

EUR Research Information Portal

A multivariate brain signature for reward

Published in:
NeuroImage

Publication status and date:
Published: 01/05/2023

DOI (link to publisher):
[10.1016/j.neuroimage.2023.119990](https://doi.org/10.1016/j.neuroimage.2023.119990)

Document Version
Publisher's PDF, also known as Version of record

Document License/Available under:
CC BY-NC-ND

Citation for the published version (APA):
Speer, S. P. H., Keyzers, C., Barrios, J. C., Teurlings, C. J. S., Smidts, A., Boksem, M. A. S., Wager, T. D., & Gazzola, V. (2023). A multivariate brain signature for reward. *NeuroImage*, 271, Article 119990.
<https://doi.org/10.1016/j.neuroimage.2023.119990>
[Link to publication on the EUR Research Information Portal](#)

Terms and Conditions of Use

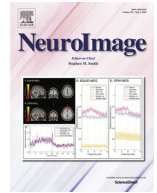
Except as permitted by the applicable copyright law, you may not reproduce or make this material available to any third party without the prior written permission from the copyright holder(s). Copyright law allows the following uses of this material without prior permission:

- you may download, save and print a copy of this material for your personal use only;
- you may share the EUR portal link to this material.

In case the material is published with an open access license (e.g. a Creative Commons (CC) license), other uses may be allowed. Please check the terms and conditions of the specific license.

Take-down policy

If you believe that this material infringes your copyright and/or any other intellectual property rights, you may request its removal by contacting us at the following email address: openaccess.library@eur.nl. Please provide us with all the relevant information, including the reasons why you believe any of your rights have been infringed. In case of a legitimate complaint, we will make the material inaccessible and/or remove it from the website.



A multivariate brain signature for reward

Sebastian P.H. Speer^{a,b}, Christian Keysers^{a,c}, Judit Campdepadrós Barrios^a, Cas J.S. Teurlings^a, Ale Smidts^d, Maarten A.S. Boksem^d, Tor D. Wager^e, Valeria Gazzola^{a,*}

^a Social Brain Lab, Netherlands Institute for Neuroscience, Amsterdam, The Netherlands

^b Princeton Neuroscience Institute, Princeton University, Princeton, NJ 08544, USA

^c Brain and Cognition, Department of Psychology, University of Amsterdam, The Netherlands

^d Rotterdam School of Management, Erasmus University, 3062 PA Rotterdam, The Netherlands

^e Department of Psychological and Brain Sciences, Dartmouth College, Hanover, NH 03755, USA

ARTICLE INFO

Keywords:

Reward
Loss
Fmri
Neural signature
Decoding
Machine learning

ABSTRACT

The processing of reinforcers and punishers is crucial to adapt to an ever changing environment and its dysregulation is prevalent in mental health and substance use disorders. While many human brain measures related to reward have been based on activity in individual brain regions, recent studies indicate that many affective and motivational processes are encoded in distributed systems that span multiple regions. Consequently, decoding these processes using individual regions yields small effect sizes and limited reliability, whereas predictive models based on distributed patterns yield larger effect sizes and excellent reliability. To create such a predictive model for the processes of rewards and losses, termed the Brain Reward Signature (BRS), we trained a model to predict the signed magnitude of monetary rewards on the Monetary Incentive Delay task (MID; $N = 39$) and achieved a highly significant decoding performance (92% for decoding rewards versus losses). We subsequently demonstrate the generalizability of our signature on another version of the MID in a different sample (92% decoding accuracy; $N = 12$) and on a gambling task from a large sample (73% decoding accuracy, $N = 1084$). We further provided preliminary data to characterize the specificity of the signature by illustrating that the signature map generates estimates that significantly differ between rewarding and negative feedback (92% decoding accuracy) but do not differ for conditions that differ in disgust rather than reward in a novel Disgust-Delay Task ($N = 39$). Finally, we show that passively viewing positive and negatively valenced facial expressions loads positively on our signature, in line with previous studies on morbid curiosity. We thus created a BRS that can accurately predict brain responses to rewards and losses in active decision making tasks, and that possibly relates to information seeking in passive observational tasks.

1. Introduction

The processing of reinforcers, such as rewards, and punishers, such as financial losses, is central to guiding our actions towards positively valenced outcomes and away from negatively valenced ones (Lutz and Widmer, 2014). Numerous functional Magnetic Resonance Imaging (fMRI) studies have investigated the neural correlates of reward processing and several meta-analyses have synthesized the findings of these studies (Bartra et al., 2013, Clithero and Rangel, 2014, Diekhof et al., 2012, Liu et al., 2011). They generally converge on two main insights: First, receiving a reward, or a loss, evokes activity in the nucleus accumbens and surrounding ventral striatum that is hypothesized to represent a positive, or negative, prediction error signal, respectively, defined as the difference between the actual outcome and the one that was expected (Diekhof et al., 2012, Galtress et al., 2012, Haber and

Knutson, 2010) (O'Doherty et al., 2004)). This signal is essential for learning as it increases the likelihood of behavior leading to better than expected outcomes (McClure et al., 2004, Schultz and Dickinson, 2000) (Yacubian, 2006)) and reduces that of behavior leading to worse than expected outcomes. Second, obtaining secondary reinforcers such as money (but also primary reinforcers such as food & nonfood consumables etc. see (Chib et al., 2009)), recruits the ventro-medial prefrontal cortex (vmPFC) (Kringelbach, 2004, Sescousse et al., 2013), the activity of which is thought to represent the subjective value of a received good (Bartra et al., 2013, Diekhof et al., 2012, Haber and Knutson, 2010, Levy and Glimcher, 2012, Peters and Büchel, 2010) and is also involved in integrating goal information and conceptual information into this value signal (Hare et al., 2008, Plassmann et al., 2007). Using MVPA, (McNamee et al., 2013) found that spatially distributed patterns in the dorsal part of the vmPFC encodes goal-value informa-

* Corresponding author.

E-mail address: cv.gazzola@nin.knaw.nl (V. Gazzola).

<https://doi.org/10.1016/j.neuroimage.2023.119990>.

Received 15 July 2022; Received in revised form 20 February 2023; Accepted 25 February 2023

Available online 5 March 2023.

1053-8119/© 2023 Published by Elsevier Inc. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>)

tion that is independent of stimulus category, whereas the more ventral part of the vmPFC encodes unique category dependent value signals in spatially distinct areas.

Most of these studies have so far used an univariate approach that aims at identifying the locations in the brain recruited while participants process rewards. In some cases, however, the aim is not to map a circuit involved in reward, but to perform reverse inference by asking whether reward processing is involved in a given task X, based on the pattern of brain activity measured at a particular moment in that task (Poldrack, 2006). It has been shown that finding activity in a particular region of the brain is a poor indicator of the recruitment of a particular mental process, because most locations are recruited while engaging a number of mental processes (Poldrack, 2006, Wager et al., 2016). In contrast, a pattern of activity across many voxels or scalp electrodes, that can include reductions and increases of BOLD and EEG signal, has been shown to be associated with a particular mental process with higher sensitivity and specificity, and therefore to provide scientists with a helpful tool to evaluate how strongly a specific mental process is recruited in a given task (Värbu et al., 2022, Wager et al., 2013, Yarkoni et al., 2011). The ability to decode the degree to which someone is receiving a reward or a loss has yet to be developed. The advantages of such a multivariate brain model are that it leads to larger effect sizes in brain-outcome association compared to more traditional local region-based approaches; makes quantitative predictions about outcomes that can be empirically falsified and can be tested and validated across studies and labs, which promotes reproducibility (for a review on brain signatures see (Kragel et al., 2018)).

Here we therefore aim to develop such a multivariate brain model for reward processing - the brain reward signature (*BRS*) - that would use distributed BOLD-based information within and across brain regions to make population-level, between-subject predictions about the strength of engagement of reinforcement/punishment processing. These predictions should ideally generalize accurately across contexts, and be able to distinguish reinforcement or punishment processing from other categories of related mental processes, such as (emotional) salience (Kragel et al., 2018). So far, few signatures for reward-related processes are available (Grosenick et al., 2013) and to our knowledge none of these have been validated on independent samples. A recent large scale challenge to predict Autism Spectrum Disorder diagnoses from fMRI (>146 team & fMRI from > 2000 individuals) highlighted the importance of validating predictive models in independent datasets because model development on a given dataset faces the risk of overfitting. Specifically, techniques such as cross-validation to measure predictive performance are not completely robust to systematic exploration of analytic choices, because the models may overfit on noise that is specific to the data set the models are trained on. Consequently, our study thus further contributes by validating the *BRS* in three independent samples.

In this study we use a predictive modeling approach (Kragel et al., 2018) that has been successfully employed to explore the neural representation of various affective processes, including the degree of physical pain (Wager et al., 2013), vicarious pain (Krishnan et al., 2016), social rejection (Woo et al., 2014), unpleasant pictures (Chang et al., 2015), basic emotions ((Kragel et al., 2016, Kragel and LaBar, 2015, Lindquist and Barrett, 2012, Wager et al., 2015), empathy (Ashar et al., 2017), guilt (Yu et al., 2020), and also faces and object categories (Haxby et al., 2001), intentions (Haynes et al., 2007, Soon et al., 2013), semantics (Huth et al., 2012, Huth et al., 2016) and clinical conditions (Arbabshirani et al., 2017, Woo et al., 2017). Our primary goal is to create a *signed relative BRS*. Specifically, the objective is to create a signature that generates more positive values for conditions associated with higher rewards, and more negative values for conditions associated with higher losses. Additionally, the signature should be specific: it should not generate high pattern responses in datasets in which reinforcement or punishment processing should be absent, but other positive or negative emotions were evoked, such as disgust or guilt. Third, it should

generalize across studies, samples and contexts where the same neurocognitive processes are engaged (i.e., be generalizable).

Based on our aim to generate a signed relative signature, we trained and tested a LASSO-PCR model (least absolute shrinkage and selection operator-regularized principal components regression; Wager et al., 2011, 2013) to predict the signed magnitude of reward received in the Monetary-Incentive-Delay task (MID, $N = 39$; see Methods) to establish the *BRS* and test its performance as quantified based on the correlation coefficient between the actual reward value and the pattern response from the neural signature. The pattern response is defined as the dot product between the *BRS* and the parameter estimates from a given condition and task plus the intercept. The MID was used because it is the most consistently used task to investigate the neural correlates of reward processing in humans (more than 200 MRI studies until now (Oldham et al., 2018) and has been designed on the basis of findings that reward anticipation engages dopaminergic neurons in the ventral tegmental area (VTA; (Knutson et al., 2000)). One strength of the MID is that it allows modeling a simple decision, which reduces the cognitive confounds that are associated with more complex decision making (Balodis and Potenza, 2015, Knutson and Greer, 2008, Lutz and Widmer, 2014), reliably. Further, the MID robustly engages the striatum, which is crucial in reward and reinforcer processing (Haber and Knutson, 2010). To further probe the performance but also the generalizability, we then applied the *BRS* to a different version of the MID (with 5 instead of three levels of reward; $N = 12$) from different participants using different scanners and scanning protocols (Srirangarajan et al., 2021). Besides that, we also tested the *BRS* in a completely different task with monetary outcomes using a block design instead of an event-related design on a large sample (1084 subjects) to thoroughly evaluate the generalizability of the predictions from our signature map. In addition, to examine the specificity of the *BRS*, we employed the novel Disgust-Delay Task (DDT, $N = 39$; Fig. 1D), which evokes neural patterns associated with disgust. In this task, we aimed at exploring whether the signature is specific to monetary rewards and losses or rewarding outcomes more generally (i.e. positive versus negative feedback) and whether it is specific to reward or generalizes to emotional salience (i.e. disgusting versus neutral outcomes). The DDT was chosen because it is similar in task structure and solely differs in the neurocognitive processes it is designed to elicit. To test specificity across a wider range of emotions and generalizability to different experimental paradigms, we evaluated the predictions of the *BRS* on the Emotion Viewing task, in which participants passively view actors expressing positive (happiness), neutral and negative (anger, disgust, fear, pain, sadness) emotions. Collectively, data from 5 different fMRI tasks and four independent samples ($N = 1169$) were used to train and test the *BRS*. It is important to note that testing specificity is an open ended process, as numerous different conditions unrelated to outcome processing can be tested, but this is a preliminary validation.

