



Universiteit
Leiden
The Netherlands

The mystery remains: breadth of attention in flanker and Navon tasks unaffected by affective states induced by an appraisal manipulation

Kolnes, M.; Gentsch, K.; Steenbergen, H. van; Uusberg, A.

Citation

Kolnes, M., Gentsch, K., Steenbergen, H. van, & Uusberg, A. (2022). The mystery remains: breadth of attention in flanker and Navon tasks unaffected by affective states induced by an appraisal manipulation. *Cognition And Emotion*, 36(5), 836-854.
doi:10.1080/02699931.2022.2056580

Version: Publisher's Version

License: [Licensed under Article 25fa Copyright Act/Law \(Amendment Taverne\)](#)

Downloaded from: <https://hdl.handle.net/1887/3512611>

Note: To cite this publication please use the final published version (if applicable).



The mystery remains: breadth of attention in Flanker and Navon tasks unaffected by affective states induced by an appraisal manipulation

Martin Kolnes, Kornelia Gentsch, Henk van Steenbergen & Andero Uusberg

To cite this article: Martin Kolnes, Kornelia Gentsch, Henk van Steenbergen & Andero Uusberg (2022) The mystery remains: breadth of attention in Flanker and Navon tasks unaffected by affective states induced by an appraisal manipulation, *Cognition and Emotion*, 36:5, 836-854, DOI: [10.1080/02699931.2022.2056580](https://doi.org/10.1080/02699931.2022.2056580)

To link to this article: <https://doi.org/10.1080/02699931.2022.2056580>



View supplementary material [↗](#)



Published online: 31 Mar 2022.



Submit your article to this journal [↗](#)



Article views: 399



View related articles [↗](#)



View Crossmark data [↗](#)



The mystery remains: breadth of attention in Flanker and Navon tasks unaffected by affective states induced by an appraisal manipulation

Martin Kolnes^a, Kornelia Gentsch^b, Henk van Steenbergen^{c*} and Andero Uusberg^{a*}

^aInstitute of Psychology, University of Tartu, Tartu, Estonia; ^bDepartment of Clinical Psychology und Neuropsychology, Institute of Psychology, Johannes-Gutenberg-University, Mainz, Germany; ^cInstitute of Psychology and Leiden Institute for Brain and Cognition, Leiden University, Leiden, The Netherlands

ABSTRACT

Affective effects on breadth of attention have been related to aspects of different components of affective states such as the arousal and valence of affective experience and the motivational intensity of action tendency. As none of these explanations fully aligns with existing evidence, we hypothesised that affective effects on breadth of attention may arise from the appraisal component of affective states. Based on this reconceptualisation, we tested the effects of conduciveness and power appraisals on two measures of breadth of attention. In two web-based experiments, we manipulated these appraisals in a 2 × 2 design using a game-like arithmetic task where participants could (1) gain or lose rewards (goal conducive vs. obstructive) based on (2) either their action or the actions of a “robot” (high vs. low power). Breadth of attention was assessed using the flanker task (Experiment 1; $n = 236$) and the Navon task (Experiment 2; $n = 215$). We found that appraisals did not directly influence breadth of attention even though high power appraisal significantly improved the overall performance in both experiments indicating successful appraisal manipulation. We discuss ways in which these findings inform future efforts to explain the origins of affective effects on attentional breadth.

ARTICLE HISTORY

Received 21 June 2021
Revised 14 February 2022
Accepted 17 March 2022

KEYWORDS

Breadth of attention;
emotions; appraisal theory;
goal conduciveness
appraisal; power appraisal

Attentional processes are sensitive to affective states that we view as multi-component states sharing some but not all components of emotions (Gross et al., 2019). In addition to concentrating attention on affective stimuli (Vuilleumier, 2015), affective states can also modulate breadth of attention. Breadth of attention refers to the size of the attended area of the visual field (Eriksen & Eriksen, 1974) or the priority of local elements vs. overall shape in stimulus processing (Navon, 1977). Affective influences on breadth of attention have been attributed to different affective components. Accounts focusing on arousal (Easterbrook, 1959) and valence (Fredrickson, 2004) attribute these effects to the experiential component. Accounts focusing on motivational

intensity attribute them to the action tendency component (Gable & Harmon-Jones, 2010b; Kaplan et al., 2012). However, the empirical evidence for each of these accounts is mixed and it remains unclear which component of affect drives its impact on attentional breadth (Clore & Huntsinger, 2007; Friedman & Forster, 2011; Huntsinger, 2013). In this study, we tested a novel idea that affective effects on breadth of attention may be ultimately driven by appraisals of the subjective meaning of the given situation.

Affective shifts in breadth of attention were first attributed to the arousal and valence dimensions of affective experience. Classic studies from Easterbrook (1959) suggested that affect with high arousal narrows breadth of attention (Mather & Sutherland,

CONTACT Martin Kolnes  martin.kolnes@ut.ee  Department of Psychology, University of Tartu, Tartu 50409, Estonia

*Both authors contributed equally to this work.

 Supplemental data for this article can be accessed online at <https://doi.org/10.1080/02699931.2022.2056580>.

This research was not preregistered. The analysis script and data are openly available in OSF repository at <https://osf.io/rq2nw/>.

© 2022 Informa UK Limited, trading as Taylor & Francis Group

2011). However, this line of research focused mainly on negative states and thus potentially confounded arousal effects with valence-arousal interactions (Steenbergen et al., 2011). Alternative models focusing on valence suggest that positive emotions, such as amusement and contentment, broaden breadth of attention while negative emotions narrow it (Fredrickson, 2004; Fredrickson & Branigan, 2005). Although some studies demonstrate this relationship (e.g. Fredrickson, 2004; Rowe et al., 2007), there are also null findings (Bruyneel et al., 2013; van Steenbergen, 2015; Vanlessen et al., 2016) and opposite findings such that negative emotions broaden breadth of attention (Gable & Harmon-Jones, 2010a; Huntsinger, 2013; von Hecker & Meiser, 2005). In summary, although arousal and valence influence breadth of attention, these dimensions of affective experience fail to fully explain the heterogeneous findings.

Affective effects on breadth of attention have also been linked to motivational intensity, defined as the strength of the motivation (Gable & Harmon-Jones, 2010b) generated as part of the action tendency component of affective states (Gable & Harmon-Jones, 2013). By this account, the intensity of both approach- and avoidance-oriented motivation narrows breadth of attention. For example, joy – an emotion with low approach motivational intensity – should broaden breadth of attention whereas desire – an emotion with high approach motivational intensity – should narrow it. Likewise, sadness (low avoidance motivational intensity) should widen breadth of attention whereas fear (high avoidance motivational intensity) should narrow it (Gable & Harmon-Jones, 2010b; Lacey et al., 2021). However, this account also cannot explain all findings (Clore & Huntsinger, 2007; Friedman & Forster, 2011; Huntsinger, 2013).

The existing literature has thus sought to attribute affective effects on breadth of attention to the experiential and to the action tendency components of affect, but neither can fully explain the empirical findings. One way to explain this pattern of results is to assume that some other process that contributes to both affective components influences breadth of attention. We propose that one such process is appraisal, a cognitive component of emotion that decodes the motivational meaning of a situation (i.e. its relation to current individual goals and concerns) and orchestrates changes in other emotion components including experience and action tendency (Moors et al., 2013; Scherer, 2009). For instance, a

person attributing an unpleasant social encounter to oneself is more likely to experience guilt whereas a person attributing a similar encounter to their partner is likely to experience anger (Siemer et al., 2007). In addition to emotional experiences, appraisals have been shown to shape facial and vocal expressions, autonomic changes, and action tendencies accompanying emotions (Moors et al., 2013). Appraisals also appear to underlie affective impacts on different attentional processes (Kolnes et al., 2019; Schimmack & Derryberry, 2005). Changes to attentional breadth may thus be among the set of functional effects appraisals have on the body and the mind.

Although the role of appraisals in modulating breadth of attention has not been systematically investigated, there is preliminary evidence consistent with this view. First, breadth of attention is sensitive to goal status: it is narrower when a goal is still being pursued compared to when it is accomplished or abandoned (Gable & Harmon-Jones, 2010b; Kaplan et al., 2012). Even though this effect was interpreted in terms of motivational intensity, goal status can also be considered an appraisal dimension. Second, shifts in breadth of attention do not necessarily require conscious affective experience (Friedman & Förster, 2010) suggesting they may arise from appraisals which can also be unconscious (Moors et al., 2013). In summary, even as researchers have not always articulated it, several existing findings are consistent with the idea that affective shifts in breadth of attention are driven by emotion-antecedent appraisal processes.

To more directly test the idea that appraisals influence breadth of attention, we focused on two appraisal dimensions – goal conduciveness and power appraisals – as probable modulating sources of breadth of attention. Goal conduciveness appraisal concerns whether an event is conducive or obstructive for reaching current goals (Scherer, 2009). Appraising events as goal-conducive may broaden breadth of attention because attentional resources can be used to explore new options in a benign situation (Carver, 2003). On the other hand, appraising events as goal-obstructive may narrow breadth of attention because attentional resources should be focused on resolving problems in a problematic situation. Traditionally, goal conduciveness has been associated with the valence dimension of the affective experience (Frijda et al., 1989). A goal conducive event, like winning in a game, elicits positive

emotions, and a goal obstructive event, like losing in a game, elicits negative emotions. Goal conduciveness appraisal may thus underlie the observed effects of valence on breadth of attention.

Power appraisal evaluates the resources at one's disposal to change contingencies and outcomes according to current goals (Scherer, 2009). High power appraisal should be associated with narrowed breadth of attention to direct attentional resources to the aspects of the situation that enable to change it. By contrast, low power appraisal should be associated with a widened breadth of attention because it implies that one should look out for new ways to increase personal influence over the situation. Power appraisals are associated with the intensity of action tendencies. Situations that are relatively easy to change and thus appraised as high in power tend to produce affective states with high motivational intensity, such as desire and disgust. Situations that are more difficult to change and are thus appraised as low in power tend to generate low motivational intensity states, like contentment and sadness. Power appraisal may thus underlie the observed effects of motivational intensity and goal status.

1.2. Present research

The present study aimed to examine whether goal conduciveness and power appraisals influence breadth of attention. In two experiments, we used a game-like task, adapted from Gentsch et al. (2013), to manipulate goal conduciveness (conductive vs. obstructive, i.e. possibility to gain vs. lose) and power appraisals (high vs. low power, i.e. self vs. robot) in a two-by-two design. We used large web-based samples and a within-subjects design to achieve high statistical power. Post-hoc power analysis showed that with 0.8 power both experiments were able to detect small effect sizes (Cohen's $d = 0.20$).

