). Accentuate the positive: Evidence that context dependent self-reference drives self-bias. *Cognition*, 240, 105600. <https://doi.org/10.1016/j.cognition.2023.105600>
- Lenth, R. V., Banfai, B., Bolker, B., Buerkner, P., Giné-Vázquez, I., Herve, M., Jung, M., Love, J., Miguez, F., Piskowski, J., Riebl, H., & Singmann, H. (2025). *emmeans: Estimated marginal means, aka least-squares means* (Version 1.1.1) [Computer software]. <https://cran.r-project.org/web/packages/emmeans/index.html>
- Lo, S., & Andrews, S. (2015). To transform or not to transform: Using generalized linear mixed models to analyse reaction time data. *Frontiers in Psychology*, 6(1171), 1–16. <https://doi.org/10.3389/fpsyg.2015.01171>
- Markus, H. R., & Kitayama, S. (1991). Cultural variation in the self-concept. In J. Strauss & G. R. Goethals (Eds.), *The self: Interdisciplinary approaches* (pp. 18–48). Springer. https://doi.org/10.1007/978-1-4684-8264-5_2
- Moray, N. (1959). Attention in dichotic listening: Affective cues and the influence of instructions. *Quarterly Journal of Experimental Psychology*, 11(1), 56–60. <https://doi.org/10.1080/17470215908416289>
- Ni, M., Lee, R. T., & Zayas, V. (2025). The bittersweetness of self-representations: Cross-cultural insights that the self triggers both positive and negative implicit evaluations. *Journal of Cross-Cultural Psychology*, 00220221251355914. <https://doi.org/10.1177/00220221251355914>
- Northoff, G., & Hayes, D. J. (2011). Is our self nothing but reward? *Biological Psychiatry, Nucleus Accumbens Neuroadaptations and Relapse in Addiction*, 69(11), 1019–1025. <https://doi.org/10.1016/j.biopsych.2010.12.014>
- R Core Team. (2023). *R: A language and environment for statistical computing* [Computer software]. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rogers, T. B., Kuiper, N. A., & Kirker, W. S. (1999). Self-reference and the encoding of personal information. In *The self in social psychology* (pp. 139–149). Psychology Press.
- Sedikides, C., & Gregg, A. P. (2008). Self-enhancement: Food for thought. *Perspectives on Psychological Science*, 3(2), 102–116. <https://doi.org/10.1111/j.1745-6916.2008.00068.x>

- Shi, K., Li, S., Du, T., Zang, X., Li, J., Tang, J., Wang, Y., & Yang, J. (2025). Implicit self-evaluation bivalence: Measurement and computation. *Personality and Individual Differences*, 247, 113409. <https://doi.org/10.1016/j.paid.2025.113409>
- Stolte, M., Humphreys, G., Yankouskaya, A., & Sui, J. (2017). Dissociating biases towards the self and positive emotion. *Quarterly Journal of Experimental Psychology*, 70(6), 1011–1022. <https://doi.org/10.1080/17470218.2015.1101477>
- Sui, J., Greenshaw, A. J., Macrae, C. N., & Cao, B. (2021). Self research: A new pathway to precision psychiatry. *Journal of Affective Disorders*, 293, 276–278. <https://doi.org/10.1016/j.jad.2021.06.041>
- Sui, J., He, X., & Humphreys, G. W. (2012). Perceptual effects of social salience: Evidence from self-prioritization effects on perceptual matching. *Journal of Experimental Psychology: Human Perception and Performance*, 38(5), 1105–1117. <https://doi.org/10.1037/a0029792>
- Sui, J., & Humphreys, G. W. (2015). The interaction between self-bias and reward: Evidence for common and distinct processes. *Quarterly Journal of Experimental Psychology*, 68(10), 1952–1964. <https://doi.org/10.1080/17470218.2015.1023207>
- Sui, J., & Humphreys, G. W. (2017). Aging enhances cognitive biases to friends but not the self. *Psychonomic Bulletin & Review*, 24(6), 2021–2030. <https://doi.org/10.3758/s13423-017-1264-1>
- Sui, J., Yankouskaya, A., & Humphreys, G. W. (2015). Super-capacity me! Super-capacity and violations of race independence for self-but not for reward-associated stimuli. *Journal of Experimental Psychology: Human Perception and Performance*, 41(2), 441–452. <https://doi.org/10.1037/a0038288>
- Sun, Y., Fuentes, L. J., Humphreys, G. W., & Sui, J. (2016). Try to see it my way: Embodied perspective enhances self and friend-biases in perceptual matching. *Cognition*, 153, 108–117. <https://doi.org/10.1016/j.cognition.2016.04.015>
- Vicovaro, M., Dalmaso, M., & Bertamini, M. (2022). Towards the boundaries of self-prioritization: Associating the self with asymmetric shapes disrupts the self-prioritization effect. *Journal of Experimental Psychology: Human Perception and Performance*, 48(9), 972–986. <https://doi.org/10.1037/xhp0001036>
- Wang, L., Song, X., Jia, F., Xue, W., & Li, L. (2023). Effectiveness of enhanced self-positivity bias training in mitigating depressive mood following negative life events. *Behavioral Science*, 13(7), 7. <https://doi.org/10.3390/bs13070534>
- Watson, L. A., Dritschel, B., Obonsawin, M. C., & Jentsch, I. (2007). Seeing yourself in a positive light: Brain correlates of the self-positivity bias. *Brain Research*, 1152, 106–110. <https://doi.org/10.1016/j.brainres.2007.03.049>
- Wentura, D., & Rothermund, K. (2014). Priming is not priming is not priming. *Social Cognition*, 32(Supplement), 47–67. <https://doi.org/10.1521/soco.2014.32.sup.47>
- Xia, R., Shao, H., Cui, L., Zhang, P., Xue, J., Zhou, A., & Li, S. (2021). Evidence for self-positivity bias in a subliminal self-cue: An event-related potential study. *Neuroscience Letters*, 744(135625), 1–6. <https://doi.org/10.1016/j.neulet.2021.135625>
- Yankouskaya, A., Humphreys, G., Stolte, M., Stokes, M., Moradi, Z., & Sui, J. (2017). An anterior–posterior axis within the ventromedial prefrontal cortex separates self and reward. *Social Cognitive and Affective Neuroscience*, 12(12), 1859–1868. <https://doi.org/10.1093/scan/nsx112>
- Yankouskaya, A., & Sui, J. (2021). Self-positivity or self-negativity as a function of the medial prefrontal cortex. *Brain Sciences*, 11(2), 264. <https://doi.org/10.3390/brainsci11020264>
- Zayas, V., Wang, A. M., & McCalla, J. D. (2022). Me as good and me as bad: Priming the self triggers positive and negative implicit evaluations. *Journal of Personality and Social Psychology*, 122(1), 106–134. <https://doi.org/10.1037/pspp0000332>
- Zhang, W., Lan, Y., Li, Q., Gao, Z., Hu, J., & Gao, S. (2023). The foreign-language effect on self-positivity bias: Behavioral and electrophysiological evidence. *Bilingualism: Language and Cognition*, 26(3), 570–579. <https://doi.org/10.1017/S1366728922000827>

SUPPORTING INFORMATION

Additional supporting information can be found online in the Supporting Information section at the end of this article.

How to cite this article: Lee, N. A., Martin, D., & Sui, J. (2026). Happy me: Asymmetric bidirectional links between the self and positivity underpin biased processing. *British Journal of Psychology*, 00, 1–16. <https://doi.org/10.1111/bjop.70100>