As we are interested in investigating the neural underpinnings of the reinforcement vs punishment processing more generally and not the neural correlates of how much exactly someone earns on the MID, our performance assessment focuses on the signature's relative performance, i.e., whether the signature can predict differences in rewards across conditions. This is because it has been consistently shown across species that value-based choice behavior is context dependent (Bateson et al., 2003, Huber et al., 1982, Shafir et al., 2002) (Bateson et al., 2003). Specifically, it has been found that how a chooser decides between any two options depends on the number or quality of other options in multi-dimensional attribute space ((Huber et al., 1982); Louie et al., 2013). This context-dependence of value based decisions is hypothesized to be implemented on the neural level by means of divisive normalization (Louie et al., 2011, Louie et al., 2013, Louie et al., 2014), where the response of a given neuron is divided by the summed activity of a larger neuronal pool (Carandini and Heeger, 2012). This divisive normalization thus produces context dependence, where the value of an option is explicitly contingent on the value of the other available options, which

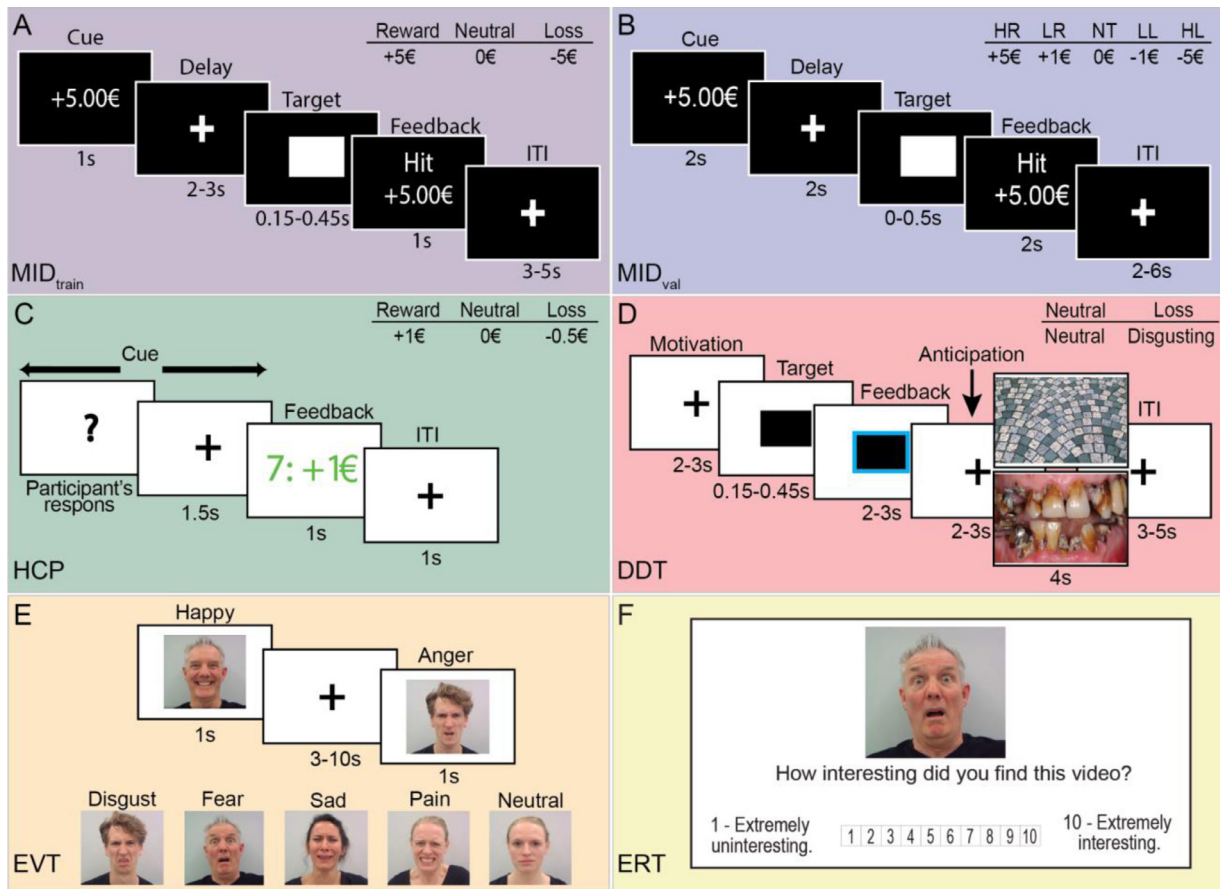


Fig. 1. A) Example trial of the MID_{train} task: Each trial started with a cue informing participants about the money that can be obtained or lost. Subsequently, participants were presented with a fixation cross for a variable amount of time (2-3 s) and the target in the form of a white square appeared for a variable amount of time. Afterwards, participants were informed whether they hit or missed and the associated monetary outcome was presented. Lastly, another fixation cross was presented for a variable amount of time (3-5 s). B) Example trial of the MID_{val} task. The differences to the MID_{train} consisted of differences in timing and number of conditions. C) Example trial of the Gambling task from the HCP: Each trial began with the presentation of the mystery card represented by the question mark and as soon as the participants responded, a fixation cross was presented. Next, participants received feedback about the outcome for 1 s. Lastly, another fixation cross was presented for 1 s. D) Example trial of the DDT: Each trial began with the presentation of a fixation cross (2-3 s) followed by a target which was presented for a duration that adapted to the participants' performance. Next, participants received feedback (2-3 s), viewed another fixation cross (2-3 s) and were then presented with a disgusting or neutral image contingent on their performance. The trials were separated by a fixation cross (3-5 s). E) Example trial of the EVT: Each trial began with the presentation of a fixation cross for a jittered period between 3-10s. Next, participants viewed a video of an actor (four different actors in total) expressing one of 6 emotions (anger, disgust, fear, happiness, pain, sadness) at two different intensities (high and low) or a neutral demeanor for 1s. F) Example trial of the Emotion Rating Task, in which participants had to rate how interesting they found the presented video on a scale from 1 to 10. Only the neutral demeanor and the high intensity emotions from the Emotion Viewing task were used. Participants did not have time constraints to give the response (median reaction time=3.7s [Q1=2.9s Q3=5.0s], and the next trial started immediately after the response was given.

allows efficient coding of information in changing environments. Therefore, our feature selection procedure was based on correlations between actual and regression-predicted rewards, to capture the relative predictive performance, not the absolute predictive performance. This focus on within-subject differences between conditions also has the advantage of being less sensitive to confounding individual differences such as vascular response properties.

2. Methods

For this project, data from five different studies were used. First, to establish the BRS, the Monetary-Incentive-Delay task (MID; Fig. 1A) with three levels of monetary outcomes (+5€, 0€, -5€) was used, which will from now on be referred to as MID_{train}. To test whether the BRS generalizes to the MID task with five levels of monetary outcomes (+5€, +1€, 0€, -1€, -5€) from different participants using different scanners and scanning parameters, openly available data from Srirangarajan and colleagues (2021) was used (Fig. 1B). This dataset will from now

on be referred to as MID validation task (MID_{val}). In addition, to investigate whether the BRS is able to predict differences in reward in a different task with monetary outcomes using a block design instead of an event related design we utilized the Gambling task (Fig. 1C) from the Human Connectome Project (HCP). Further, to assess the construct validity and test whether the signature is specific to monetary reward, we applied it to a moral conflict learning paradigm, in which participants learn that one action leads to high-monetary rewards for themselves and a high-shock to someone else, while the other action leads to low monetary rewards for self and a low-shock to someone else. The results of this analysis are reported elsewhere (Fig. 8 and Supplementary Table 9; Fornari et al., under revision, and included in the reply to reviewers). Finally, to assess whether the BRS generalizes to negative emotional salience, we employed the novel Disgust-Delay Task (DDT; Fig. 1D). While all the above task explore the processing of outcomes that are contingent on participant's choices, to characterize the tuning of the BRS during passive viewing of other people's emotions, we applied it to an unpublished dataset from our laboratory that includes a

novel Emotion Viewing task (EVT; Fig. 1E). As we found that viewing negative facial expressions loaded positively on the *BRS*, a finding that has been associated with information seeking, we performed an online behavioral study (Emotion Rating Task, ERT; Fig. 1F) that confirmed that participants find negative and positive facial expressions more interesting than neutral facial expressions.

2.1. Participants

For the MID_{train} and the DDT task the same 40 participants were used which were collected from a university sample. One participant had a hit rate of zero in both tasks, indicating that the participant never experienced reward. We thus excluded this participant from the analysis. The remaining 39 participants ($M_{age} = 23.62$, $SD_{age} = 3.17$; 28 females) were right-handed with normal or corrected to normal vision, spoke English fluently, were not on any psychoactive medication influencing cognitive function, and had no record of neurological or psychiatric illness. The study was approved by the Erasmus Research Institute of Management (ERIM; Protocol NR: 2018/02/06-61976ssp) internal review board and was conducted according to the Declaration of Helsinki.

For the MID_{val}, nineteen subjects completed the MID task while being scanned with a multi-band acquisition protocol. According to the pre-registered exclusion criteria, data from three subjects were excluded due to excessive motion during at least one of the three task runs, while data from four subjects were excluded due to equipment failure (i.e., faulty response registration by a new button box), leaving twelve subjects total for analyses. For the justification of the sample size and details about participants see the paper by Srirangarajan and colleagues (2021) or contact the authors (Srirangarajan and colleagues).

For the HCP gambling task, task-based fMRI recordings were used from 1206 participants (HCP All Family Subjects). Out of these 1206 participants, 1084 had complete fMRI data for both runs of the Gambling task. Additional behavioral measures on the individual participants can be downloaded from the project website (Van Essen et al., 2012). Individual demographic data is unavailable for these datasets due to data-privacy concerns (Van Essen et al., 2012), but summary demographic data for the 1206 participants (of which only 1084 performed the gambling task and are used here) were reported as $M_{age} = 29.31$, $SD_{age} = 3.67$; 657 females. For the EVT, 34 participants were used which were collected from a university sample. Three subjects were excluded due to missing data and four additional subjects were removed due to excessive image distortion. The remaining 27 participants ($M_{age} = 23.04$, $SD_{age} = 2.29$; 12 females) were right-handed with normal or corrected to normal vision, spoke English fluently, were not on any psychoactive medication influencing cognitive function, and had no record of neurological or psychiatric illness. The study was approved by the University of Amsterdam (UvA; Protocol 2019-EXT-11148) Ethics Review Board of the Faculty of Social and Behavioral Sciences and was conducted according to the Declaration of Helsinki. All participants were reimbursed for their participation with 10€/hour.

For the ERT 200 participants ($M_{age} = 26.33$, $SD_{age} = 9.439$; 96 females) completed five separate experiments that included 40 individuals each. Participants were recruited online from Prolific (<https://www.prolific.co/>), with previous participation in other studies from the lab as exclusion criteria. The study was approved by the University of Amsterdam (UvA; Protocol 2022-EXT-15474) Ethics Review Board of the Faculty of Social and Behavioral Sciences and was conducted according to the Declaration of Helsinki. Participants' remuneration was 7.5€/hour.

2.2. Task and Stimuli

2.2.1. Monetary-Incentive-Delay task (MID_{train})

The MID_{train} consisted of 108 trials of approximately 9 s each. During each trial, participants saw one of three cues (cue phase, 1 s), were then asked to fixate on a crosshair as they waited a variable interval (delay phase, 2000–3000 ms), and then responded to a white target square

that appeared for a variable length of time (target phase, 150–450 ms) with a button press (Fig. 1A). Feedback (outcome phase, 1 s), which followed the disappearance of the target, notified participants whether they had won or lost money during that trial. On incentivized trials, participants could win or avoid losing money by pressing the button during target presentation. On neutral trials, no money could be won or lost. Task difficulty, in the form of the length of time the target was presented, was set adaptively throughout the task such that each participant should succeed on 66% of his or her target responses. This was done to make subjects with different performance levels comparable and prevent participants from getting frustrated. Cues signaled potential reward (+ 5.00 €), potential loss (- 5.00 €), or no monetary outcome (0 €). Trial types were pseudo-randomly ordered within each session (Knutson et al., 2000). Participants were instructed that at the end of the experiment one trial would randomly be chosen and that the performance on this trial would determine their remuneration. In the MID task we focus on the feedback phase as we are interested in the neural response associated with receiving a monetary outcome. We acknowledge that we re-used text for the task description from (Knutson et al., 2000).

2.2.2. MID_{val}

Since the main goal of the study by Srirangarajan and colleagues (2021) was to examine whether acquiring fMRI data with multi-band versus single-band scanning protocols compromises detection of mesolimbic activity during reward processing, the fMRI data was collected in three runs. Importantly, the MID task was identical across all three runs. The MID_{val} was similar to the MID_{train} with some exceptions. First, the MID_{val} included six task trial conditions: a large gain condition (+5.00 \$); a medium gain condition (+1.00 \$); a no gain condition (+ \$0.00); a no loss condition (- \$0.00); medium loss condition (- 1.00 \$); and a large loss condition (-5.00 \$). Each trial condition was repeated 12 times in a pseudorandom order, totalling 72 trials. Furthermore, timing differed slightly. The cue phase was now 0–2 s, the delay phase was 2–4 s, the target phase appeared briefly between 4–4.5 s, the outcome phase lasted 6–8 s, and the Inter-Trial Interval lasted 2, 4, or 6 s. Thus, each trial lasted an average of 12 s (including the ITI). As before, adaptive timing of target duration within condition ensured that subjects succeeded in “hitting” targets on approximately 66% of the trials (Knutson et al., 2005). Thus, each MID task run lasted 864 s in total (approximately 14.4 min), and all three runs were acquired during a single session, but with counterbalanced ordering across subjects. We acknowledge that we re-used text for the task description from Srirangarajan and colleagues (2021).