Considering that different breadth of attention measures might not reflect the same aspects of this process (Dale & Arnell, 2013), we used two well-known tasks. In Experiment 1, we used the flanker task where participants have to respond to a central arrow that is flanked by pairs of same or different kinds of arrows on both sides (Eriksen & Eriksen, 1974). In Experiment 2, we used a modified Navon task (Navon, 1977) where the stimuli were large letters (e.g. "H", the global level of the stimulus) made from small letters (e.g. "F", the local level of

the stimulus). On each trial, participants needed to decide whether one of two target letters ("H" or "T") was present on the local or the global level of the Navon stimulus.

To further isolate appraisal effects on cognitive processing of the tasks, we used drift-diffusion modelling (DDM; Ratcliff, 1978) that is designed to disentangle different response strategies in simple two-choice reaction time tasks (Ratcliff & McKoon, 2008). DDM assumes that decision making is a dynamic process where information in favour of one of two alternatives accumulates over time until it reaches a decision boundary (Ratcliff & McKoon, 2008). We will focus on the *drift rate* parameter of DDM that describes the rate of information accumulation which should change as a function of appraisal.

According to our hypotheses, breadth of attention should be narrower in goal obstructive trials (i.e. when it is possible to lose points) and in high power trials (i.e. when it is possible to influence the outcome). We expect that narrow breadth of attention should enhance processing of the central target in the Flanker task and the local level stimuli within Navon letters. Thus, goal obstructive trials and high power trials should exhibit (a) faster reaction times, (b) fewer errors, and (c) higher drift rates for flanker trials where the central arrows were flanked by different rather than the same arrows and for local Navon trials where the target letter was on the local rather than the global level. On the other hand, breadth of attention should be broader in goal conducive trials (i.e. possibility to win) and low power trials (i.e. not possible to influence the outcome). These effects should manifest in (a) faster reaction times, (b) fewer errors, and (c) higher drift rates for flanker trials where the central arrows were flanked by same rather than different arrows and for global Navon trials where the target letter was on the global level.

2. Experiment 1

In Experiment 1, we asked whether manipulations of goal conduciveness and power appraisals influence breadth of attention in the Flanker task (e.g. Liu et al., 2016). In the task, the participant had to respond to a central arrow that was flanked by a pair of same or different arrows on both sides, resulting in response-compatible and response-incompatible trials, respectively. The response-incompatible trials typically increase response times and error rates

presumably because the participant must overcome interference from the flanking arrows to arrive at a correct decision. This should also be reflected in a lower drift rate. Narrower breadth of attention should decrease that interference while broader breadth of attention should increase it. Smaller difference between response-compatible and response-incompatible flanker trials, therefore, indicates a narrower breadth of attention, while a larger difference between the two types of flanker trials suggests a wider breadth of attention.

2.1. Method

2.1.1. Participants

Initially, 236 participants (age: $M = 27.62$, $SD = 9.44$; 166 females) completed the web-based experiment. We excluded 27 participants from the final sample based on unreliable flanker task performance: 5 had a high error rate (over 30%) and 22 participants had less than 50% trials left in at least one condition after removing outlying response times. The final sample comprised 209 participants (age: $M = 27.19$, $SD = 8.82$; 147 females), of whom 200 also completed the post-experiment questionnaire.

The participants were recruited through community and campus mailings lists, and advertisement via Facebook. The recruitment text explained that in the experiment participants can test their arithmetic skills in a game-like task. For additional motivation, participants were told that they will be included in a raffle of a 100€-gift card if their final score was among the top 20 at the end of the study. The study was approved by the Institutional Review Board of University of Tartu.

2.1.2. Procedure

The whole study was web-based. Participants were instructed to choose a quiet time and place for completing the study. They started with the informed consent form and a short demographic questionnaire presented on Google Forms. They were then directed to the experiment programmed with Psychopy 3 (Peirce et al., 2019) and hosted on Pavlov.org. Upon completion, they were routed to the post-experiment questionnaire.

At the start of the experiment, participants were asked to sit 60 cm from their monitor (average reported distance $M = 56.77$ cm, $SD = 11.00$ cm) and to resize a rectangle on the screen until it matched a credit card or a similarly sized card. Next,

participants had to choose a female or a male icon to represent themselves throughout the experiment. Participants were instructed to acquire as many points as possible by responding correctly to the arithmetic task. To facilitate fast responding, they were also instructed to hold their hand over the arrow keys.

In the experiment, the flanker task was embedded within an arithmetic task designed to manipulate appraisals (Figure 1). Goal conduciveness appraisal was manipulated through available outcomes in the arithmetic task. On goal conducive trials, the participant could either win or not win, but never lose points. On goal obstructive trials, participants could either lose or not lose, but never win points. Power appraisal was manipulated through the possibility to influence the outcome of the trial. On high-power trials, the outcome depended on the participant. On low-power trials, the outcome depended on the choice that was made by the program, represented by a robot icon.

The experiment consisted of 128 trials divided into four randomly sequenced and otherwise identical blocks of 32 trials. Overall, eight trial types were used (all appraisal combinations for both response-compatible and response-incompatible flanker trials, e.g. goal conducive and high power appraisal in a flanker response-compatible trial). Each experiment block contained four trials of each type, and the entire experiment contained 16 trials of each type.

At the beginning of each trial, information about goal conduciveness and power appraisals was presented by different cues: positive vs. negative number of points, human vs. robot icon, background colour, and descriptive text. Goal conduciveness was indicated by the number of points that could be gained or lost on the current trial (e.g. +30 for goal conducive and -30 for goal obstructive trial; the points ranged randomly between plus or minus 25 to 35 points, respectively). Power was indicated by the icon presented in the middle of the screen. The robot icon represented low power trials where the response in the arithmetic task was made by the program. The human icon represented high power trials where participants could respond in the arithmetic task. Additionally, each trial type was colour coded using four counterbalanced colour pairings. Blue or pink signalled goal conduciveness, and lighter or darker shade of the given colour of the trial signalled power appraisal. Finally, each trial type was characterised by a short textual reminder.

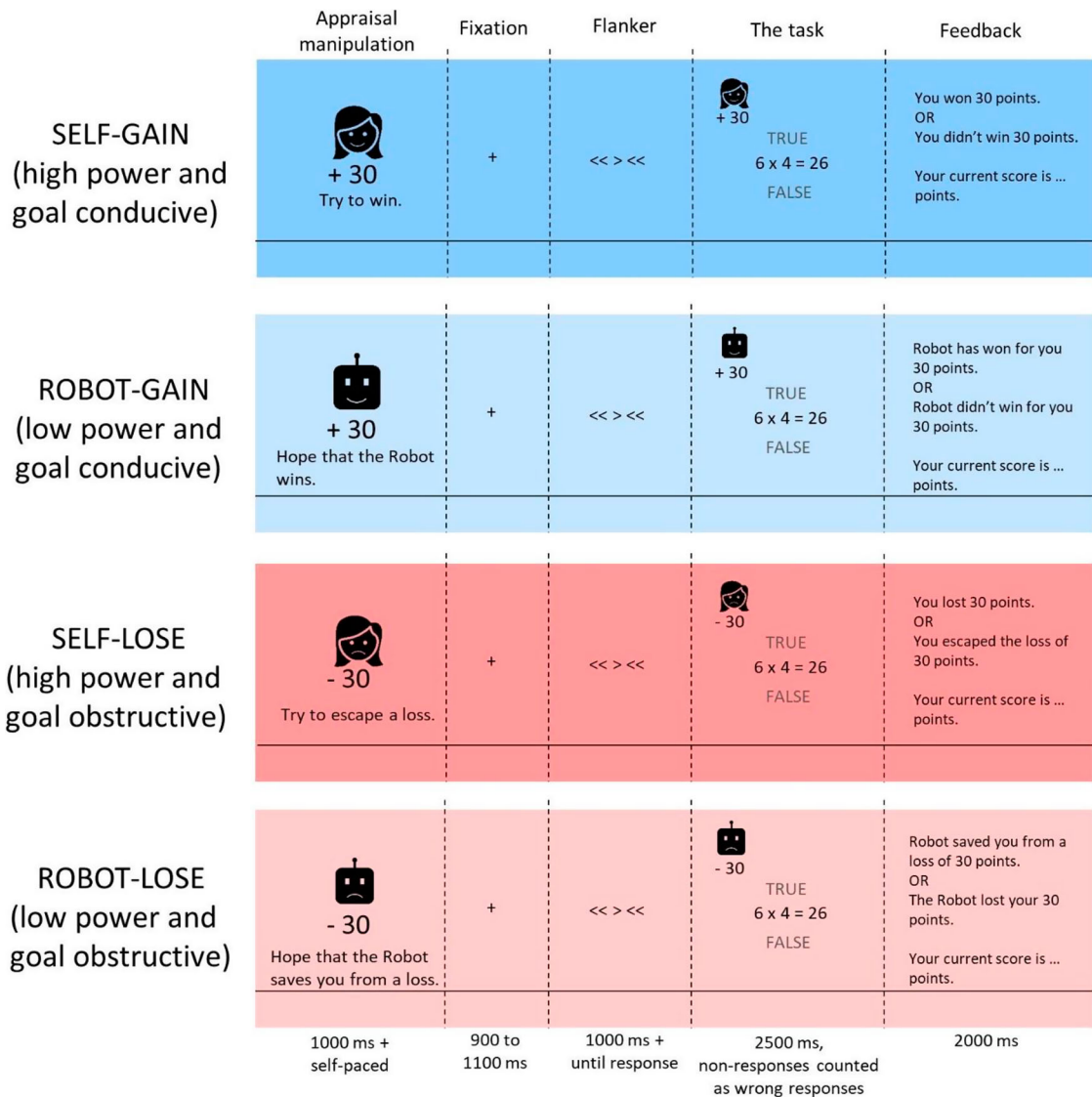


Figure 1. The trial structure of Experiment 1. The experiment had a 2×2 design: (1) goal conduciveness (gain vs. loss); (2) power appraisal (self vs. robot). Goal conduciveness determined whether participants played for a win or for avoiding a loss during the trial. Power appraisal determined whether participants were able to respond in the arithmetic task or the response was made by the robot. A correct response in the arithmetic task resulted in gaining points or not losing points, in goal conducive and goal obstructive trials, respectively. A false response resulted in not gaining or losing points, in goal conducive and goal obstructive trials, respectively. The colour coding of the background was counterbalanced between participants.