2.2.3. Gambling task from the Human Connectome Project (HCP)

This task was adapted from the Gambling task developed by Delgado and Fiez (Delgado et al., 2000). Participants played a card guessing game where they were asked to guess the number on a mystery card (represented by a “?”) in order to win or lose money (Fig. 1C). Participants were told that potential card numbers ranged from 1–9 and were asked to indicate whether they expected the mystery card number to be more or less than 5 by pressing one of two buttons on the response box. Feedback was the number on the card generated by the program as a function of whether the trial was a reward, loss or neutral trial, and could result in: 1) a green up arrow with “\$1” for reward trials, 2) a red down arrow next to -\$0.50 for loss trials; or 3) the number 5 and a gray double headed arrow for neutral trials. The “?” was presented for up to 1500 ms (if the participant responds before 1500 ms, a fixation cross was displayed for the remaining time), followed by feedback for 1000 ms. There was a 1000 ms ITI with a “+” presented on the screen. The task was presented in blocks of 8 trials that are either mostly reward (6 reward trials pseudo-randomly interleaved with either 1 neutral and 1 loss trial, 2 neutral trials, or 2 loss trials) or mostly loss (6 loss trials pseudo-randomly interleaved with either 1 neutral and 1 reward trial, 2 neutral trials, or 2 reward trials). In each of the two runs, there were 2 mostly

reward and 2 mostly loss blocks, interleaved with 4 fixation blocks (15 s each). This experiment was designed to be analyzed in blocks of mainly reward blocks and mainly loss blocks. As a consequence, here we do not analyze a specific period within each trial, but the average activation across several trials within each block type. We acknowledge that we re-used text for the task description from (Delgado et al., 2000).

2.2.4. Disgust-Delay Task (DDT)

A new paradigm termed the Disgust-Delay-Task (DDT) inspired by the MID task (Knutson et al., 2000) was developed (Fig. 1D). In this task, participants had to press a button during the presentation of a target stimulus, i.e., a black rectangle. They were then informed, during the feedback phase, about whether the trial was a success or not. However, instead of winning money, or avoiding losing money, during the outcome phase, participants then either saw a disgusting image or a neutral image depending on their performance. Disgusting images were selected based on a pretest that ensured that these images evoked disgust specifically and no other negatively valenced emotions (see Appendix 1). On each trial of the DDT, participants were first presented with a fixation cross for 2-3s (Fig. 1D). Subsequently, the target stimulus was presented for 150-450 ms depending on the participants' performance. As in the MID tasks above, an adaptive algorithm was implemented which varies the duration to ensure an equal number of successful and unsuccessful trials (50% each). Afterwards, the participants received feedback whether or not they hit the target in time for a period that varied between 2-3 s. This was followed by another fixation cross that varied between 2-3 s. The trial ended with the presentation of either a neutral image or a disgusting image for 4 s depending on whether the participant hit or missed the target. Next, participants had to wait for a period jittered between 3-5 s. Participants completed 72 trials of the DDT. Here, we can thus analyse two periods of interest. During the feedback period, we can investigate the impact of a non-financial reinforcer (i.e., success or failure feedback) on brain activity. During the outcome phase, we can investigate the impact of neural response to the experience of disgust triggered by the disgusting images.

2.2.5. Emotion Viewing Task (EVT)

The EVT will be subject to a full publication, and will only be described briefly here. Participants viewed actors expressing six different emotions at two different intensities. Specifically, there were four different actors that differed in age and gender (females aged 29, 27; males aged: 24, 54). Each of the four actors expressed anger, disgust, fear, happiness, pain and sadness at high and low intensity. In addition, the actors also expressed neutral demeanor, with blinking as the only deliberate movement. Thus, leading to a total 4 actors x 6 emotions x 2 intensities + 4 actors x 1 neutral = 52 videos (one more condition, a neutral facial expression with deliberate facial movements was included in the study but will not be analyzed or reported here). The task was administered in eight runs, in which the 52 videos were repeated in random order. Participants were instructed to view the videos while feeling with the actor, without moving facial muscles or silent verbalization of the emotion name. On each trial of the EVT, participants were presented with a fixation cross, jittered between 3-10 s, and subsequently saw a video for 1s. In total, the whole task lasted for approximately 75 minutes. The stimuli had been selected from a larger pool of recorded facial expressions based on an online validation. In the validation, participants reported how much of each emotion (anger, disgust, fear, happiness, pain and sadness) was visible in each video on a ten point scale. We then selected the neutral and the low and high intensity movies for each emotion so that (a) the average rating did not differ across the emotions (e.g. there was as a high an angry rating for the Anger video as painful rating in the Pain video) for a given intensity, (b) the rating was higher for the high than the low intensity movies and (c) higher in the low intensity than in the neutral movies, and (d) ratings on off-target emotions

was minimal (e.g. a High-Disgust movie would not have high ratings on any other emotion).

2.2.6. Emotion Rating Task

It has been argued that negatively valenced stimuli activate reward processing, because they are more interesting than neutral stimuli, and participants are motivated to seek information (Oosterwijk et al., 2020). To verify the hypothesis that our positive and negative emotional facial expressions are more interesting to watch than our neutral facial expressions, a separate pool of participants (online) were asked to rate 'How interesting did you find this video' on a Likert scale from 1 (Extremely uninteresting) to 10 (Extremely interesting) on part of the movies used in the Emotion Viewing Task. Only high intensity movies were used, together with the neutral condition. Because there was only one positive emotion (happiness) and one neutral, but several negative emotions (anger, fear, disgust, pain, and sadness), an imbalance that may bias reports, we asked five separate pools of participants to view a balanced set of three types of videos each: the neutral, the high intensity positive and one of the high intensity negative facial expressions, for a total of 3 categories x 4 actors = 12 ratings each, in randomized order.

2.3. fMRI acquisition

For MID_{train} and DDT, the fMRI images were collected using a 3T Siemens Verio MRI system. Functional scans were acquired by a T2*-weighted gradient-echo, echo-planar pulse sequence in descending interleaved order (3.0 mm slice thickness, 3.0 × 3.0 mm in-plane resolution, 64 × 64 voxels per slice, flip angle = 75°). TE was 30 ms, and TR was 2,030 ms. A T1-weighted image was acquired for anatomical reference (1.0 × 0.5 × 0.5 mm resolution, 192 sagittal slices, flip angle = 9°, TE = 2.26 ms, TR = 1900 ms).

For MID_{val}, all data were acquired on a 3 Tesla General Electric scanner with a 32-channel head coil at the Stanford Center for Cognitive and Neurobiological Imaging (CNI). Structural (T1-weighted) scans were first acquired for all participants. Functional (T2*-weighted) images for single-band and multi-band scans were then acquired using the following common parameters: TE = 25 ms, FOV = 23.8 × 23.8 cm; 2 acquisition matrix = 70 × 70, no gap, phase encoding = PA, voxel dimensions = 3.4 × 3.4 × 3.4 mm. Additional parameters that varied between scanning protocols included: (1) multi-band factor = 1, TR = 2000 msec, flip angle = 77°, number of slices = 41; (2) multi-band factor = 4, TR = 500 msec, flip angle = 42°, number of slices = 32; (3) multi-band factor = 8, TR = 500 msec, flip angle = 42°, number of slices = 41. All fMRI data were reconstructed using 1D-GRAPPA (Blaimer et al., 2013). For more information about the scanning protocol please refer to the paper by Srirangarajan and colleagues (2021).

For the HCP project, the data was collected using a customized 3T Siemens Connectome Skyra with a standard 32-channel Siemens receiver head coil and a body transmission coil. T1-weighted high-resolution structural images were acquired using a 3D MPRAGE sequence with 0.7 mm isotropic resolution (FOV = 224 × 224 mm, matrix = 320 × 320, 256 sagittal slices, TR = 2400 ms, TE = 2.14 ms, TI = 1000 ms, FA = 8°) and used to register functional MRI data to a standard brain space. Functional MRI data were collected using gradient-echo echo-planar imaging (EPI) with 2.0 mm isotropic resolution (FOV = 208 × 180 mm, matrix = 104 × 90, 72 slices, TR = 720 ms, TE = 33.1 ms, FA = 52°, multiband factor = 8, 253 frames, ~3 m and 12 s/run).

For the EVT, the fMRI images were collected using a 7T Phillips MRI system equipped with an Tx8/Rx32 rf-coil (Nova Medical). Functional scans were acquired by a T2*-weighted gradient-echo 3D echo-planar imaging (EPI; 1.6 mm slice thickness, 1.6 × 1.6 mm in-plane resolution, 128 × 128 voxels per slice, flip angle = 13°). TE was 19.45 ms, and TR was 1816 ms. A T1-weighted image (MPRAGE) was acquired for anatomical reference (0.8 × 0.8 × 0.8 mm resolution, 232 sagittal slices, flip angle = 8°, TE = 3.29 ms, TR = 3000 ms).

2.4. Preprocessing

For the MID_{train} , MID_{val} and the DDT, the fMRI data were preprocessed using fMRIPrep version 1.0.8, a Nipype based tool (Gorgolewski et al., 2011). We chose fMRIPrep because it addresses the challenge of robust and reproducible preprocessing as it automatically adapts a workflow based on best-in-class algorithms to virtually any dataset, enabling high-quality preprocessing without the need of manual intervention (Esteban et al., 2019). Each T1w volume was corrected for intensity nonuniformity and skullstripped. Spatial normalization to the International Consortium for Brain Mapping 152 Nonlinear Asymmetrical template version 2009c (Esteban et al., 2016) was performed through nonlinear registration, using brain-extracted versions of both T1w volume and template. Brain tissue segmentation of cerebrospinal fluid (CSF), white matter (WM), and gray matter was performed on the brain-extracted T1w. Field map distortion correction was performed by coregistering the functional image to the same-subject T1w image with intensity inverted (Caballero-Gaudes and Reynolds, 2017) constrained with an average field map template (Tustison et al., 2010). This was followed by coregistration to the corresponding T1w using boundary-based registration (Smith et al., 2002) with 9 degrees of freedom. Motion correcting transformations, field distortion correcting warp, blood oxygen level-dependent images-to-T1w transformation, and T1w to template Montreal Imaging Institute (MNI) warp were concatenated and applied in a single step using Lanczos interpolation. Physiological noise regressors were extracted using CompCor (Cox and Hyde, 1997). Principal components were estimated for the two CompCor variants: temporal (tCompCor) and anatomical (aCompCor). Six tCompCor components were then calculated including only the top 5% variable voxels within that subcortical mask. For aCompCor, six components were calculated within the intersection of the subcortical mask and the union of CSF and WM masks calculated in T1w space, after their projection to the native space of each functional run. Frame-wise displacement (Treiber et al., 2016) was calculated for each functional run using the implementation of Nipype. For more details of the pipeline, see <https://fmripred.org/en/latest/workflows.html>. After the preprocessing the voxel size of the images is $3 \times 3 \times 3.5$ mm.

For the HCP data, preprocessing of the images included motion correction, distortion correction, co-registration and normalized to MNI space as described in the HCP 1200 Subjects Release (Glasser et al., 2013).

For the EVT, the whole brain fMRI data was preprocessed and analyzed using SPM12 (7771; Wellcome Trust Centre for Neuroimaging, UCL, UK) with MATLAB R2020b version 9.9.0 (The MathWorks Inc., Natick, USA). The preprocessing pipeline was organized as follows: realignment to the first image of every run and then to the estimated average (two-pass), co-registration of anatomical images to the mean functional image (rigid body transformation, $DOF = 6$), segmentation of the anatomical scan that yields normalization parameters that were then used to bring the EPI images to MNI-space and voxel sizes were resampled to $1 \times 1 \times 1$ mm for the functional images. In order to completely incorporate the entire brain (including cerebellar areas in all scans, the bounding box settings were changed to $[-90 -126 -72; 90 90 108]$, as SPM's default settings have been reported as a risk to omit some of the cerebellar areas (Gazzola and Keysers, 2009).

2.5. Statistical analyses

2.5.1. MID_{train} & MID_{val}

To model all possible outcomes of the MID tasks for every participant, we estimated a general linear model (GLM) using regressors for onsets of the outcome phase for successful high reward trials (HR-won: received +5.00 €), unsuccessful high reward trials (HR-lost: did not receive +5.00 €), successful low reward trials (LR-won: received +1.00€; for MID_{val} only), unsuccessful low reward trials (LR-lost: did not receive +1.00€; for MID_{val} only), successful neutral trials (NT-won: 0 €; for the

MID_{val} the neutral gain, i.e. +0 €, and neutral loss trials, i.e. -0 € were combined), unsuccessful neutral trial (NT-lost: 0€), successful low loss trials (LL-won: did not lose 1.00 €; for MID_{val} only), unsuccessful low loss trials (LL-lost: did lose 1.00 €; for MID_{val} only), successful high loss trials (HL-won: did not lose 5.00€) and unsuccessful high loss trials (HL-lost: lost 5.00€). The duration of the epoch for the outcome phase was 1 s, and the beginning of the outcome phase was used as onset time. Average background, WM and CSF signal, framewise displacement, six head motion regressors, and six aCompCor (which are component based noise correction regressors) regressors, all obtained from fMRIPrep, were entered as regressors of no interest. First, a smoothing kernel of 5 mm full width at half maximum was applied. For consistency, the same smoothing procedure was applied to all other datasets as well. Subsequently, all regressors of interest (but not regressors of no interest) were convolved with the canonical hemodynamic response function. Linear contrasts were computed between HR-won and HR-lost trials, LR-won and LR-lost trial, NT-won and NT-lost trials, LL-lost and LL-won trial, HL-lost and HL-won trials. These contrasts were chosen to isolate the effect of receiving or losing money by means of comparing each regressor with the regressor of opposite outcome within the same condition. As a consequence, only neural activation related to receiving or losing money should remain as all other aspects of the contrasted trials are the same. The resulting subject level t-maps were then converted to z-maps. Here, we use the z-maps as the primary input to our multivariate pattern analysis because z-maps represent effect-sizes in units of variance, that should be more comparable across experiments and designs than the simple difference between the parameter estimates, which are in arbitrary units, or the t-maps that depend on the sample size in terms of acquired volumes. As the purpose of the study by Srirangarajan and colleagues (2021) was to test whether acquiring fMRI data with multi-band versus single-band scanning protocols compromises detection of mesolimbic activity during reward processing, the fMRI data was collected in three runs. For this study we were however not interested in the effects of scanning protocols. As a consequence, we averaged over the z-maps for each subject across the three runs to increase the signal to noise ratio.