For goal conducive high power trials, the text was “try to win”, for goal obstructive high power trials “try to avoid a loss”, for goal conducive low power trials “hope that the Robot wins”, and for goal obstructive low power trials “hope that the Robot saves you from a loss”. The initial trial screen had a presentation time of 1 s, after which participants were able to press the “up” arrow key on their keyboard to continue to the next phase of the trial.

After the initial trial screen, a fixation cross was presented at the middle of the screen for 0.9 to 1.1 s. Next, the flanker stimulus was presented in the middle of the screen for up to 1 s until a response was registered. If no response was given within 1 s, the screen turned blank until the participant responded. Within the flanker stimulus, a middle arrow was the target stimulus that was flanked by four response-compatible or response-incompatible

stimuli, two arrows on both sides pointing either to the left or the right. Participants were instructed to press either the left or right arrow key according to the middle arrow as fast as possible while making as few mistakes as possible. The flanker stimuli were black, lowercase Arial bold font and were presented on the current trial type's background. The size of the flanker stimuli depended on the calibration results from the beginning of the experiment. Successful calibration resulted in a 3-cm-wide flanker stimulus array.

After the flanker task, the arithmetic task was presented along with visual reminders of the trial type for 2.5 s. The multiplication problem consisted of an equation with two single-digit multiplicands and an answer that was either correct or off by two in either direction. Participants had to indicate whether the answer was correct or false by pressing the "up" or "down" arrow key on their keyboard, respectively. Participants were able to respond to the task in the high power trials and were not able to respond in the low power trials where the response was made automatically by the program after 1 s of presentation time. The chosen response option was highlighted by the text colour change from grey to black. In the goal conducive ("potential to win") trial, making a correct decision about the equation yielded points and an incorrect decision yielded nothing. In the goal obstructive ("potential to lose") trial a correct decision yielded nothing (i.e. avoidance of loss) and an incorrect decision yielded a loss. The gains and losses accumulated throughout the task.

The multiplication problems were assigned into four difficulty categories based on an online dataset (*Testing Times*, 2013). Each category was represented equally within experimental conditions and blocks. The robot was accurate in 75% of the multiplication trials and made mistakes only for multiplication tasks in the most difficult category.

At the end of the trial, the trial outcome was presented for 2 s. In high power trials, participants saw whether they won points, did not win points, lost points or did not lose points. In the low power trials, participants saw whether the robot won points for them, did not win points for them, lost their points or did not lose their points. For example, when 30 points were in play the respective descriptions were as follows: "Robot has won 30 points for you", "Robot did not win 30 points for you", "Robot lost your 30 points", "Robot saved you from losing 30

points". Additionally, their current score was presented on the screen.

Between each block, participants were given feedback about their current standing, overall accuracy in the arithmetic task, and false response percentage in the flanker task. Based on pilot studies, a cut-off value of 240 points per block that was easily achievable for most of the participants was used to generate feedback on current standing. Specifically, participants who achieved at least 240 points per block (92.8% of the participants) were told that their score was among the top 20 and they would be included in the raffle for the 100 €-gift card. Participants scoring below the cut-off, were told that their score was currently not in the top 20 and were encouraged to focus more on the task.

After the experiment, participants completed a short questionnaire designed to measure the effectiveness of the appraisal manipulation. In the questionnaire, the four different pictures that were shown at the beginning of each trial (see [Figure 1](#)) were presented one by one ([Figure 2](#)). For each picture, participants were asked to assess: (1) relevance appraisal - "By seeing this icon, I felt that the result of this trial is relevant for me"; (2) goal conduciveness appraisal - "By seeing this icon, I felt that the result of this trial is beneficial for my score"; (3) power appraisal - "By seeing this icon, I felt that I can control the outcome of the trial"; (4) valence - "This icon was positive for me"; and (5) motivation - "This icon gave me motivation to act". The responses were given on a scale ranging from 0 (agree not at all) to 7 (completely agree).

2.1.3. Analyses

Data were analysed with R (R Core Team, 2014) and RStudio (RStudio Team, 2015) software. Linear mixed models analyses were implemented in R packages lme4 (Bates et al., 2015) and lmerTest (Kuznetsova et al., 2020). The Benjamini and Hochberg (1995) method was used to adjust the *p*-values for multiple comparisons.

For all flanker analyses, we removed 7.67% of responses that were faster or slower than three times the condition's median absolute deviation (Leys et al., 2013). For response time analysis, we also removed 2.18% of trials with an incorrect response. The response times were analysed using linear mixed-model regressions with experiment conditions as fixed factors and participants and the experiment block as random intercept effect.

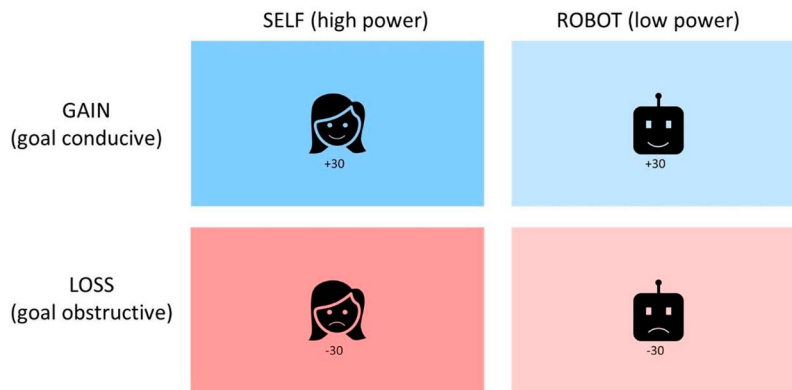


Figure 2. Example icons used in the post-experiment questionnaire. The icons were presented one at a time.

The drift-diffusion analysis used pre-processed flanker data with error trials. We used the EZ diffusion model that is effective with small number of trials ($N < 100$; Lerche et al., 2017) and low error rate (Wagenmakers et al., 2007). The EZ DDM model estimates three parameters for each experimental condition: drift rate (v), non-decision time (t), and decision separation boundary (a). Drift rate describes the speed of information accumulation for reaching the decision boundaries. Non-decision time reflects the duration of processes other than the decision process, like stimulus encoding and motor execution. Decision separation boundary indicates the information required to reach one of the decisions and start the response execution. The EZ DDM was fitted to participant's response time and response accuracy data from the eight different task conditions (2 goal conduciveness levels $\times$ 2 power levels $\times$ flanker congruence levels). To assess how each condition influenced the DDM parameters, we used linear mixed effect models with participants as random intercept effects.

To estimate the validity of our findings, we also computed Bayes factors for the main effects of each condition with Bayesian paired t -tests in JASP (JASP Team, 2021). We considered two different prior distributions (e.g. Gronau et al., 2017). First, we used default prior which does not assume clear prior knowledge about the effect (Cauchy distribution, centred at zero, with scale $1/\sqrt{2}$). Second, we used an informed prior, the Oosterwijk's prior, which represents small-to-medium effects sizes common in the field of psychology (t -distribution, centred at 0.35, with a scale of .102 and 3 df; Quintana & Williams, 2018).

2.2. Results

2.2.1. Manipulation check

We analysed participants' responses from the post-experiment questionnaire to check whether the operationalisation of goal conduciveness and power appraisal was successful. See Table S1 for descriptive statistics of all the post-experiment stimulus assessment questions.

To assess the effectiveness of the operationalisation of goal conduciveness appraisal, we analysed whether goal conduciveness ratings were predicted by the goal conduciveness appraisal (conductive vs. obstructive) and power appraisal (high vs. low). The analysis showed that the main effect of goal conduciveness appraisal was significant, with higher ratings for goal conducive conditions ($\beta = 1.42$, $SE = 0.15$, $t(597) = 9.41$, $p < .001$). The main effect of power appraisal was also significant, with high power stimuli rated more conducive than low power stimuli ($\beta = 0.91$, $SE = 0.15$, $t(597) = 6.02$, $p < .001$). This may arise from the general pleasantness of being in a position with power (Shuman et al., 2013). The interaction effect of conduciveness and power appraisals was not significant. The results suggest that the operationalisation of goal conduciveness appraisal was successful, as the goal conduciveness ratings for goal conducive conditions were higher than for goal obstructive conditions.

To assess the effectiveness of the operationalisation of power appraisal, we analysed whether power ratings were predicted by power appraisal (high vs. low) and conduciveness appraisal (conductive vs. obstructive). The analysis showed a significant main effect of power appraisal, with higher ratings for

high-power stimuli ($\beta = 3.60$, $SE = 0.15$, $t(597) = 24.77$, $p < .001$). The main effect of goal conduciveness appraisal and the interaction effect between conduciveness and power appraisals were not significant. The results suggest that the manipulation of power appraisal was as intended and distinct from goal conduciveness, as the power ratings for high power conditions were higher than for low power condition (Figure 3).

2.2.2. Appraisal effects on error rates and response times

Based on the results of the self-report data, we conclude the operationalisation of both appraisals worked well. Next, we addressed the central question whether the manipulation of goal conduciveness and power appraisals influenced breadth of attention. The descriptive statistics for each condition are presented in Table 1. See Table S2 for complete results.

First, we tested whether flanker congruence, conduciveness appraisal, and power appraisal predict the average error rate of the participants in the flanker task (Figure 4A). The analysis showed a significant main effect of power appraisal, with high power trials decreasing the overall error rate ($\beta = -3.08$, $SE = 0.77$, $t(1463) = -3.99$, $p < .001$). Additionally, the analysis showed a significant main effect of flanker congruence, with response-compatible trials decreasing the overall error rate ($\beta = -3.50$, $SE = 0.77$, $t(1463) = -4.53$, $p < .001$). The error rate analysis did not show any significant effect of appraisals on the breadth of attention as there were no interaction effects between flanker congruence and appraisal manipulations.

Next, the single trial reaction times in the flanker task were submitted to a similar 2 (flanker congruence: response-compatible vs. response-incompatible) $\times$ 2 (goal conduciveness appraisal: conducive vs. obstructive) $\times$ 2 (power appraisal: high vs. low) mixed-model analysis (Figure 4B). The analysis showed a significant main effect of power appraisal, with high power trials decreasing the overall reaction time ($\beta = -0.013$, $SE = 0.002$, $t(24830) = -7.32$, $p < .001$). The main effect of goal conduciveness appraisal was also significant, with goal conducive trial decreasing the overall reaction time ($\beta = -0.004$, $SE = 0.002$, $t(24830) = -2.36$, $p = .018$). Additionally, the main effect of flanker congruence was significant, with response-compatible trials decreasing the overall reaction time ($\beta = -0.100$, $SE = 0.002$, $t(24830) = -54.86$, $p < .001$). Similar to the error rate analysis,

the results did not show a significant interaction effect of appraisal manipulations with the breadth of attention.

To test the null hypothesis that there is no difference in the breadth of attention among the experimental conditions, we used Bayesian paired t -tests. For this aim, we calculated the flanker interference score for each participant as a difference between reaction times on response-incompatible trials and response-compatible trials. First, we analysed the main effects of goal conduciveness and power appraisals by using the default prior. Both analyses suggest that the data provide moderate and strong support for the null hypothesis for goal conduciveness ($BF_{10} = 0.12$) and power appraisal ($BF_{10} = 0.09$), respectively. Next, we analysed the same effects by using the informed prior assuming small-to-medium effect size. Results of both analyses show moderate support for the null hypothesis ($BF_{10} = 0.03$; $BF_{10} = 0.03$).

2.2.3. Drift diffusion analysis

The DDM parameters for each condition are presented in Table 1. See Table S3 for complete results. First, we found that the drift rate was significantly higher in response-compatible flanker trials compared to response-incompatible flanker trials ($\beta = 0.050$, $SE = 0.008$, $t(1456) = 6.56$, $p < .001$), and higher in high power trials compared to low power trials ($\beta = 0.046$, $SE = 0.008$, $t(1456) = 6.06$, $p < .001$).

Second, we found that non-decision time was smaller in response-compatible flanker trials ($\beta = -0.095$, $SE = 0.003$, $t(1456) = -34.24$, $p < .001$) and higher in high power trials compared to low power trials ($\beta = 0.006$, $SE = 0.003$, $t(1456) = 2.19$, $p = .03$).

Third, we found that the decision boundary was lower in high power trials, $\beta = -0.006$, $SE = 0.001$, $t(1456) = -5.30$, $p < .001$. In addition, it was influenced by a two-way interaction effect of flanker task and power appraisal, $\beta = -0.003$, $SE = 0.002$, $t(1456) = -2.07$, $p = .04$. Post-hoc tests showed that there was a difference between high power response compatible and response-incompatible flanker trials ($p < .01$), but there was no difference in the respective low power flanker trials ($p = .06$).

In conclusion, DDM showed that faster response times in high power appraisal trials were due to higher drift rate, higher non-decision time, and lower decision boundary in those trials. In addition, faster responses for response-compatible flanker trials were due to higher drift rate and lower non-decision time.

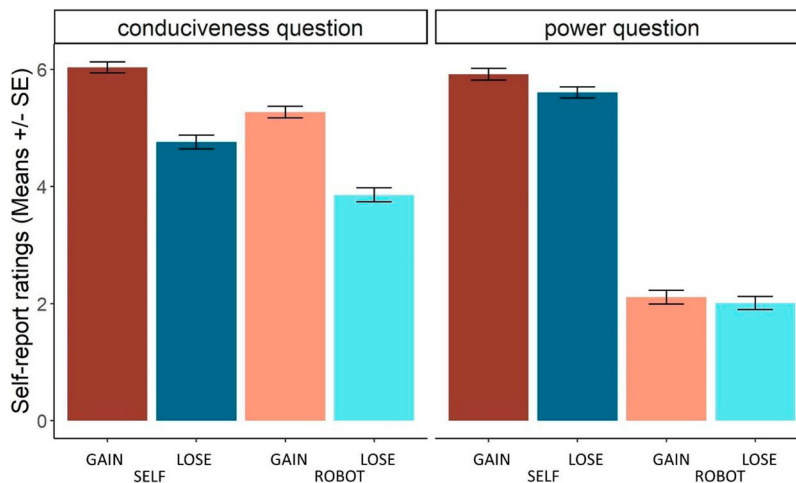


Figure 3. Experiment 1, post-experiment self-report ratings. Conduciveness question = I felt that the result of this trial is beneficial for my score; power question = I felt that I can control the outcome of the trial; self = high power trials; robot = low power trials; gain = goal conducive trials, potential to gain victory points; lose = goal obstructive trials, potential to lose victory points.

2.3. Discussion

In Experiment 1, we examined whether goal conduciveness and power appraisals influence breadth of attention operationalised in a flanker task. We did not find evidence for appraisal effect on breadth of attention. The Bayes factors indicated moderate to strong support for the null hypothesis.

One reason for this null result may be the limited sensitivity of the flanker task to capture changes in breadth of attention. Besides attentional breadth (Liu et al., 2016; Rowe et al., 2007), the flanker task has also been linked to executive control because of the inherent response conflict in the task (Posner & Rothbart, 2007). It is possible that the response conflict overshadowed attentional broadening effects (Bruyneel et al., 2013; Vanlessen et al., 2016). Thus, the absence of appraisal effects on breadth of attention should be replicated with another paradigm. This was the aim of Experiment 2.

Even as appraisals did not influence breadth of attention in Experiment 1, we found that high power appraisal improved overall task performance. Participants made significantly fewer errors and responded faster in the high compared to the low power condition. The DDM suggested that this improved performance can be attributed to a higher drift rate (reflecting faster processing of choice-relevant information), higher non-decision time (reflecting sensory and motor processes) as well as lower decision boundary (indicating a more liberal decision

strategy). The confidence in the DDM results is qualified by the relatively low error rate and small number of trials that were available in this study. Nevertheless, taken together, these findings suggest that the power appraisal may have improved task performance through numerous cognitive pathways.

3. Experiment 2

The aim of Experiment 2 was to investigate whether conduciveness and power appraisals influence breadth of attention when it is assessed using a task with lower cognitive control demands. To this end, we used the Navon task (Navon, 1977) where participants have to respond to a letter (e.g. "T") made out of smaller letters (e.g. "L"). The paradigm assumes that a wide breadth of attention facilitates the processing of the large letter (i.e. global features of the stimulus) whereas a narrow breadth of attention facilitates the processing of the smaller letters (i.e. local features of the stimulus). To remove the response conflict that was present in the flanker task, we used the undirected version of the Navon task (e.g. Gable & Harmon-Jones, 2008), where only one of the target letters ("T" or "H" in this study) could appear either on the global (large "T" made out of small "Ls") or the local level (large "L" made out of small "Ts"). In addition, to reduce the overall global dominance of the Navon letters (whereby it is easier to respond to target letters that appear on global level), we

Table 1. Descriptive statistics of Experiment 1.

Flanker task	Appraisal manipulation						DDM parameters					
	Power	Goal-conduciveness	Error rate		Reaction time		Drift rate		Non-decision time		Decision boundary	
			<i>M</i> (%)	<i>SD</i>	<i>M</i> (s)	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>
Response compatible	Self	Gain	3.86	7.82	0.49	0.07	0.47	0.07	0.41	0.05	0.07	0.01
	Self	Loss	3.35	7.58	0.49	0.08	0.47	0.08	0.41	0.05	0.07	0.01
	Robot	Gain	5.68	9.86	0.50	0.09	0.42	0.09	0.40	0.06	0.08	0.01
	Robot	Loss	5.65	9.55	0.50	0.10	0.41	0.09	0.40	0.05	0.08	0.02
Response incompatible	Self	Gain	6.82	10.41	0.58	0.09	0.40	0.10	0.50	0.06	0.08	0.02
	Self	Loss	6.07	9.92	0.59	0.09	0.40	0.10	0.50	0.06	0.08	0.01
	Robot	Gain	10.50	13.20	0.60	0.11	0.35	0.12	0.49	0.05	0.08	0.02
	Robot	Loss	9.15	12.18	0.60	0.11	0.36	0.11	0.50	0.06	0.08	0.02

M = mean, *SD* = standard deviation, DDM = drift diffusion modelling, Self = high power appraisal, Robot = low power appraisal, Gain = goal conducive, Loss = goal obstructive.

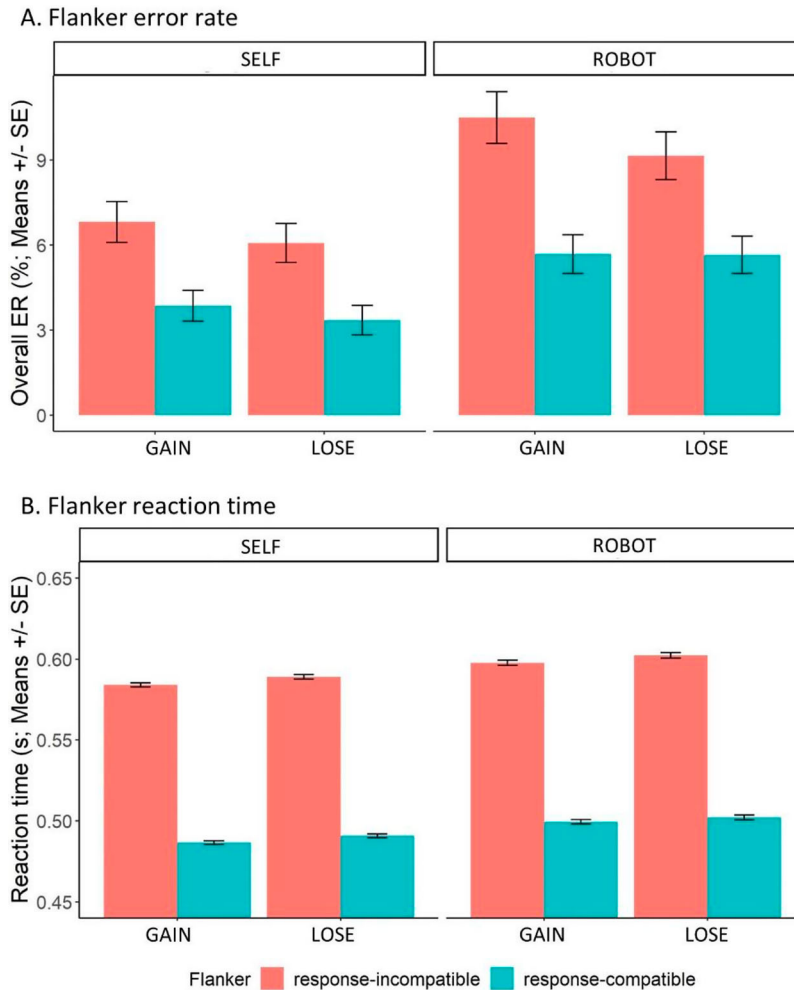


Figure 4. (A) Flanker error rate results. (B) Flanker reaction time results. Self = high power trials; robot = low power trials; gain = goal conducive trials, trials with possibility to win victory points; lose = goal obstructive trials, trials with possibility to lose points.

increased the sparsity in the Navon stimuli to balance the mean response times to global and local levels (Martin, 1979; see Method section below).