2.5.2. DDT

To model the experience of disgust and the experience of viewing neutral images we estimated a GLM using regressors for onsets of the picture presentation phase of the DDT for the presentation of disgusting images and neutral images. The duration of the epoch for the picture presentation phase was 4 s, and the beginning of the picture presentation phase was used as onset time (see Fig. 1D).

In addition, to explore whether the BRS predicts monetary outcomes specifically or generalizes to rewarding versus loss outcomes more generally, we modeled the feedback phase of the DDT. As the structure of the MID and the DDT are very similar the only difference here is that instead of monetary outcome the feedback is purely motivational. The duration of the epoch for the feedback phase was 2 s since this was the minimum of time it lasted on every trial. We defined the feedback phase by counting back two seconds from the onset of the Anticipation phase (see Fig. 1D). Lastly, to have a neutral period to compare the neural patterns associated with disgusting and neutral images to, we modeled the neural activation of viewing the fixation cross at the beginning of each trial (Motivation Delay). This period was chosen because it was most distant in time from the picture presentation phase. The duration of the epoch for the motivation delay was 2 s since this was the minimum of time it lasted on every trial (see Fig. 1D). As above, average background, WM and CSF signal, framewise displacement, six head motion regressors, and six aCompCor regressors, all obtained from fMRIPrep, were entered as regressors of no interest. First, a smoothing kernel of 5 mm full width at half maximum was applied. Next, all regressors of interest (but not the nuisance regressors) were convolved with the canonical hemodynamic response function. Linear contrasts were computed between the presentation of disgusting images and the fixation period and

the presentation of a neutral image and the fixation period. As before, the subject level t-maps were converted to z-maps to render them more comparable across experiments.

2.5.3. HCP

Since the HCP gambling task was administered in a block design and the ITIs between trials were short we employed a GLM using regressors for onsets of the reward blocks, loss blocks and fixation blocks. The duration of the reward and loss blocks were 28s each whereas the fixation period was 15s. Twelve motion regressors (x translation in mm, y translation in mm, z translation in mm, x rotation in degrees, y rotation in degrees, z rotation in degrees, derivative of x translation, derivative of y translation, derivative of z translation, derivative of x rotation, derivative of y rotation, derivative of z rotation), the absolute root mean square (RMS) motion and the relative RMS motion, obtained from the HCP preprocessing pipeline, were added as regressors of no interest. Different nuisance regressors were applied here as the data was obtained in preprocessed format from the HCP website and only the 14 regressors mentioned in the previous sentence were available. As before, as a first step, a smoothing kernel of 5 mm full width at half maximum (FWHM) was applied. Afterwards, all regressors of interest (but not the regressors of no interest) were convolved with the canonical hemodynamic response function. Linear contrasts were computed between the reward block and the fixation block, the loss block and the fixation block and the fixation block and the baseline. Again, the resulting subject level t-maps were subsequently converted to z-maps.

2.5.4. EVT

Data were analyzed using a GLM that contained one regressor per stimulus category (i.e. one for Anger High, one for Anger Low, one for Pain High, one for Pain Low, ..., one for neutral), modeled as a boxcar of duration 1s aligned on each movies onset, then convolved with the hemodynamic response function. All six head motion regressors obtained from the preprocessing pipeline, were entered as regressors of no interest. A smoothing kernel of 5 mm full width at half maximum was applied to the EPI images. Linear contrasts were computed to sum the parameter estimates for each video type across the runs. As before, the subject level t-maps were converted to z-maps to render them more comparable across experiments.

2.5.5. ERT

For each of the five pools of participants separately, ratings were analyzed with paired samples t-tests that compared the negative emotion presented for that group against happy and neutral. Student t-tests were used for parametric data, and Wilcoxon signed-rank for non-parametric data. Both p-values and Bayes Factor values were calculated.

2.6. Multivariate pattern analyses

2.6.1. Creation of the BRS

We used the normalized and smoothed (5mm FWHM) z-maps to develop population-level reward-predictive patterns, as previous studies suggested that smoothing could improve inter-subject functional alignment while retaining sensitivity to mesoscopic activity patterns that are consistent across subjects (Etzel et al., 2011, Op de Beeck, 2010, Shmuel et al., 2010). A LASSO-PCR model (least absolute shrinkage and selection operator-regularized principal components regression; (Wager et al., 2011, Wager et al., 2013)) was then trained on the whole-brain maps from the subject level z-maps derived from the analyses described above. The rationale behind using LASSO-PCR is twofold. First, to deal with the fact that fMRI datasets contain many voxels with correlated signals that are challenging for regression analysis, LASSO-PCR does not use each voxel as an individual predictor, but applies principal component analysis to the fMRI data to summarize the data using orthogonal components. Next, to focus on the most informative components, a LASSO regression is used to predict the outcome variable

(reward magnitude in our case) from component scores, which adds a penalty term to the model to shrink less important principal component coefficients to zero. Specifically, the LASSO-PCR model was trained on the z-maps (HR-won > HR-lost, NT-won > NT-lost; HL-lost > HL-won) from the MID_{train} to predict the 3 different levels of monetary outcome (+5.00 €, 0.00 € & -5.00 €). For feature selection, we identified voxels that correlated more strongly with reward rather than salience. As explained in the introduction, this was done to maximize relative prediction performance rather than absolute prediction, because reward processing has been found to be context dependent ((Bateson et al., 2003, Huber et al., 1982, Louie et al., 2013, Simonson, 1989) and there are no absolute values assigned to individual options. Specifically, given the three parameter estimate images for each participant (High Reward: HR-won > HR-lost, Neutral: NT-won > NT-lost, High Loss: HL-lost > HL-won), we can consider two codings: one for outcome (1, 0, -1) and one for salience (1, 0, 1). We can then compute the Spearman correlation between the parameter estimates V_j at each voxels j and the outcome and salience coding separately for each subject within the cross validation loop. As we know that the spacing is uncertain, because rewards might not be equidistant from zero as losses (Kahneman, 2011), we use the Spearman instead of the Pearson correlation. We then selected voxels such that $r(V_j, Outcome) \neq 0$ and $|r(V_j, Outcome)| > |r(V_j, Salience)|$. At the group level, to do this, we first performed a two-sided Wilcoxon signed-rank test on the correlation between voxel values and outcome coding $r(V_j, Outcome)$ and then a one-sided Wilcoxon signed-rank test on the difference between absolute values of the correlation between voxel values and outcome and voxel values and salience $|r(V_j, Outcome)| > |r(V_j, Salience)|$. We then selected all voxels for which $p_{r(V_j, Outcome)} \neq 0 < \alpha$ and $p_{|r(V_j, Outcome)| > |r(V_j, Salience)|} < \alpha$ where α was chosen permissively at $\alpha=0.5$ to allow for a reasonable amount of voxels to enter the LASSO-PCR model. More conservative thresholds were also applied to test the robustness of the findings (see Appendix 2). To reiterate the feature selection procedure, we correlated for each subject the parameter estimates for each of the three conditions (High Reward, Neutral & High Loss) with the two codings (outcome and salience) at each voxel, to select the voxels that correlate more strongly with the outcome coding than with the salience coding, while making sure that the voxels respond to the outcome coding. This was done on each iteration of the cross-validation on the training set to only allow voxels to enter the LASSO-PCR model that respond stronger to outcome than to salience.

The feature selection and model fitting were implemented using a 5-fold cross-validation procedure during which all participants were randomly assigned to 5 different subsamples while ensuring that all images from an individual subject remained within a subsample and does not spread across subsamples. We always used 4 subsamples for training and one for testing. As a result, out-of-sample prediction is always done on new individuals, which prevents dependence across images from the same participants invalidating predictive accuracy. We obtained predicted values by computing the dot product of the weight map computed over 4 of the subsamples (at each iteration) and the z-maps of the left out subsamples and adding the intercept computed over the four subsamples for each subject and condition (at each iteration). To evaluate the predictive accuracy of the model, the Spearman correlation between the predicted monetary outcome levels and the actual outcomes for the left-out subsample were computed at each fold, and then the correlations were averaged across folds. In accordance with the mass-univariate analyses and to identify which brain regions made reliable contributions to the model (Wager et al., 2013, Zhou et al., 2020), the pattern maps were thresholded at $p < 0.001$ (two-tailed; uncorrected) using bootstrap procedures with 5000 samples. The result was a spatial pattern of regression weights across the whole brain that significantly contributed to the prediction of monetary out-of-sample outcomes in the MID_{train}. To test for robustness, we also applied a more conservative threshold at FDR $p < 0.05$ (two-tailed) and a procedure in which we first selected only voxels that were non-zero in at least 90% of the bootstrap iteration and then applied FDR correction at $p < 0.05$ (see Ap-

pendix 2). We also computed the Bayes-Factor for the correlation between predicted and actual monetary outcome values to also be able to test for evidence of absence of an effect (Keyzers et al., 2020). To calculate the Bayes-Factor for the correlation, Jeffreys exact Bayes Factor was used (Ly et al., 2016) as implemented in the Pingouin python package (Vallat, 2018). In addition, we evaluated whether the *BRS*'s predictions within a given condition (High Reward, Neutral, High Loss) are significantly different from zero, by means of a one sample t-test against zero. Since not all of the predictions across conditions and experiments were normally distributed we used the Wilcoxon signed-rank test and the associated Bayes factors were computed as proposed by (van Doorn et al., 2020), with a Cauchy prior with the scale $\frac{1}{\sqrt{2}}$. To compare the *BRS*'s predictions between conditions Wilcoxon rank-sum tests were employed and to compute the Bayes Factors we again used the procedure proposed by (van Doorn et al., 2020).

We also conducted within-person forced-choice discrimination, where two activation maps from the same participant were compared, and the image with the higher overall signature response (i.e., the stronger expression of the signature pattern) was classified as associated with higher reward. If this matched the actual labels the classification was labeled as correct and otherwise as incorrect. Subsequently the average across subjects was computed to determine forced-choice accuracy. We conducted these forced-choice tests for all combinations of conditions (i.e., HR vs. NT, NT vs HP & HR vs HP). The advantage of the forced-choice test is that it is 'threshold free' in the sense that an absolute decision threshold across individuals is not required; zero is used as the threshold for the difference between the two paired alternatives (Wager et al., 2013). Thus, individual differences in the shape and amplitude of the blood oxygen level dependent (BOLD) fMRI response do not add noise in this kind of test. To test for significance, permutation tests were used where the order of conditions was permuted ($N = 10000$) and the accuracy was computed again. The empirical classification accuracy was then compared to the null distribution of accuracies based on permuted values to obtain p-values.

2.6.2. Validation on the *MID_{val}*

To test how well the *BRS* generalizes to new data involving monetary outcomes the *MID_{val}* was used. Specifically, we tested whether the *BRS* generalizes to a MID task with five levels of monetary outcomes (+5 €, +1 €, 0 €, -1 €, -5 €) from different participants using different scanners and scanning parameters. To this end, we obtained pattern expression values by computing the dot product of the cross-validated weightmap (averaged across folds) of the reward pattern (created on the *MID_{train}*) and the z-maps and adding the intercept (averaged across folds) for each subject and condition from the *MID_{val}*. For the *MID_{val}* the High Reward (HR-won > HR-lost), Low Reward (LR-won > LR-lost), Neutral (NT-won > NT-lost), Low Loss (LL-lost > LL-won) and High Loss (HL-lost > HL-won) contrasts were used. The resulting pattern expression represents scalar response values, which constitute the predicted monetary outcome for the given condition. The pattern expression values were then tested for differences between experimental conditions. We calculated the Spearman correlation between the pattern expression values and the actual monetary outcome values for each of the conditions (+5 €, +1 €, 0 €, -1 €, -5 €), with higher correlations representing higher predictive accuracy, in the sense of variance of rewards explained by the pattern expression values. Specifically, the predicted monetary outcome values obtained from the dot multiplication (5 conditions * 12 subjects = 60 predicted monetary outcome values) were correlated with the 60 actual monetary outcomes in a single correlation. To estimate significance of the predictive performance, a permutation test ($N = 5000$) was performed where the true monetary outcome values were shuffled and the procedure was repeated. To assess the robustness of the estimation of significance we also repeated the permutation tests with the root mean squared error as a predictive performance evaluation metric ($N = 5000$). To test whether the predictions made by the *BRS* in the different conditions were different from zero and whether

predictions between conditions were significantly different from each other, the same procedure as detailed above was used. As before, we conducted within-person forced-choice discrimination to further assess the predictive accuracy of the *BRS*. As above, permutation testing was used to evaluate statistical significance of classification accuracies.