As the Bayes factors in Experiment 1 provided strong support for the null hypothesis, we expected to replicate this result of no appraisal effects on breadth of attention. In addition, we expected to replicate the overall power appraisal effect on task performance.

3.1. Method

3.1.1. Participants

Initially, 215 participants (age: $M = 24.71$, $SD = 7.25$; 64 females) completed the experiment. We excluded 19 participants from the final sample based on unreliable Navon task performance: 15 had a high error rate (over 30%) and 4 participants had less than 50% trials left in at least one condition after removing outlying response times. The final sample contained data from 196 participants (age: $M = 24.51$, $SD = 7.28$; 58 females), of whom 191 also completed the post-experiment questionnaire.

Participants were recruited among Leiden's University students through the recruitment website SONA (<https://www.sona-systems.com>) and globally through Prolific Academic (www.prolific.co). Participants recruited through SONA received course credits and participants recruited through Prolific Academic received monetary compensation (3.50 €). Similar to the first experiment, for additional motivation, one randomly chosen participant received a 100 € gift card at the end of the study. The study was approved by the Institutional Review Board of Leiden University.

3.1.2. Procedure

The procedure was identical to Experiment 1 except for the following aspects. FormR survey framework (Arslan et al., 2020) was used for the post-experiment questionnaire. Instead of the flanker task to measure breadth of attention, we used the Navon task. Additionally, we added a short blank screen (1 s) to the trial structure after the Navon task (Figure 5).

We used the undirected version of the Navon task (Gable & Harmon-Jones, 2008) where participants have to identify one of the target letters in the Navon stimulus. The target letters ("T" and "H") were half of the time on the global level and half of the time on the local level. When the target was presented on the global level, a large H or T was constructed from small L-s, F-s, E-s or U-s. When the

target was presented on the local level, the large letters L, F, E or U were constructed of H-s or T-s. To reduce the global precedence that is usually present in the traditional Navon stimuli (Navon, 1977), we constructed the Navon letters with fewer local elements (Martin, 1979). The letters were black and in upper-case Arial font and were presented on the current trial type's background. The size of the Navon stimuli depended on the calibration result from the beginning of the experiment. After a successful calibration, the global shape of the Navon stimuli was approximately 50×35 mm, and each local element 5×4 mm. At the start of the experiment, participants were asked to sit 60 cm from their monitor (reported distances: $M = 56.38$ cm, $SD = 16.00$ cm).

Participants were asked to use their arrow keys to indicate as quickly as possible which target letter was present on the screen. The pairing of letters with response keys was counterbalanced between participants. Faster responses to targets on the global level vs. the local level indicated a global focus of attention. The opposite indicated a local focus of attention.

Like Experiment 1, Experiment 2 consisted of 4 blocks and 128 trials in total. Overall, in 16 trials, 8 trial types were presented (all appraisal combinations for both global and local Navon trials, e.g. high power and goal conducive appraisal in a global Navon trial), each experimental block contained four trials of each type.

3.1.3. Analyses

We removed responses faster or slower than three times the condition's median absolute deviation (8.06%). For reaction time analysis, we also removed the error trials (5.43% of the data). Taken together, the reaction time analysis was carried out on the remaining 87.09% of data and the DDM was carried out on 92.52% (including error trials).

3.2. Results

3.2.1. Manipulation check

We analysed the data from the post-experiment questionnaire identically to Experiment 1. See Table S4 for descriptive statistics of all the post-experiment stimulus assessment questions. To assess the effectiveness of the goal conduciveness manipulation, we analysed whether stimulus conduciveness ratings were predicted by the goal conduciveness appraisal (conductive vs. obstructive) and power appraisal (high vs. low). The analysis showed a significant main effect

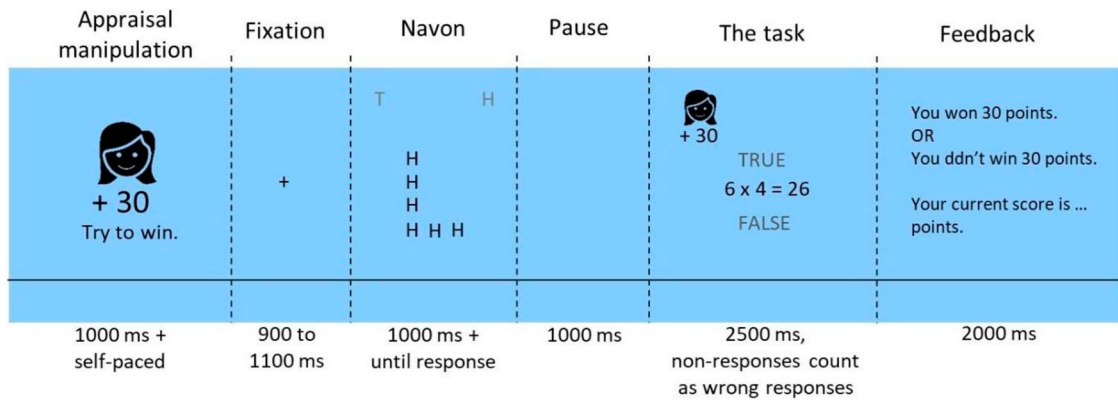


Figure 5. Experiment 2 trial structure. Goal conducive trial.

of conduciveness appraisal, with higher ratings for goal conducive conditions ($\beta = 2.28$, $SE = 0.16$, $t(570) = 14.35$, $p < .001$). The main effect of power appraisal was also significant, with higher ratings for high power stimuli ($\beta = 0.37$, $SE = 0.16$, $t(570) = 2.31$, $p = .02$). The interaction effect of goal conduciveness and power appraisal on self-report ratings was insignificant.

To assess the effectiveness of the power appraisal manipulation, we analysed whether self-report assessments of stimulus power were predicted by the goal conduciveness appraisal (conductive vs. obstructive) and power appraisal (high vs. low). The analysis showed a significant main effect of power appraisal, with higher ratings for high power stimuli ($\beta = 3.58$, $SE = 0.15$, $t(570) = 23.26$, $p < .001$). The interaction

effect between power appraisal and conduciveness appraisal was almost significant ($\beta = 0.42$, $SE = 0.22$, $t(570) = 1.95$, $p = .05$). Taken together, these findings replicate the findings from Experiment 1 confirming the suitability of the paradigm for operationalising the two appraisals (Figure 6).

3.2.2. Appraisal effects on error rates and response times

Based on the results of the self-report data, we conclude the operationalisation of both appraisals worked as intended. Next, we tested whether the manipulation of goal conduciveness and power appraisal influenced breadth of attention. The descriptive statistics for each condition are presented in Table 2. See Table S5 for complete results.

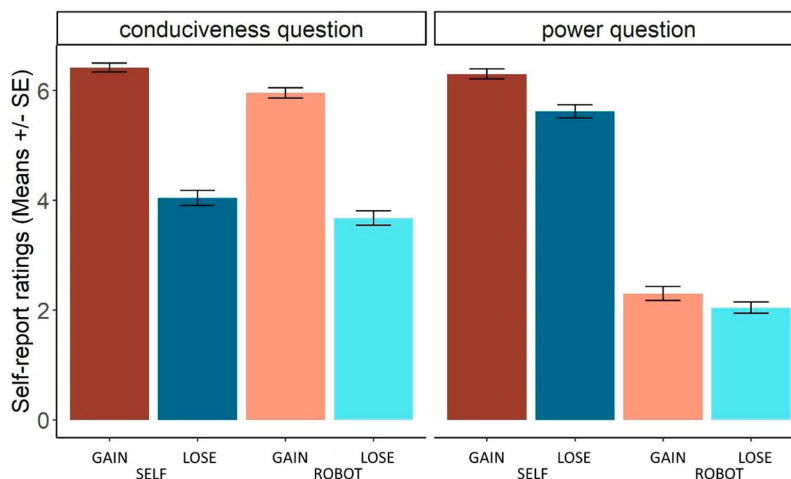


Figure 6. Experiment 2 post-experiment self-report ratings. Conduciveness question = I felt that the result of this trial is beneficial for my score; power question = I felt that I can control the outcome of the trial; self = high power trials; robot = low power trials; gain = goal conducive trials, potential to gain victory points; lose = goal obstructive trials, potential to lose victory points.

First, we analysed whether Navon level, conduciveness appraisal, and power appraisal predict the average error rate of the participants in the Navon task. The analysis showed no significant effects.

Next, single-trial reaction times were submitted to a 2 (Navon level: local, global) $\times$ 2 (goal conduciveness appraisal: conducive, obstructive) $\times$ 2 (power appraisal: high, low) mixed-model analysis. Replicating the findings of Experiment 1, the analysis showed a significant main effect of power appraisal, with decreased reaction time in high power trials ($\beta = -0.029$, $SE = 0.005$, $t(21640) = -5.60$, $p < .001$).

To assess the support for the null hypothesis, we used Bayesian paired t -tests. For this aim, we calculated the Navon global precedence (global trials – local trials) score for each participant based on their reaction times. First, we analysed the main effects of conduciveness and power appraisal by using the default prior. These analyses suggest that there was strong and moderate evidence for the null hypothesis for conduciveness ($BF_{10} = 0.09$) and power appraisals ($BF_{10} = 0.18$), respectively. Next, we analysed the same effects by using the informed prior assuming small-to-medium effect size. These analyses again suggest that there was strong and moderate evidence for the null hypothesis for conduciveness ($BF_{10} = 0.06$) and power appraisals ($BF_{10} = 0.20$), respectively (Figure 7).

3.2.3. Drift diffusion analysis

As in Experiment 1, the DDM parameters were calculated for each participant with the EZ DDM algorithm (see Method section). A linear mixed-effect model was used to assess to what extent Navon level (global vs. local), goal conduciveness appraisal (conductive vs. obstructive), and power appraisal (high vs. low) influenced drift rate, non-decision time, and decision boundary. The DDM parameters for each condition are presented in Table 2. See Table S6 for complete results. Similar to Experiment 1, we found that the drift rate was significantly higher in high power trials compared to low power trials ($\beta = 0.013$, $SE = 0.005$, $t(1365) = 2.34$, $p = .02$). In addition, drift rate was higher in local Navon trials compared to global Navon local trials ($\beta = 0.012$, $SE = 0.005$, $t(1365) = 2.16$, $p = .03$). The decision boundary lower in high power trials ($\beta = -0.005$, $SE = 0.002$, $t(1365) = -2.51$, $p = .01$). Meanwhile, non-decision time was not influenced by any of the analysed factors.