2.6.3. Validation on the HCP Gambling task

To investigate whether our *BRS* generalizes to a completely different task involving monetary outcomes, the HCP Gambling task was used. Specifically, we tested whether the *BRS* generalizes to the Gambling with three different levels of monetary outcomes (+1 €, 0 €, -0.5 €) that were not symmetrically distributed around zero. As before, we obtained pattern expression values by computing the dot product of the cross-validated weightmap (averaged over folds) of the reward pattern (created on the *MID_{train}*) and the z-maps and adding the intercept (averaged over folds) for each subject and condition from the HCP Gambling task and then tested the predictive performance using the Spearman correlation between actual monetary outcomes and predicted monetary outcome values (3 conditions * 1084 subjects = 3252 predicted monetary outcome values). As above, permutation tests were used to estimate significance. To test whether the predictions made by the *BRS* in the different conditions were different from zero and whether predictions between conditions were significantly different from each other, the same procedure as detailed above was used. Again, we conducted within-person forced-choice discrimination, to further assess the predictive accuracy of the *BRS*. As above, permutation testing was used to evaluate statistical significance of classification accuracies. To evaluate the test-retest reliability of the HCP Gambling task, we also computed the pattern response to the first and second run separately and then calculated the Pearson, Spearman and intraclass correlation between the pattern responses for the two runs. We chose to assess test-retest reliability for the HCP specifically because it was the only sample large enough to get meaningful estimates of test-retest reliability.

2.6.4. Testing the specificity on the DDT task

For specificity, the signature expression should not significantly differ from zero when applied to z-maps from tasks involving other types of emotionally salient outcomes. To assess the specificity of our *BRS* we employed the DDT task. We explored whether the *BRS* also predicts disgusting (coded as -1) versus neutral outcomes (coded as 0). In addition, we also tested whether the *BRS* would be able to predict positive or negative feedback in the disgust delay task. This was done to explore whether the *BRS* predicts monetary outcomes specifically or generalizes to rewarding versus loss outcomes more generally. As before, we obtained pattern expression values by computing the dot product of the cross-validated weight map (averaged over folds) of the reward pattern (created on the *MID_{train}*) and the z-maps and adding the intercept (averaged over folds) for each subject and condition from the DDT task and then tested the predictive performance using the Spearman correlation between actual emotional outcomes (neutral vs disgusting images) and predicted emotional outcomes (2 conditions * 39 subjects = 78 predicted emotional outcome values). As above, permutation tests were used to estimate significance. To test whether the predictions made by the *BRS* in the different conditions were different from zero and whether predictions between conditions were significantly different from each other, the same procedure as detailed above was used. As above, we conducted within-person forced-choice discrimination, to further assess the predictive accuracy of the *BRS*. As above, permutation testing was used to evaluate statistical significance of classification accuracies.

2.6.5. Testing the specificity and generalizability of the EVT task

To further characterize the specificity and generalizability to other experimental task structures of our *BRS* we employed the EVT task. We investigated whether simply viewing the emotions of others loads onto the *BRS*, compared to neutral facial expressions, or whether the *BRS* is specific for first person rewards or losses. In addition, we explore

Table 1
Coding of outcome and salience for feature selection.

Parameter estimate image	Outcome	Salience
High Reward (HR)	1	1
Neutral (N)	0	0
High Loss (HL)	-1	1

whether viewing negatively valenced facial expressions (actors expressing anger, disgust, fear, pain, sadness) would load negatively on the BRS, in analogy to first person losses, or whether they would load positively on the BRS as has been proposed in studies of morbid curiosity where the information content of other people's negative emotions is rewarding (Oosterwijk et al., 2020). As before, we obtained pattern expression values by computing the dot product of the cross-validated weight map (averaged over folds) of the reward pattern (created on the MID_{train}) and the z-maps and adding the intercept (averaged over folds) for each subject and condition from the EVT task. To test whether the predictions made by the BRS in the different conditions were different from zero and whether predictions between conditions were significantly different from each other, the same procedure as detailed above was used. As above, we conducted within-person forced-choice discrimination, to further assess the predictive accuracy of the BRS. The difference to the analysis above is that here we computed two-sided p-values since there were no clear predictions about whether observing negative emotions of others should load positively or negatively on the BRS compared to the neutral expressions. As above, permutation testing was used to evaluate statistical significance of classification accuracies.

3. Results

3.1. Within-task prediction

To create a generalizable BRS we first trained and tested our LAS-SOPCR model on the MID_{train} using 5-fold cross validation and a threshold of $p < 0.5$ (threshold was applied within the cross-validation loop) for the feature selection procedure. The analysis revealed that outcomes in the left-out cross-validation folds in the MID_{train} could be significantly predicted by the BRS ($RMSE = 2.89$, $p_{perm} < 0.001$, $r = 0.72$, $p_{perm} < 0.001$, $BF_{10} > 1000$). The feature selection procedure selected 39% of voxels across the whole brain (Fig. 2A). Using the bootstrap procedure, we observed that particularly voxels in the dorsal striatum and the ventromedial prefrontal cortex (vmPFC) significantly contributed to the predictive success of our model (at $p < 0.001$; Fig. 2B and Table 2; for other thresholds see Appendix 2). Fig. 3A shows the signature values obtained when multiplying the z-maps of the individual participants with the thresholded ($p_{bootstrap} < 0.001$; see methods) BRS. For the forced choice analysis we observed significant classification accuracies for all tests. However, classification accuracy was substantially higher between rewarding and loss conditions and neutral and loss trials than between reward and neutral conditions (see Table 3).

Table 2
Clusters for the significant voxels identified by the bootstrap procedure.

Region	peak_x	peak_y	peak_z	peak_value	volume_mm	nr_voxels
R Dorsal Striatum	24	14	-2	722.053	3456	432
L Dorsal Striatum	-20	12	-8	97.371	3152	394
R Occipital Pole	16	-92	-8	679.417	1464	183
vmPFC	2	44	-4	576.394	1248	156

L = Left; R = Right; vmPFC = ventromedial prefrontal cortex. Only clusters of at least 50 voxels are shown, a complete list can be found in Supplementary Table S2. All voxels were used in the analysis. The table was generated using the python package Atlasreader (Notter et al., 2019)

Table 3
Forced choice accuracies (%) for the MID_{train}.

	HR	N	HL
HR	-		
N	67***	-	
HL	92***	95***	-

HR = High Reward, N = Neutral; HL = High Loss; * = $p_{perm} < 0.05$; ** = $p_{perm} < 0.01$; *** = $p_{perm} < 0.001$.

3.2. Meta-analytic decoding of the BRS map

To functionally characterize the BRS, the Neurosynth (Yarkoni et al., 2011) decoder function was used to assess its similarity to the reverse inference meta-analysis maps generated for the entire set of terms included in the Neurosynth dataset. Here the unthresholded z-map obtained through the bootstrap procedure was used, since the neurosynth decoder works best on unthresholded whole brain maps. The most relevant features were 'reward' and 'monetary' for the top 50 terms (excluding anatomical terms) ranked by the correlation strengths between the BRS map and the meta-analytic maps (see word cloud, size of the font scaled by correlation strength, Fig. 2C).

3.3. Testing the generalizability on the MID_{val}

To test the generalizability of the BRS map we tested the prediction performance on the MID_{val}. This allowed us to evaluate how well the BRS is able to predict relative reward magnitude based on activation patterns in new participants from a different scanner and with a different number of levels of monetary outcomes. Using the significant voxels from the BRS map in Fig. 2B we observed a significant prediction of the relative monetary outcomes on the MID_{val} ($RMSE = 2.97$, $r = 0.75$, $p_{perm} < 0.001$, $BF_{10} > 1000$; Fig. 3B). To test the robustness of this finding the prediction was also repeated using all voxels, and the FDR-corrected map ($p < 0.05$; see Appendix 2), and a map derived from first selecting the most consistent voxels and correcting using FDR (see Methods). The robustness checks revealed very similar significant predictions on the MID_{val} (see Appendix 2). For the forced choice analysis we observed significant classification accuracies for the tests comparing the reward to the loss condition and the neutral to the loss condition. No significant classification accuracies were observed when contrasting rewarding and neutral trials. In addition, no significant classification accuracy was observed when comparing high and low loss trials (see Table 4).

3.4. Testing the generalizability on the HCP gambling task

To further test the generalizability of the BRS map we assessed the prediction performance on the HCP gambling task. This enabled us to test how well the BRS is able to predict on a much larger set of participants, from a different scanner, on a different task using a different experimental design (block vs event-related) and with different asymmetric levels of monetary outcomes. Using the significant voxels from

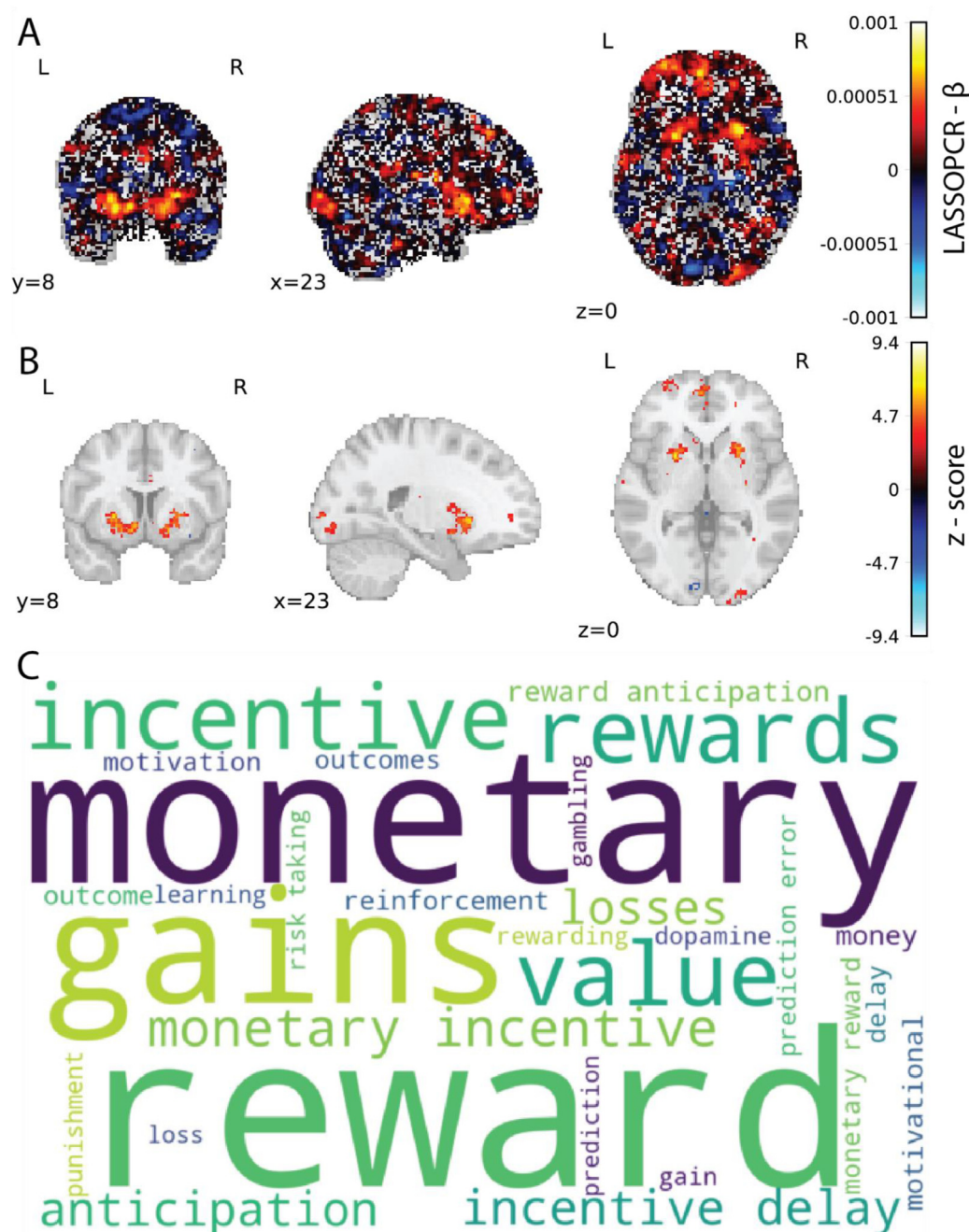


Fig. 2. A) Mean weights for the out-of sample prediction on the MID_{train}. B) Voxels significantly contributing to the out-of-sample prediction identified using the bootstrap procedure ($p < 0.001$). C) Word cloud showing the top 50 relevant terms (excluding anatomical terms) for the meta-analytic decoding of the BRS map. The size of the font was scaled by correlation strength ($r_{\min} = 0.11$, $r_{\max} = 0.22$).

the BRS map shown in Fig. 2B, we observed a significant prediction of the monetary outcomes on the HCP gambling task ($RMSE = 0.7$, $p_{\text{percm}} < 0.001$, $r = 0.21$, $p_{\text{percm}} < 0.001$, $BF_{10} > 1000$; Fig. 3C). To test the robustness of this finding the prediction was also repeated using all voxels, and FDR-corrected map ($p < 0.05$) and a map derived from first selecting the most consistent voxels and correcting using FDR (see Methods). The robustness checks revealed very similar significant predictions on the HCP gambling task (see Appendix 2). For the forced choice analysis we observed significant classification accuracies for all tests. However, as for the MID tasks, the classification accuracy was substantially higher between rewarding and loss trials and neutral and loss trials than between reward and neutral trials (see Table 5). The analysis of the test-retest reliability revealed that there is a significant correlation between the

patterns responses of the first and the second run of the HCP ($r_{\text{pearson}} = 0.24$, $p_{\text{perm}} < 0.001$; $r_{\text{spearman}} = 0.23$, $p_{\text{perm}} < 0.001$; $r_{\text{ICC}} = 0.24$, $p_{\text{perm}} < 0.001$; $\text{BF}_{10} > 1000$).

3.5. Testing the specificity on the DDT

In order to evaluate the specificity of the *BRS* map we assessed the prediction performance on the outcome phase of the DDT, in which participants see disgusting or neutral images. This enabled us to investigate whether the *BRS* map predicts differences in emotional salience more generally or whether it more specifically captures differences in reward. Using the significant voxels from the *BRS* map (see Fig. 2 middle) we did not observe a significant prediction of the differences in outcomes

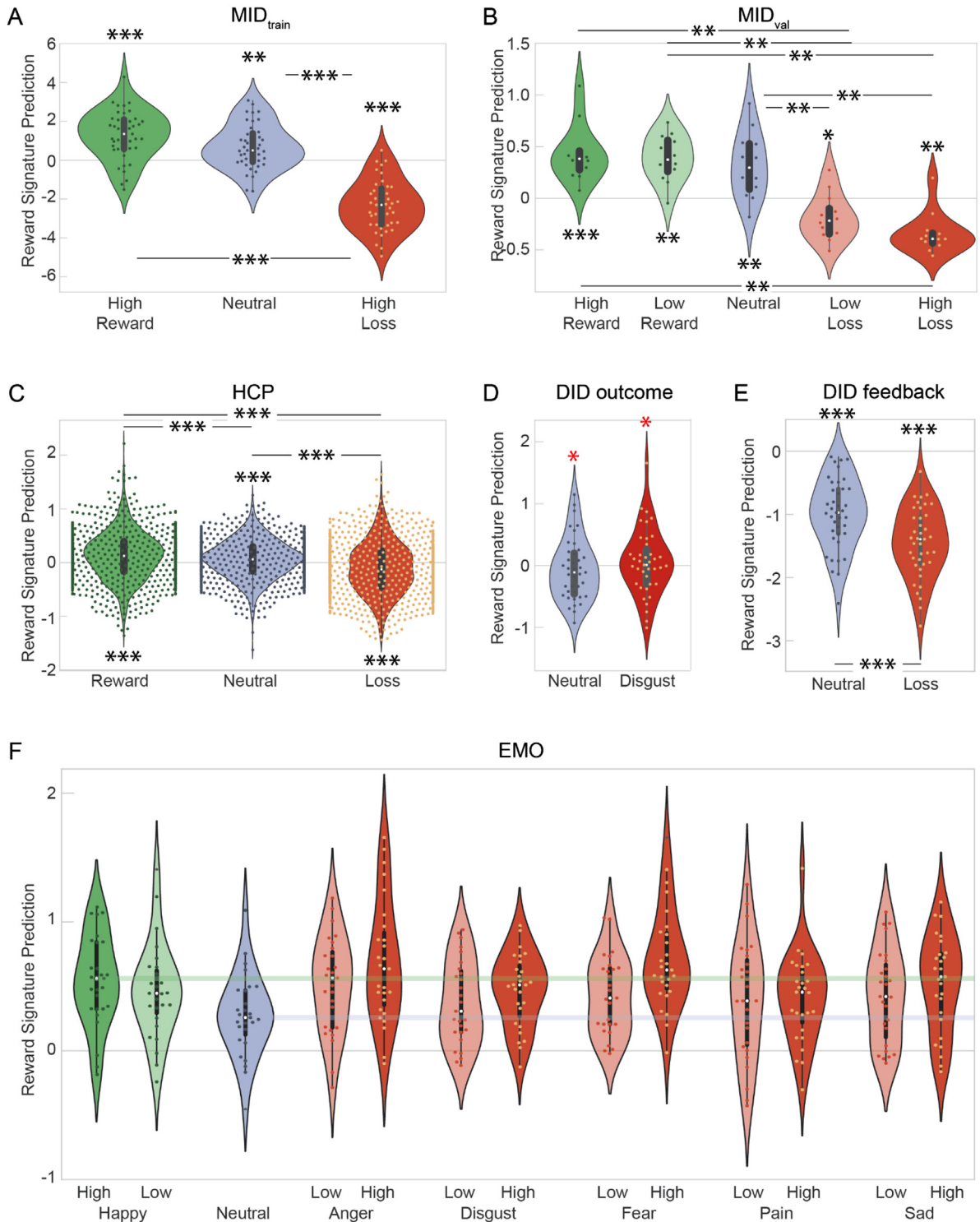


Fig. 3. A) Violinplot for the predicted monetary outcomes across conditions in the MID_{train}. B) Violinplot for the predicted monetary outcomes across conditions in the MID_{val}. C) Violinplot for the predicted monetary outcomes across conditions in the HCP gambling task. Because the HCP data contained 1084 subjects only 30% of actual data points could be plotted. The darker dots at the edges are due to several points overlapping with each other. D) Violinplot for the predicted disgusting versus neutral outcomes in the DDT for the outcome phase. E) Violinplot for the predicted positive versus negative feedback in the DDT for the feedback phase. F) Violinplot for the BRS values for the EVT. The gray and green thick horizontal lines show the median of the High Happy and Neutral condition for visual comparison against the other conditions. Note that not to overload the panel, statistical comparisons are omitted, but are mentioned in the text and in Table 6. For all panels: In the violin plots the circles represent individual observations arranged so that they do not overlap. Within the box and whisker plots, the white point represents the median, the box represents the lower and upper quartiles, the whiskers represent the 1.5 interquartile range. DID = Disgust Incentive Delay Task; MID = Monetary Incentive Delay Task; HCP = Human Connectome Project Gambling Task. *:BF₁₀ > 3; **:BF₁₀ > 10; ***:BF₁₀ > 100; red*:BF₁₀ < 0.33. Note that we use BF₁₀ values rather than p-values for the stars in the Fig. to provide evidence for the null or alternative hypothesis. For A-E, the stars above the violin represent the BF obtained from one-sample t-tests against zero, whereas the stars above the bars between violins represent the BFs obtained from Wilcoxon rank-sum tests comparing predictions. For F, all loadings were significantly different from zero (all BF₁₀>229, all $p<0.001$), and comparisons across conditions are detailed in Table 6, using forced choice statistics.

Table 4
Forced choice accuracies (%) for the MID_{val}.

Column1	HR	LR	N	LL	HL
HR	-				
LR	58	-			
N	58	50	-		
LL	92**	92**	92**	-	
HL	92**	92**	100**	75	-

HR = High Reward, LR = Low Reward, N = Neutral; LL = Low Loss; HL = High Loss; * = $p_{\text{perm}} < 0.05$; *** = $p_{\text{perm}} < 0.001$; ** = $p_{\text{perm}} < 0.01$.

Table 5
Forced choice accuracies (%) for the HCP.

	HR	N	HL
HR	-		
N	53**	-	
HL	73***	63***	-

HR = High Reward, N = Neutral; HL = High Loss; * = $p_{\text{perm}} < 0.05$; ** = $p_{\text{perm}} < 0.01$; *** = $p_{\text{perm}} < 0.001$.

in the DDT, and most importantly, found Bayesian evidence for the absence of such differentiation ($RMSE = 0.9$, $p_{\text{perm}} = 0.84$, $r = -0.13$, $p_{\text{perm}} = 0.28$, $BF_{10} = 0.23$; Fig. 3D). To test the robustness of this finding the prediction was also repeated using all voxels, and FDR-corrected map ($p < 0.05$), and a map derived from first selecting the most consistent voxels and correcting using FDR (see Methods). The robustness checks did not reveal any significant prediction on the DDT either (see Appendix 2). For the forced-choice analysis we found that the neutral trials could not be significantly distinguished from disgusting trials in the outcome phase (33%, $p = 0.98$).

To further assess the specificity of the BRS we also tested the feedback phase of the DDT (see Fig. 1D), in which participants are informed whether they successfully performed the task or not. Using the significant voxels from the BRS map shown in Fig. 2B, we found a significant prediction of feedback in the DDT ($RMSE = 0.92$, $p_{\text{perm}} < 0.001$, $r = 0.38$, $p_{\text{perm}} < 0.001$, $BF_{10} > 1000$; Fig. 3E). To test the robustness of this finding the prediction was also repeated using all voxels, an FDR-corrected map ($p < 0.05$) and a map derived from first selecting the most consistent voxels and correcting using FDR (see Methods). The robustness checks revealed very similar significant predictions on the feedback phase of the DDT (see Appendix 2). The forced-choice analysis revealed that the successful trials could be significantly discriminated from unsuccessful trials in the feedback phase (92%, $p < 0.001$). Notably, both distributions in the feedback phase are significantly below zero (see Fig. 3E). This may be due to the fact that the BRS is developed to maximize relative predictive performance. As a consequence, the predictive values for the DDT feedback are all negative because relative to receiving a monetary reward and succeeding at a given trial, just succeeding at a given trial is experienced as less rewarding.

To test whether the neural activation elicited by the outcome phase (when viewing the pictures) may have a bleed-over effect on the motivational delay period and to assess whether the results presented above are robust to different modeling choices we ran a robustness check with average activity as baseline. Results were replicated for both the outcome and the feedback phase. As before, for the outcome phase no significant prediction of differences in outcomes were observed ($RMSE = 1.03$, $p_{\text{perm}} = 0.94$, $r = -0.14$, $p_{\text{perm}} = 0.22$, $BF_{10} = 0.42$). For the feedback phase, we again found a significant prediction of feedback in the DDT ($RMSE = 0.78$, $p_{\text{perm}} < 0.001$, $r = 0.47$, $p_{\text{perm}} < 0.001$, $BF_{10} > 1000$).

3.6. Testing the specificity and generalizability on the EVT

In order to further characterize the BRS, we investigated how the brain response to watching facial expressions of other people's emotions would load on the BRS. Visual inspection of Fig. 3F illustrates that overall, witnessing emotional facial expressions, be they positive, negative or neutral, leads to positive loading on the BRS (one sample t-test against zero, all $BF_{10} > 229$, all $p < 0.001$, $t > 4.5$) with loading higher for the more intense expressions (two intensity $\times$ 6 emotion repeated ANOVA: main effect intensity: $F(1,26) = 17.53$, $p = 0.0003$, $BF_{\text{incl}} = 6.99$). The forced-choice analysis (Table 6) confirms that against the neutral stimulus (NT) all high-intensity emotions, be they positive (HH) or negative (SH, AH, FH, DH) generated significant discrimination performance, except for Pain (PH > NT in 59% of cases, n.s.); and that for all negative emotions, except for Pain, the low intensity video led to lower values on the BRS than the high intensity video. Finally, the High Happy video did not lead to BRS loading that could be discriminated from that of High Sadness, High Anger or High Disgust. This positive loading was expected for the Happy condition and also for the Neutral conditions as we observed the same result in the MID tasks and in the HCP. However, positive loading for all negative emotions, and more positive loading than for the neutral facial expression is perhaps less expected. Research on a phenomenon called morbid curiosity has however shown that people actually choose to view negatively valenced images over neutral images, and that doing so was associated with activation in reward related brain regions (Oosterwijk, 2017, Oosterwijk et al., 2020). This preference for negatively valenced material is thought to arise from a motivation to approach informative stimuli, with negative material having a higher information content than neutral material. To test whether our loading on the BRS may reflect a similar process, we performed two additional analyses. First, we compared the topology of our BRS with the activation pattern found by Oosterwijk and colleagues (Appendix 4). This analysis revealed significant similarities (i.e., positive correlation between the two maps: $r = 0.25$, $p < 0.001$). Second, we asked a different set of participants, recruited online, to rate our high intensity videos and our neutral videos on how interesting they found them. Fig. 4 shows that our happy, angry, disgusted and fearful facial expressions were indeed reported to be more interesting than our neutral facial expressions (all $p < 0.001$, all $BF_{10} > 46$), which could help explain why watching them may involve information-approaching-related processes that are similar-enough to the reward signals that we trained the BRS to capture. The only incongruence we find is that viewing the sad facial expressions did load significantly more than the neutral faces on the BRS, while online participants failed to find the sad facial expressions more interesting.

3.7. Using Neurosynth masks related to monetary outcomes for feature selection

To compare our data-driven feature selection approach to a more theory driven feature selection approach we also used two Neurosynth maps ((Yarkoni et al., 2011); see Table 8) related to monetary outcomes for feature selection within the cross-validation loop. Specifically, we used a meta-analytic map created based on the term *monetary reward* (Association test, FDR corrected for multiple comparisons at $p < 0.01$) and on the term *outcome* (Association test, FDR corrected for multiple comparisons at $p < 0.01$). A similar pattern of results as for the data-driven feature selection approach reported in the main text was found. Again the BRS significantly predicted monetary outcomes in the MID_{val} and the HCP gambling task, but did not significantly predict outcomes in the DDT. Performance on the HCP was slightly higher, whereas performance on the MID_{val} was slightly lower, which was expected as the data-driven feature selection was trained on another version of the MID task, and consequently was more likely to perform higher on a similar task. In contrast, the theory-driven approach was

Table 6
Forced choice accuracies (%) for the EVT.