3.3. Discussion

In Experiment 2, we replicated the null results of Experiment 1. We did not find any appraisal effects on breadth of attention measured by the Navon task. In addition, the reaction time analysis replicated the power appraisal effect on the overall performance: participants responded faster in high power appraisal trials. The results from the DDM analysis were also in line with the previous experiment showing that the power appraisal effect could be attributed to a higher drift rate and lower decision boundary. The fact that there was no power appraisal effect on non-decision time suggests that the performance improvement associated with high power is related to increased efficiency of cognitive processing of the Navon letters, rather than to enhancements in simple sensory or motor processes.

4. General discussion

Aiming to examine whether goal conduciveness and power appraisals drive affective changes in breadth of attention, we carried out two web-based experiments with the different breadth of attention measures: the flanker task and the Navon task. Overall, the results showed that the two appraisal dimensions did not directly influence breadth of attention. However, high power appraisal increased overall task performance in both experiments.

Previous research has mainly attributed affective effects on breadth of attention to the experiential or the action tendency components of affect. However, the empirical support for the roles of these components is mixed. We proposed that the ambiguous findings so far may mean that the origins of affective impacts on breadth of attention lie further upstream within the dynamic emotion process, at the appraisal component. According to appraisal theory, the appraisal component drives the responses in other emotion components such as subjective feeling and action tendency. Following this account, we investigated the effects of goal conduciveness and power appraisals on attentional breadth. Based on the association between goal conduciveness appraisal and valence, we expected that goal conducive trials (potential to win points) would broaden breadth of attention while goal obstructive trials (potential to lose points) would narrow it. Based on associations between the power appraisal and motivational intensity, we expected that high power trials (possibility to influence the

Table 2. Descriptive statistics of Experiment 2.

Navon level	Appraisal manipulation						DDM parameters					
	Power	Goal-conduciveness	Error rate		Reaction time		Drift rate		Non-decision time		Decision boundary	
			<i>M</i> (%)	<i>SD</i>	<i>M</i> (s)	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>
Global	Self	Gain	11.86	12.64	0.83	0.23	0.23	0.07	0.59	0.14	0.12	0.02
	Self	Loss	12.85	13.36	0.82	0.23	0.22	0.07	0.58	0.14	0.12	0.02
	Robot	Gain	13.36	13.42	0.84	0.24	0.21	0.06	0.58	0.14	0.13	0.02
	Robot	Loss	14.35	13.70	0.85	0.25	0.21	0.07	0.58	0.15	0.13	0.03
Local	Self	Gain	11.77	12.18	0.83	0.23	0.23	0.06	0.59	0.15	0.12	0.02
	Self	Loss	11.77	11.57	0.83	0.23	0.23	0.07	0.58	0.15	0.13	0.02
	Robot	Gain	13.68	12.77	0.85	0.24	0.22	0.06	0.58	0.15	0.13	0.02
	Robot	Loss	13.68	14.08	0.85	0.24	0.22	0.07	0.58	0.14	0.13	0.02

M = mean, *SD* = standard deviation, DDM = drift diffusion modelling, Self = high power appraisal, Robot = low power appraisal, Gain = goal conducive, Loss = goal obstructive.

outcome) narrow breadth of attention, and low power trials (no possibility to influence the outcome) broaden breadth of attention.

To thoroughly test these hypotheses, we used two different tasks to measure breadth of attention and two different ways to analyse the data. The two

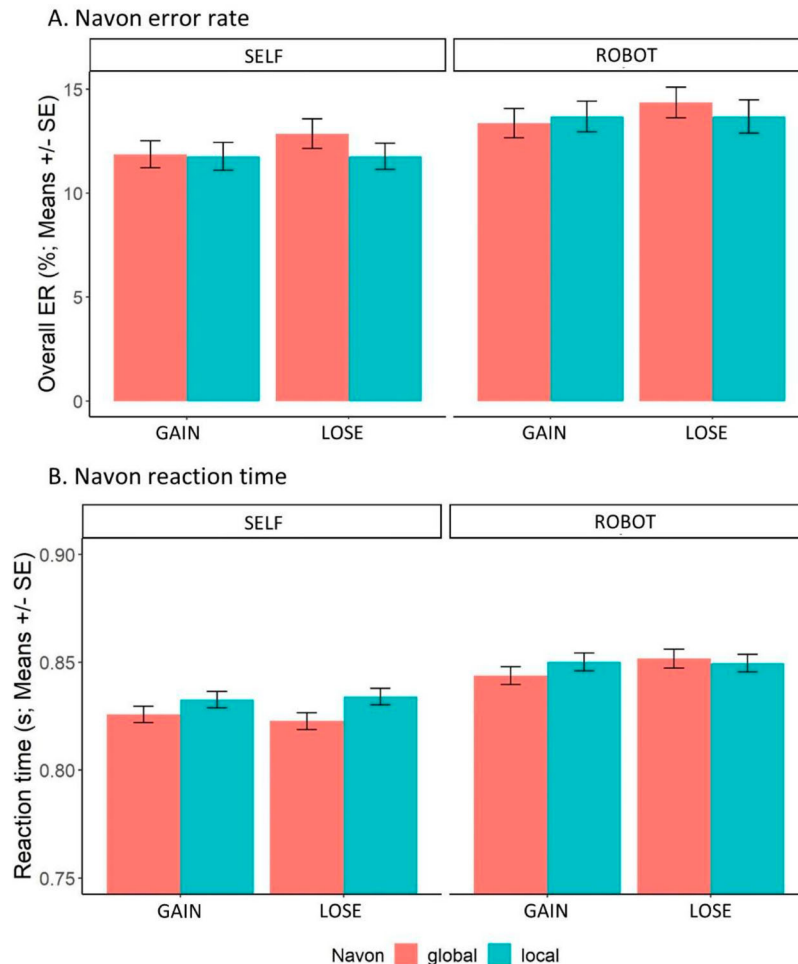


Figure 7. (A) Navon error rate results. (B) Navon reaction time results. self = high power trials; robot = low power trials; gain = goal conducive trials, trials with the possibility to win victory points; lose = goal obstructive trials, trials with the possibility to lose points.

tasks were selected to tap into different defining features of breadth of attention. The flanker task should be more sensitive to changes in the size of the focus of attention whereas the Navon task should be more sensitive to changes in the priority of global relative to local elements of a stimulus. In terms of analysis strategies, we first used an atheoretical approach by analysing the mean error rates and reaction times. In addition, we applied the DDM (Ratcliff & McKoon, 2008) that allowed us to examine the effects of the appraisal manipulation on the chosen response strategies and thus on the overall cognitive efficiency.

Across both tasks and both analyses, no evidence emerged supporting our hypotheses that appraisals directly influence breadth of attention. Thus, our findings challenge our expectations that goal conduciveness and power appraisals shift breadth of attention. Furthermore, the Bayes factors indicated moderate-to-strong support for the null hypothesis.

It is possible that alternative aspects of affect such as valence and motivational intensity better explain the present findings. However, as analysis presented in the supplementary material indicated, neither the self-report rating of stimulus valence ("This icon was positive for me.") nor of motivational intensity ("This icon gave me motivation to act.") revealed significant effects. Therefore, we conclude that these alternative accounts cannot sufficiently explain our results.

Both experiments, however, showed that high power appraisal significantly improved overall performance indicated by faster response times. DDM revealed that in high power trials the drift rate was higher and the decision boundary was lower. The former means that information for a decision was gathered more quickly, and the latter that less information was needed to make the decision. The findings are in line with motivational impacts on cognitive control (Botvinick & Braver, 2015; Pessoa & Engelmann, 2010). According to this perspective, allocation of cognitive control is based on an automatic weighing of the anticipated costs and benefits of investing high vs. low amounts of cognitive control into a task. When anticipated rewards outweigh the anticipated costs, it is more beneficial to direct resources to the task. In our experiments, in high power trials, participants were able to influence the outcome of the trial themselves and thus the favourable outcome was more achievable by their own actions than in the low power trials. High power appraisals may influence cognitive resource allocation which manifest in improved task performance.

When interpreting the present findings it is worth considering that the absence of direct appraisal effect on breadth of attention was unlikely to arise from a lack of statistical power, low validity of web-based data, unsuccessful manipulation of the appraisals, or low validity of the dependent variable measures. First, it is unlikely that our study was underpowered as we collected large samples for both experiments, and used a within-subject design which has higher statistical power than between-subject designs (Charness et al., 2012). Second, we took several steps to assure valid data were collected via the internet. In both studies, we ensured via a standard calibration procedure that the size of the stimuli presented in both tasks was similar for all participants. Also in the post-experiment questionnaire, most participants reported that they complied with the instructions. Controlling for the reported distances, statistical models did not alter any of the substantive findings (see Supplementary materials). In addition, web-based experiments, including experiments that are hosted on Pavlovia, allow precise measurement in the range of milliseconds (Bridges et al., 2020; van Steenbergen & Bocanegra, 2016). To conclude, there are solid reasons to believe that the web-based paradigms worked as intended.

Third, the likelihood that the manipulation of goal conduciveness and power appraisals failed also seems to be unlikely. Both post-experiment questionnaires confirmed our expectations that participants rated goal conducive trials more highly on the conduciveness scale than goal obstructive trials. Similarly, participants rated high power trials more highly on the power appraisal scale than low power trials. Our confidence that appraisals operationalisation worked as intended is also supported by the robust main effect of power appraisal on the performance measures in both experiments.

Even if the goal conduciveness and power appraisal operationalisations worked, it is possible that the affective manipulation more broadly failed because the relevance of the experimental conditions was insufficient to induce affective reactions. In appraisal theories, goal relevance appraisal is considered to be a necessary gateway appraisal for other appraisals to be assessed and consequently to drive affective reactions (Scherer, 2009). By analysing the post-experiment relevance question ("By seeing this icon, I felt that the result of this trial is relevant for me"), we found that all experimental conditions were on average rated higher than the midpoint of the scale,

thus indicating that participants agreed with the statement (see Supplementary materials). This finding indicates that the different experimental conditions were perceived as sufficiently relevant, and thus were likely appraised by the subsequent appraisal dimensions.