	SH	AH	FH	DH	PH	AL	DL	SL	FL	PL	NT	HL	HH
SH	-	41	44	59	63	44	70***	67*	59	59	81***	64	44
AH	59	-	44	70**	74**	70**	85**	67*	74**	74**	85***	63	59
FH	56	56	-	74***	78**	70**	89***	74**	78***	78***	93***	74***	67*
DH	41	30***	26***	-	59	37	67	56	56	56	74***	48	37
PH	37	26***	22***	41	-	52	59	44	52	44	59	48	30***
AL	56	30***	30***	63	48	-	63	67*	63	70***	67*	63	41
DL	30***	15***	11***	33*	41	37	-	44	44	48	59	52	33*
SL	33*	33*	26***	44	56	33*	56	-	41	48	59	41	41
FL	41	26***	22***	44	48	37	56	59	-	56	56	48	37
PL	41	26***	22***	44	56	30***	52	52	44	-	52	37	48
NT	19***	15***	7***	26***	41	33*	41	41	44	48	-	30***	26***
HL	37	37	26***	52	52	37	48	59	52	63	70***	-	44
HH	56	41	33*	63	70***	59	67*	59	63	52	74***	56	-

Numbers indicate the proportion of subjects in which the condition mentioned over the row is numerically larger than the one indicated in the column. The p values were assessed using a label permutation statistics. SH = Sad High, AH = Anger High, FH = Fear High, DH = Disgust High, PH = Pain High, AL = Anger low, DL = Disgust Low, SL = Sad Low, FL = Fear Low, PL = Pain Low, NT = Neutral; HL = Happy Low, HH Happy High; * = $p_{\text{perm}} < 0.05$; ** = $p_{\text{perm}} < 0.01$; *** = $p_{\text{perm}} < 0.001$.

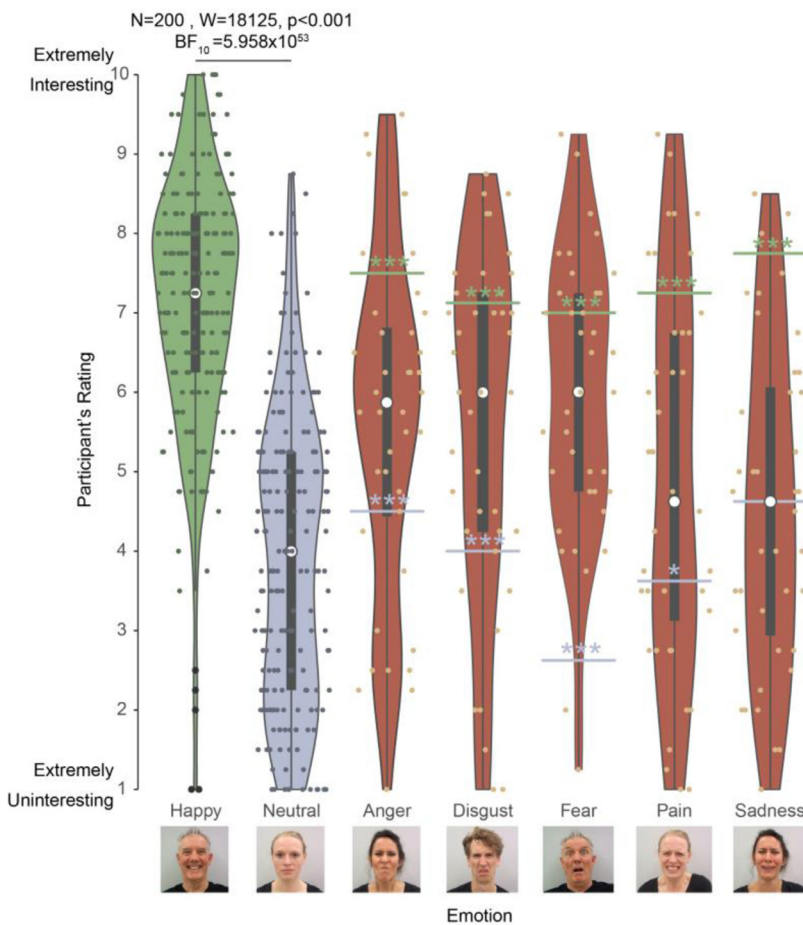


Fig. 4. Violinplot for the ratings across positive, neutral and negative emotions in the emotion rating task. Circles represent individual observations arranged so that they do not overlap. The box plot includes a white point for the median, a box for the quartiles, and whiskers for the max and min. Since the data was collected in five batches of participants (each containing videos of one negative emotion while always containing the same Happy and Neutral videos) we pooled the observations for Happy and Neutral videos for visualization purposes only in the leftmost violins. As a consequence, there are less observations for the negative emotions than for the Happy and Neutral videos. A Bayesian repeated ANOVA with Happy and Neutral video ratings as within factor, and batch as between, showed that the difference between Happy and Neutral did not change across batches (BF_{incl} main effect of video = 5.66×10^{13} , BF_{incl} main effect of batch = 0.278; BF_{incl} interaction = 0.209). Green and gray horizontal lines represent the median value of the happy and neutral facial expressions for that pool of participants, respectively. The black dots on the happy violin represent outliers, values above $Q3 + 1.5 \times \text{IQR}$ or below $Q1 - 1.5 \times \text{IQR}$. For the negative emotions, stars over the green or gray lines represent the significance of a paired comparison of the rating for the negative emotion against those for the happy or neutral facial expressions, as detailed in Table 7.

more task independent and more likely to perform similarly well across tasks.

4. Discussion

In the current study, we developed a multivariate brain model, the BRS, to allow us to decode the *relative* degree of reward across conditions in active decisions tasks. In particular, using the correlation between actual and decoded reward in the MID and HPC gambling task, we show the ability of the BRS to explain a significant proportion of the variance in the reward magnitude involved. This BRS is not only able to

predict variance in the monetary outcome in unseen subjects from the same sample but also generalizes to different samples using a different version of the same task and also to entirely different tasks. Further, this signature was found to not only predict monetary outcomes, but also rewarding outcomes in the form of positive versus negative feedback more generally. Relatedly, the BRS was found to load positively on prediction errors for money and for avoiding painful stimuli to others (a negative reinforcer) in a study currently under review elsewhere (but see Reply to Reviewer) investigating learning under moral conflict (Fornari et al., 2022). Importantly, the BRS values were appropriately signed, with wit-

Table 7
Comparison of interest ratings across conditions.

	Anger	Disgust	Fear	Pain	Sad
Vs Neutral	$t_{(39)} = 3.724$ $p = 6.192e-04$ $BF_{10} = 46.843$	$t_{(39)} = 3.811$ $p = 4.168e-05$ $BF_{10} = 58.945$	$t_{(39)} = 7.209$ $p = 1.11e-08$ $BF_{10} = 1.165e+6$	$t_{(39)} = 2.186$ $p = 0.035$ $BF_{10} = 1.433$	$t_{(39)} = 0.190$ $p = 0.850$ $BF_{10} = 0.173$
Vs Happy	$t_{(39)} = -4.111$ $p = 1.96e-04$ $BF_{10} = 131.922$	$W = 64.5$ $p = 1.522e-03$ $BF_{10} = 890.132$	$t_{(39)} = -3.46$ $p = 1.323e-03$ $BF_{10} = 23.811$	$t_{(39)} = -5.304$ $p = 4.772e-06$ $BF_{10} = 3982.874$	$t_{(39)} = -7.302$ $p = 8.28e-09$ $BF_{10} = 1.536e+6$

Table 8
Neurosynth maps for monetary reward and outcome.

Network	Studies	Date of	Link to download
Monetary Reward	97	04.10.2021	https://neurosynth.org/analyses/terms/monetary%20reward/
Outcome	385	04.10.2021	https://neurosynth.org/analyses/terms/outcome/

nessing shocks loading negatively on the BRS and receiving money, positively, which further speaks to the generalizability of its sensitivity to reinforcers as outcomes in decision-making tasks. With regards to such outcomes, this BRS was found to be specific to rewarding outcomes and did not generalize to emotionally salient (disgusting) images (when are the result of a decision). However, when passively viewing facial expressions of other individuals, rather than outcomes arising from the participant's own actions, viewing all facial expressions yielded positive BRS values, with almost all of them being larger than the neutral facial expressions. This suggests that the BRS, when used in context in which participants do not need to make choices, may capture a wider set of processes that future experiments will need to further characterize.

To create the BRS that is sensitive to the neurocognitive underpinnings of reward processing, we trained a LASSOPCR model on the MID, which is the most consistently used task to evoke the neural mechanisms associated with processing monetary outcomes (Oldham et al., 2018). To ensure that the BRS predicts reward specifically and not salience in general, we only selected voxels for predictions that correlated more strongly with outcomes (i.e., voxels that differentiate between reward, neutral and loss outcomes), than with salience (i.e., voxels that differentiate only between neutral and consequential, reward or loss, outcomes). We found that clusters of voxels in the bilateral dorsal striatum, the vmPFC and the right occipital pole significantly decoded monetary outcomes in novel participants from the same sample. We subsequently tested whether the observed clusters indeed reflect reward processing areas by means of using the Neurosynth (Yarkoni et al., 2011) decoder. This decoder compared our BRS to the entire set of terms included in the Neurosynth database and found that the highest ranked associations were *reward* and *monetary*, providing converging evidence that the BRS predicts rewarding outcomes.

The finding that activation patterns in the dorsal striatum are predictive of rewarding outcomes aligns well with previous fMRI studies that found that the striatum encodes prediction error signals (Diekhof et al., 2012, Galtress et al., 2012, Haber and Knutson, 2010, O'Doherty et al., 2004). The striatum has been consistently linked to both the anticipation and evaluation of rewarding outcomes (Oldham et al., 2018). In addition, abnormal activity in the striatum and connectivity between the striatum and the limbic system have been linked to impaired reward processing in obesity and bipolar disorder (Caseras et al., 2013, Nummenmaa et al., 2012, Yip et al., 2015). Similarly, the observation that a cluster of voxels in the vmPFC is predictive of rewarding outcomes is in accordance with previous fMRI research on economic decisions and reward processing, as it has been associated consistently with the receipt of reward or loss and the computation of subjective value (Bartra et al., 2013, Diekhof et al., 2012, Haber and Knutson, 2010, Kringelbach, 2004, Levy and Glimcher, 2012, Peters and Büchel, 2010, Sescousse et al., 2013). It is relevant to note that while we found reward

to be positively associated in our BRS, this does not preclude the existence of circuits and ensembles that encode loss and aversive processes and conversely exhibit decreased activation in response to reward.

As a next step, we tested the generalizability of the BRS on two different samples. Firstly, we tested the relative predictive accuracy of the BRS on a different version of the MID, with five levels of monetary outcomes instead of three, from a different sample and found that we could again decode monetary outcomes significantly with high accuracy, as assessed using the correlation between decoded and actual reward magnitude. Secondly, we assessed the predictive performance of the BRS on a large sample ($N = 1084$) with a different task, namely a gambling task from the Human Connectome Project. Again, we found that the BRS was able to significantly predict monetary outcomes. Together, these results highlight the generalizability of the predictions of the BRS. The observation that predictive accuracy dropped in comparison to the other two samples can be explained by the fact that this task differed from the MID task in two ways: In contrast to the MID, the gambling included rewards that were not symmetrically distributed around zero. In addition, the gambling task was developed for analysis using a block design (averaging over several trials of the same condition) whereas the MID used an event related design (modeling specific phases within a trial individually).

While our feature selection procedure, which removed voxels that primarily responded to salience, and training the BRS on a well-established reward processing task provided a good fundament for ensuring the specificity of predictions, we also wanted to empirically test this specificity. To this end, we also evaluated the predictions of the BRS on two phases of the DDT, a novel task designed to evoke disgust as a negative outcome. First, we tested the outcome phase to test whether predictions are specific to reward or generalize to other emotionally salient outcomes such as disgust. The analysis provided evidence in favor of the absence of an effect. Stated differently, the BRS generated predictions that did not differ between participants viewing a disgusting image or a neutral image. Second, we tested predictions during the feedback phase which provided a success/failure feedback to the participants, and could therefore be triggering neurocognitive processes associated with reward/loss that are non-monetary in nature. Here we found a significant predictive performance of the BRS, suggesting that the BRS decodes reward and loss processing more generally and is not limited to monetary outcomes alone. This was also confirmed by the signed loading of prediction errors for shocks on the BRS in the moral conflict task, where less intense shock than expected, a negative reinforcer, led to positive loading on the BRS. These two findings suggest that the BRS predicts reinforcing outcomes, be they positive (financial or otherwise) or negative reinforcers (withholding a shock) with some specificity that does not generalize to the other emotionally salient outcome (disgust in the DDT). Importantly, this specificity for reinforcers

holds in the context of tasks in which the outcome is the result of a participant's decision.