Turning to the dependent measures, while the Navon task is a widely used measure of breadth of attention, the interpretation of the flanker task results is less straightforward. In addition to being sensitive to the breadth of attention, performance in the flanker task also depends on cognitive control required for overcoming the interference generated by opposite-direction flankers (Posner & Rothbart, 2007). Thus, any observed effects on flanker interference may be attributed to some combination of breadth of attention and cognitive control. For instance, negative stimuli have been shown to reduce the flanker interference effect (Birk et al., 2011) which may reflect both an improvement in cognitive control and narrowing of breadth of attention. Furthermore, breadth of attention could be one factor that modulates the link between affective states and executive control (Cohen & Henik, 2012). Future studies can employ techniques to disentangle the cognitive control and breadth of attention effects in the flanker task by including the manipulation of the distance between target and flankers (e.g. Rowe et al., 2007) or use another attentional breadth measure, such as functional field of view task (Pringle et al., 2001) or the attentional breadth task (Grol & Raedt, 2014). In the present study, however, the question of disentangling cognitive control and breadth of attention effects is less relevant given that we did not find any expected effects on flanker task performance.

In the Navon task, we reduced the role of cognitive control by not including trials where a response conflict would be generated by the stimuli. In addition, we increased the sparsity in the Navon stimuli to successfully remove the global dominance effect that is present in traditional Navon stimuli (Navon, 1977). It is important to note that there is still debate about the global and local processing differences in the Navon task (Kimchi, 2015). For example, the global dominance effect has been ascribed to different stimulus processing stages. Some authors have highlighted early perceptual-organisational processes (e.g. Han & Humphreys, 2002). Others have emphasised the role of sensory factors, such as faster processing of low than high

spatial frequencies (e.g. Badcock et al., 1990). Yet, other authors have argued that the global and local elements are processed in parallel and the difference in response times stems from post-perceptual processes (e.g. Boer & Keuss, 1982). However, there are reasons to believe that broad vs. narrow attention can operate independently of the mechanisms that facilitate global vs. local perception. For example, Sasaki et al. (2001) showed that attending to global vs. local elements in a hierarchical stimulus produced distinct brain activity that is consistent with the zoom lens model of attention (Eriksen & James, 1986). Attending to the global elements of the hierarchical stimulus, occipital cortex activity was lower in amplitude and topographically more scattered than the activity produced by attending to the local elements of the stimulus. Similar neuroimaging results have been shown with other measures of attentional breadth (Müller et al., 2003).

The present findings raise several important hypotheses for future research. First, although the (motivational) relevance appraisal indicated that the different conditions were assessed to be relevant, it could be argued that the experiments did not induce affective states strong enough to elicit changes in breadth of attention. Thus, future studies should test whether a certain critical level of affective intensity must be achieved to elicit changes in breadth of attention. For example, the design could be implemented in a laboratory setting using more immediate rewards such as food or punishments such as electric shocks.

Finally, the affective states in our experiments were quite short-lived. In both experiments, the minimal time from the start of the appraisal manipulation to the dependent variable measure was around 2 s. By contrast, many previous studies used longer-lasting affective manipulations, such as viewing emotional pictures for several seconds, watching short video clips, or recalling emotional memories (Mauss & Robinson, 2009). For example, one study showed that viewing emotional pictures for 5 s influenced the early sensory processing of flanker stimuli (Moriya & Nittono, 2011). Possibly, affective states need to last a certain amount of time for their effects on breadth of attention to unfold. Future studies could investigate the effects of different durations of affective stimuli on attentional breadth. For instance, the temporal distance between affect manipulation and the measure of breadth of attention could be systematically manipulated.

Future studies could also be designed to overcome the limited validity of manipulation check measures. Instead of assessing the success of appraisal manipulations with a post-experiment questionnaire, future studies could deploy more objective and immediate measures of the appraisal process such as the electroencephalogram (van Peer et al., 2014).

5. Conclusion

In conclusion, using two large web-based studies with the different breadth of attention measures, we did not find goal conduciveness or power appraisal effects on breadth of attention. However, we reliably showed that high power appraisal improved overall task performance suggesting that the power appraisal may direct cognitive resource allocation which results in optimised response management. The present study also indicates that the intensity and duration of affective states may need to be considered more systematically in the ongoing search for the true origins of affective effects on breadth of attention.

Disclosure statement

No potential conflict of interest was reported by the author(s).

Funding

This research was supported by the Estonian Research Council grant PSG525.

References

- Arslan, R. C., Walther, M. P., & Tata, C. S. (2020). Formr: A study framework allowing for automated feedback generation and complex longitudinal experience-sampling studies using R. *Behavior Research Methods*, 52(1), 376–387. <https://doi.org/10.3758/s13428-019-01236-y>
- Badcock, J. C., Whitworth, F. A., Badcock, D. R., & Lovegrove, W. J. (1990). Low-frequency filtering and the processing of local-global stimuli. *Perception*, 19(5), 617–629. <https://doi.org/10.1068/p190617>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1–48. <https://doi.org/10.18637/jss.v067.i01>
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: A practical and powerful approach to multiple testing. *Journal of the Royal Statistical Society. Series B (Methodological)*, 57(1), 289–300. <https://doi.org/10.1111/j.2517-6161.1995.tb02031.x>
- Birk, J. L., Dennis, T. A., Shin, L. M., & Urry, H. L. (2011). Threat facilitates subsequent executive control during anxious mood. *Emotion (Washington, D.C.)*, 11(6), 1291–1304. <https://doi.org/10.1037/a0026152>
- Boer, L. C., & Keuss, P. J. (1982). Global precedence as a postperceptual effect: An analysis of speed–accuracy tradeoff functions. *Perception & Psychophysics*, 31(4), 358–366. <https://doi.org/10.3758/BF03202660>
- Botvinick, M., & Braver, T. (2015). Motivation and cognitive control: From behavior to neural mechanism. *Annual Review of Psychology*, 66(1), 83–113. <https://doi.org/10.1146/annurev-psych-010814-015044>
- Bridges, D., Pitiot, A., MacAskill, M. R., & Peirce, J. W. (2020). The timing mega-study: Comparing a range of experiment generators, both lab-based and online. *PeerJ*, (8), e9414. <https://doi.org/10.7717/peerj.9414>
- Bruyneel, L., van Steenbergen, H., Hommel, B., Band, G. P. H., De Raedt, R., & Koster, E. H. W. (2013). Happy but still focused: Failures to find evidence for a mood-induced widening of visual attention. *Psychological Research*, 77(3), 320–332. <https://doi.org/10.1007/s00426-012-0432-1>
- Carver, C. (2003). Pleasure as a sign you can attend to something else: Placing positive feelings within a general model of affect. *Cognition and Emotion*, 17(2), 241–261. <https://doi.org/10.1080/026999303022294>
- Charness, G., Gneezy, U., & Kuhn, M. A. (2012). Experimental methods: Between-subject and within-subject design. *Journal of Economic Behavior & Organization*, 81(1), 1–8. <https://doi.org/10.1016/j.jebo.2011.08.009>
- Clore, G. L., & Huntsinger, J. R. (2007). How emotions inform judgment and regulate thought. *Trends in Cognitive Sciences*, 11(9), 393–399. <https://doi.org/10.1016/j.tics.2007.08.005>
- Cohen, N., & Henik, A. (2012). Do irrelevant emotional stimuli impair or improve executive control? *Frontiers in Integrative Neuroscience*, 6, 33. <https://doi.org/10.3389/fnint.2012.00033>
- Dale, G., & Arnell, K. M. (2013). Investigating the stability of and relationships among global/local processing measures. *Attention, Perception, & Psychophysics*, 75(3), 394–406. <https://doi.org/10.3758/s13414-012-0416-7>
- Easterbrook, J. A. (1959). The effect of emotion on cue utilization and the organization of behavior. *Psychological Review*, 66(3), 183–201. <https://doi.org/10.1037/h0047707>
- Eriksen, B. A., & Eriksen, C. W. (1974). Effects of noise letters upon the identification of a target letter in a nonsearch task. *Perception & Psychophysics*, 16(1), 143–149. <https://doi.org/10.3758/BF03203267>
- Eriksen, C. W., & James, J. D. S. (1986). Visual attention within and around the field of focal attention: A zoom lens model. *Perception & Psychophysics*, 40(4), 225–240. <https://doi.org/10.3758/BF03211502>
- Fredrickson, B. L. (2004). The broaden-and-build theory of positive emotions. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 359(1449), 1367–1377. <https://doi.org/10.1098/rstb.2004.1512>
- Fredrickson, B. L., & Branigan, C. (2005). Positive emotions broaden the scope of attention and thought-action repertoires. *Cognition & Emotion*, 19(3), 313–332. <https://doi.org/10.1080/02699930441000238>
- Friedman, R. S., & Förster, J. (2010). Implicit affective cues and attentional tuning: An integrative review. *Psychological Bulletin*, 136(5), 875–893. <https://doi.org/10.1037/a0020495>
- Friedman, R. S., & Forster, J. (2011). Limitations of the motivational intensity model of attentional tuning: Reply to Harmon-Jones, Gable, and Price (2011). *Psychological Bulletin*, 137(3), 513–516. <https://doi.org/10.1037/a0023088>

- Frijda, N. H., Kuipers, P., & ter Schure, E. (1989). Relations among emotion, appraisal, and emotional action readiness. *Journal of Personality and Social Psychology*, 57(2), 212–228. <https://doi.org/10.1037/0022-3514.57.2.212>
- Gable, P., & Harmon-Jones, E. (2010a). The Blues Broaden, but the Nasty Narrows: Attentional Consequences of Negative Affects Low and High in Motivational Intensity. *Psychological Science*, 21(2), 211–215. <https://doi.org/10.1177/0956797609359622>
- Gable, P., & Harmon-Jones, E. (2010b). The motivational dimensional model of affect: Implications for breadth of attention, memory, and cognitive categorisation. *Cognition & Emotion*, 24(2), 322–337. <https://doi.org/10.1080/02699930903378305>
- Gable, P., & Harmon-Jones, E. (2013). Does arousal per se account for the influence of appetitive stimuli on attentional scope and the late positive potential? *Psychophysiology*, 50(4), 344–350. <https://doi.org/10.1111/psyp.12023>
- Gable, P. A., & Harmon-Jones, E. (2008). Approach-motivated positive affect reduces breadth of attention. *Psychological Science*, 19(5), 476–482. <https://doi.org/10.1111/j.1467-9280.2008.02112.x>
- Gentsch, K., Grandjean, D., & Scherer, K. R. (2013). Temporal dynamics of event-related potentials related to goal conduciveness and power appraisals. *Psychophysiology*, 50(10), 1010–1022. <https://doi.org/10.1111/psyp.12079>
- Grol, M., & Raedt, R. D. (2014). Effects of positive mood on attentional breadth for emotional stimuli. *Frontiers in Psychology*, 5, Article 1277. <https://doi.org/10.3389/fpsyg.2014.01277>
- Gronau, Q. F., Van Erp, S., Heck, D. W., Cesario, J., Jonas, K. J., & Wagenmakers, E.-J. (2017). A Bayesian model-averaged meta-analysis of the power pose effect with informed and default priors: The case of felt power. *Comprehensive Results in Social Psychology*, 2(1), 123–138. <https://doi.org/10.1080/23743603.2017.1326760>
- Gross, J. J., Uusberg, H., & Uusberg, A. (2019). Mental illness and well-being: An affect regulation perspective. *World Psychiatry*, 18(2), 130–139. <https://doi.org/10.1002/wps.20618>
- Han, S., & Humphreys, G. W. (2002). Segmentation and selection contribute to local processing in hierarchical analysis. *The Quarterly Journal of Experimental Psychology Section A*, 55(1), 5–21. <https://doi.org/10.1080/02724980143000127>
- Huntsinger, J. R. (2013). Does Emotion Directly Tune the Scope of Attention? *Current Directions in Psychological Science*, 22(4), 265–270. <https://doi.org/10.1177/0963721413480364>
- JASP Team. (2021). *JASP (Version 0.16)*. <https://jasp-stats.org/>
- Kaplan, R. L., Van Damme, I., & Levine, L. J. (2012). Motivation Matters: Differing Effects of Pre-Goal and Post-Goal Emotions on Attention and Memory. *Frontiers in Psychology*, 3, Article 404. <https://doi.org/10.3389/fpsyg.2012.00404>
- Kimchi, R. (2015). The perception of hierarchical structure. In J. Wagemans (Ed.), *The Oxford Handbook of perceptual organization* (pp. 129–149). Oxford University Press.
- Kolnes, M., Naar, R., Allik, J., & Uusberg, A. (2019). Does goal congruence dilate the pupil over and above goal relevance? *Neuropsychologia*, 134, 107217. <https://doi.org/10.1016/j.neuropsychologia.2019.107217>
- Kuznetsova, A., Brockhoff, P. B., Christensen, R. H. B., & Jensen, S. P. (2020). *lmerTest: Tests in Linear Mixed Effects Models* (3.1-3) [Computer software]. <https://CRAN.R-project.org/package=lmerTest>
- Lacey, M. F., Wilhelm, R. A., & Gable, P. A. (2021). What is it about positive affect that alters attentional scope? *Current Opinion in Behavioral Sciences*, 39, 185–189. <https://doi.org/10.1016/j.cobeha.2021.03.028>
- Lerche, V., Voss, A., & Nagler, M. (2017). How many trials are required for parameter estimation in diffusion modeling? A comparison of different optimization criteria. *Behavior Research Methods*, 49(2), 513–537. <https://doi.org/10.3758/s13428-016-0740-2>
- Leys, C., Ley, C., Klein, O., Bernard, P., & Licata, L. (2013). Detecting outliers: Do not use standard deviation around the mean, use absolute deviation around the median. *Journal of Experimental Social Psychology*, 49(4), 764–766. <https://doi.org/10.1016/j.jesp.2013.03.013>
- Liu, F., Ding, J., & Zhang, Q. (2016). Positive affect and selective attention: Approach-motivation intensity influences the early and late attention processing stages. *Acta Psychologica Sinica*, 48(7), 794–803. <https://doi.org/10.3724/SP.J.1041.2016.00794>
- Martin, M. (1979). Local and global processing: The role of sparsity. *Memory & Cognition*, 7(6), 476–484. <https://doi.org/10.3758/BF03198264>
- Mather, M., & Sutherland, M. R. (2011). Arousal-biased competition in perception and memory. *Perspectives on Psychological Science: A Journal of the Association for Psychological Science*, 6(2), 114–133. <https://doi.org/10.1177/1745691611400234>
- Mauss, I. B., & Robinson, M. D. (2009). Measures of emotion: A review. *Cognition & Emotion*, 23(2), 209–237. <https://doi.org/10.1080/02699930802204677>
- Moors, A., Ellsworth, P. C., Scherer, K. R., & Frijda, N. H. (2013). Appraisal theories of emotion: State of the art and future development. *Emotion Review*, 5(2), 119–124. <https://doi.org/10.1177/1754073912468165>
- Moriya, H., & Nittono, H. (2011). Effect of mood states on the breadth of spatial attentional focus: An event-related potential study. *Neuropsychologia*, 49(5), 1162–1170. <https://doi.org/10.1016/j.neuropsychologia.2011.02.036>
- Müller, N. G., Bartelt, O. A., Donner, T. H., Villringer, A., & Brandt, S. A. (2003). A physiological correlate of the 'Zoom Lens' of visual attention. *The Journal of Neuroscience: The Official Journal of the Society for Neuroscience*, 23(9), 3561–3565. <https://doi.org/10.1523/JNEUROSCI.23-09-03561.2003>
- Navon, D. (1977). Forest before trees: The precedence of global features in visual perception. *Cognitive Psychology*, 9(3), 353–383. [https://doi.org/10.1016/0010-0285\(77\)90012-3](https://doi.org/10.1016/0010-0285(77)90012-3)
- Peirce, J., Gray, J. R., Simpson, S., MacAskill, M., Höchenberger, R., Sogo, H., Kastman, E., & Lindeløv, J. K. (2019). PsychoPy2: Experiments in behavior made easy. *Behavior Research Methods*, 51(1), 195–203. <https://doi.org/10.3758/s13428-018-01193-y>
- Pessoa, L., & Engelmann, J. B. (2010). Embedding reward signals into perception and cognition. *Frontiers in Neuroscience*, 4(17), 1–8. <https://doi.org/10.3389/fnins.2010.00017>
- Posner, M. I., & Rothbart, M. K. (2007). Research on attention networks as a model for the integration of psychological science. *Annual Review of Psychology*, 58(1), 1–23. <https://doi.org/10.1146/annurev.psych.58.110405.085516>
- Pringle, H. L., Irwin, D. E., Kramer, A. F., & Atchley, P. (2001). The role of attentional breadth in perceptual change detection.

- Psychonomic Bulletin & Review*, 8(1), 89–95. <https://doi.org/10.3758/BF03196143>
- Quintana, D. S., & Williams, D. R. (2018). Bayesian alternatives for common null-hypothesis significance tests in psychiatry: A non-technical guide using JASP. *BMC Psychiatry*, 18(1), 178. <https://doi.org/10.1186/s12888-018-1761-4>
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59–108. <https://doi.org/10.1037/0033-295X.85.2.59>
- Ratcliff, R., & McKoon, G. (2008). The diffusion decision model: Theory and data for two-choice decision tasks. *Neural Computation*, 20(4), 873–922. <https://doi.org/10.1162/neco.2008.12-06-420>
- R Core Team. (2014). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <http://www.R-project.org/>.
- Rowe, G., Hirsh, J. B., & Anderson, A. K. (2007). Positive affect increases the breadth of attentional selection. *Proceedings of the National Academy of Sciences of the United States of America*, 104(1), 383–388. <https://doi.org/10.1073/pnas.0605198104>
- RStudio Team. (2015). *RStudio: Integrated Development for R*. RStudio. <http://www.rstudio.com/>.
- Sasaki, Y., Hadjikhani, N., Fischl, B., Liu, A. K., Marret, S., Dale, A. M., & Tootell, R. B. H. (2001). Local and global attention are mapped retinotopically in human occipital cortex. *Proceedings of the National Academy of Sciences*, 98(4), 2077–2082. <https://doi.org/10.1073/pnas.98.4.2077>
- Scherer, K. R. (2009). Emotions are emergent processes: They require a dynamic computational architecture. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 364(1535), 3459–3474. <https://doi.org/10.1098/rstb.2009.0141>
- Schimmack, U., & Derryberry, D. (2005). Attentional interference effects of emotional pictures: Threat, negativity, or arousal? *Emotion*, 5(1), 55–66. <https://doi.org/10.1037/1528-3542.5.1.55>
- Shuman, V., Sander, D., & Scherer, K. R. (2013). Levels of Valence. *Frontiers in Psychology*, 4, Article 216. <https://doi.org/10.3389/fpsyg.2013.00261>
- Siemer, M., Mauss, I., & Gross, J. (2007). Same situation–different emotions: How appraisals shape our emotions. *Emotion*. <https://doi.org/10.1037/1528-3542.7.3.592>
- Steenbergen, H., Band, G., & Hommel, B. (2011). Threat but not arousal narrows attention: Evidence from pupil dilation and saccade control. *Frontiers in Psychology*, 2, Article 281. <https://doi.org/10.3389/fpsyg.2011.00281>
- Testing times: Which times tables do kids find the hardest? (2013, May 31). *The Guardian*. <http://www.theguardian.com/news/datablog/2013/may/31/times-tables-hardest-easiest-children>.
- Vanlessen, N., De Raedt, R., Koster, E. H. W., & Pourtois, G. (2016). Happy heart, smiling eyes: A systematic review of positive mood effects on broadening of visuospatial attention. *Neuroscience and Biobehavioral Reviews*, 68, 816–837. <https://doi.org/10.1016/j.neubiorev.2016.07.001>
- van Peer, J., Grandjean, D., & Scherer, K. (2014). Sequential unfolding of appraisals: EEG evidence for the interaction of novelty and pleasantness. *EMOTION*, 14(1), 51–63. <https://doi.org/10.1037/a0034566>
- van Steenbergen, H. (2015). Affective modulation of cognitive control: A biobehavioral perspective. In G. H. E. Gendolla, M. Tops, & S. L. Koole (Eds.), *Handbook of biobehavioral approaches to self-regulation* (pp. 89–107). Springer. https://doi.org/10.1007/978-1-4939-1236-0_7
- van Steenbergen, H., & Bocanegra, B. R. (2016). Promises and pitfalls of web-based experimentation in the advance of replicable psychological science: A reply to Plant (2015). *Behavior Research Methods*, 48(4), 1713–1717. <https://doi.org/10.3758/s13428-015-0677-x>
- von Hecker, U., & Meiser, T. (2005). Defocused attention in depressed mood: evidence from source monitoring. *Emotion*, 5(4), 456–463. <https://doi.org/10.1037/1528-3542.5.4.456>
- Vuilleumier, P. (2015). Affective and motivational control of vision. *Current Opinion in Neurology*, 28(1), 29–35. <https://doi.org/10.1097/WCO.0000000000000159>
- Wagenmakers, E.-J., Van Der Maas, H. L. J., & Grasman, R. P. P. (2007). An EZ-diffusion model for response time and accuracy. *Psychonomic Bulletin & Review*, 14(1), 3–22. <https://doi.org/10.3758/BF03194023>