To perform a first step along a necessarily long route to characterize the BRS in task structures that differ fundamentally from those in which it was trained, by identifying the boundaries of the stimuli that may load onto it, we used the EVT, in which participants witness the anger, disgust, fear, happiness, pain and sadness of others. Importantly, in the EVT, the stimuli were not a reward or punishment contingent on the participant's actions, acting as reinforcers, but stimuli presented in random sequence independently of the participants performance. Based on the literature on empathy, we might have expected that viewing positively valenced facial expressions (happy) may have triggered positively valenced feelings in the viewer and hence positive BRS values. This prediction was confirmed. Empathy may however also predict that viewing negatively valenced facial expressions (angry, fearful, disgusted, sad, painful) may have triggered negatively valenced feelings and hence negative BRS values. This prediction was not confirmed, as viewing most of the negatively valenced facial expressions yielded positive BRS values. In contrast, the literature on morbid curiosity predicts that participants are motivated to approach information independently of its valence. (Bennett et al., 2016) found that humans intrinsically value information in a way that is inconsistent with normative accounts of decision-making under uncertainty and are willing to incur considerable monetary costs to obtain information even if it is irrelevant to the task at hand. Oosterwijk and colleagues (Oosterwijk, 2017) have shown that participants actively approach negatively valenced material rather than neutral materials, with participants specifically choosing to approach social negative information over nonsocial negative information and Kashdan and Silva (Kashdan and Silvia, 2009) have found that there are significant long term benefits to be reaped from being curious about the emotions of others. Given this established motivation to approach even negatively valenced informative social stimuli, and that our negative emotional facial expressions were rated as interesting, that viewing all but the painful facial expression yielded positive loadings on the BRS may thus capture the satisfaction of this information-seeking motivation - a signal similar enough to reward-signals to be captured by a signature trained to capture reward vs. loss processing. Neuroimaging evidence supporting the neural similarity of morbid curiosity and reward-signals stems from a recent fMRI study that showed that choosing negative stimuli is associated with the activation of reward related structures (Oosterwijk et al., 2020, Scrivner, 2021). (Oosterwijk et al., 2020)) used the same Neurosynth decoder approach that we employed and found that the top 3 terms associated with their activation when participants choose to view negatively valenced materials, were 'reward', 'task' and 'monetary', which is almost identical with our results. To further explore that similarity, we correlated our BRS with Oosterwijk and colleagues' whole brain map for morbid curiosity and found a significant positive correlation, further supporting a significant similarity between the neural processes related to reward processing and morbid curiosity (see Appendix 4). In the light of these considerations, our initially counterintuitive finding that the BRS yields positive values for viewing positive and negative facial expressions, is perhaps less surprising. The information content of these stimuli may have satisfied a motivation to seek social information, triggering a signal reinforcing the processing and approach of these stimuli that resemble that involved in reward and loss sufficiently, to be captured by our signature trained to quantify the latter. Indeed, the box office successes of movies that showcase strong negative emotions in its protagonists are perhaps a token to the reinforcing value of even negative facial expressions.

Negative and positive facial expressions are also more visually salient than the neutral ones, inviting us to consider the possibility that the BRS may simply capture such salience. That both high gain and high losses are salient in the MID tasks, but load in opposite directions onto the BRS speaks against such an unsigned salience signal. Accordingly, our data suggests a more nuanced working hypothesis that the BRS captures a signed reward vs. punishment signal when participants receive

reinforcers contingent on their actions (be they financial or otherwise) and information seeking signals when they are exposed to stimuli that are not contingent on their actions. As mentioned in the introduction, characterizing the specificity of a neural signature is an ongoing process that will continue to develop as a signature is applied to a wider array of paradigms, and testing the working hypothesis of a signal that may reinforce the exploration of informative stimuli will require further studies that contrast stimuli that participants will voluntarily approach for their information content from those they will voluntarily avoid, with the prediction that the former should produce positive and the latter negative BRS loadings.

Since previous research suggested that reward may be encoded specifically in the striatum ((Haber and Knutson, 2010, Knutson et al., 2001), 2005), we also tested whether a broader circuit (i.e. including the vMPFC) is needed to decode reward. To this end we applied a more theory driven approach where we used a meta-analytic map created based on the term *monetary reward* and on the term *outcome*. These maps only included voxels in the striatum (and not in the vMPFC). Similar to the data-driven feature selection approach reported in the main text, the BRS significantly predicted monetary outcomes in the MID_{val} and the HCP gambling task, and DDT feedback phase, but did not significantly predict outcomes in the DDT outcome phase. Performance on the HCP and DDT feedback phase was higher for this theory driven approach as compared to the data driven BRS. In contrast, performance on the MID_{val} was slightly lower for the theory driven approach, which was expected as the data-driven BRS involved feature selection trained on a version of the MID task similar in nature and consequently was more likely to perform higher on a similar task. In contrast, the theory-driven approach was more task independent and more likely to perform similarly well across tasks. This aligns well with the notion that the striatum may encode reward and losses quite generally within the decision making framework.

To test the robustness of our findings, we also repeated the reported analyses in the main text, using different thresholds for the feature selection procedure and for the correction for multiple comparison (see Appendix 2). These robustness checks validated the findings from the main text. For all feature selection and multiple comparison correction thresholds, the predictions within the MID_{train}, MID_{val}, gambling task and feedback phase of the DDT remained significant. Only on the DDT outcome phase (testing for specificity for reward processing), when using all voxels instead of selecting only voxels that were significant in predicting monetary outcomes on the MID_{train}, there was not enough evidence to support the hypothesis that the BRS was unable to differentiate between disgusting and neutral images. This may be due to voxels contributing to the prediction that are not specific to predicting monetary outcome but also encode emotional salience in general. Since we used a lenient threshold for the feature selection algorithm some voxels coding for salience may have been included in the model and thus lowered the evidence in favor of the absence of an effect. This finding suggests that users should use the signature only including significant voxels when applying the BRS to other sets.

In future studies, this BRS could be employed to differentiate and compare the contribution of various emotions and cognitive processes to complex (social) decisions. For instance, we applied it to the case of moral decisions, and found that outcomes for others may enter decision making by mapping onto the BRS that was developed to capture first-person reinforcers (Fornari et al., under revision). The BRS could be applied in combination with neural signatures for vicarious pain ((Caspar et al., 2020, Krishnan et al., 2016, Zhou et al., 2020) (Caspar et al., 2020) and for guilt (Yu et al., 2020) to capture the contribution of first-person reward circuitry in social decision-making.

A limitation of our signature to consider when interpreting applications of our BRS, is that while its expression values correlate with the reward outcome obtained by participants in the MID and gambling task, the BRS failed to identify neutral outcomes as such. Specifically, in the MID tasks (MID_{train} and MID_{val}), while we correctly find the

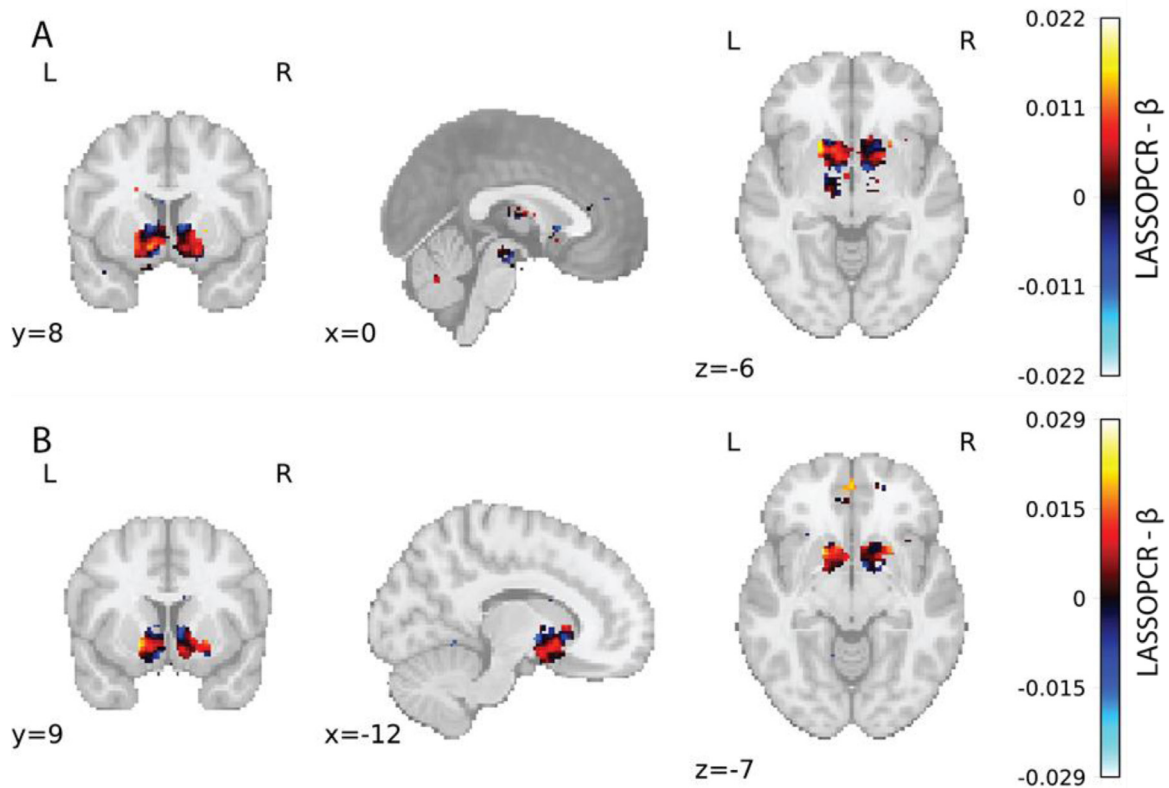


Fig. 5. A: Prediction weights derived from the feature-selection approach based on the *monetary reward* meta-analytic map. B: Prediction weights derived from the feature-selection approach based on the *outcome* meta-analytic map.

gain conditions to generate values significantly above zero, and the loss conditions to generate negative values, the neutral conditions generate values that are also slightly positive (Fig. 3A,B). In addition, the BRS failed to discriminate high and low reward conditions in the MID_{val} and also did not differentiate the two rewarding conditions from the neutral condition accurately. Conceptually, this may be explained by the observation that people generally seem to be loss averse (i.e., losses loom larger than gains; (Kahneman, 2011)). Thus, a loss of the same (monetary) value as a reward will be experienced as more severe and may thus be encoded as more distant from zero (the neutral condition) than a reward of an equal amount. This would explain why our algorithm was not always able to significantly discriminate rewarding from neutral trials, but always achieved significant discrimination between neutral and loss trials. Methodologically, this observation may be explained by the fact that when training a linear model on a dependent variable with only three levels the model will be mostly influenced by its extreme points, whereas the middle point will be less influential in determining parameter estimates. Further, our feature selection algorithm was designed to maximize *relative* prediction performance rather than absolute prediction, because value based computations and associated outcome processing have been found to be context dependent (Bateson et al., 2003, Huber et al., 1982, Louie et al., 2013, Shafir et al., 2002, Simonson, 1989) and that decisions do not reflect absolute valuations assigned to individual alternatives. This however means that our signature should not be applied to the z-values of a single condition to determine if any reward processing was triggered, but rather on multiple conditions to test whether they differ in reward processing.

Another limitation pertains to the fact that several constructs related to reward processing have been associated with the striatum and vmPFC contained in our BRS, such as the outcome value, anticipated outcome, goal value and prediction error (Diekhof et al., 2012, Galtres et al., 2012, Haber and Knutson, 2010, Knutson et al., 2005, O'Doherty et al.,

2004, Rutledge et al., 2010) and we can't precisely disentangle which of these processes are captured by our signature. The positive loading of all our facial expressions on the BRS further illustrates this point, negative facial expressions may not be intuitively considered to be rewarding, even if an emerging literature shows that people decide to approach and explore them. This invites us to refine our understanding of the function of these networks. Our understanding of the relationship between brain and cognition may benefit by engaging in such iterative multivariate cycles in which one can attempt to tease apart the neural signatures of cognitive constructs that we assume to be distinct, and vice versa, challenge our understanding of the nature of these constructs by exploring what activates the signatures that are meant to isolate them.

In summary, we created a BRS to predict monetary outcomes across decision tasks and several large samples. Within an experimental framework in which outcomes are contingent on our participants' actions, the BRS appears to perform that aim well: it is specific to rewarding outcomes (in the sense of positive and negative reinforcers) and does not seem to generalize to disgusting outcomes. However, it is important to be aware that outside of this action-outcome framework, such specificity for rewards is not supported. That passively viewing facial expressions, including many of negative valence, yields to positive BRS values indicates that it may also capture less explored signals reinforcing the exploration of informative stimuli that may satisfy morbid curiosity. These curiosity-related signals may share enough neural substrates with reward processing to be captured by our BRS that was trained to capture these latter reward processes (Oosterwijk et al., 2020).

The benefit of brain signature over the univariate approach is that it integrates distributed information from regions across the whole brain into a single optimized prediction which can then be tested across conditions on new and independent individuals and samples. As a consequence, this approach often circumvents the need for multiple comparisons and provides unbiased estimates of effect size (Reddan et al.,

Table 9
Prediction performance for feature selection based on neurosynth maps.

Analysis type	Monetary reward association	Outcome association
MID _{train} CV	0.63 (0.7 %, $p_{\text{corr}} < 0.001$, RMSE = 3.23, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)	0.68 (0.4 %; $p_{\text{corr}} < 0.001$, RMSE = 2.97, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)
HCP Gambling	0.26 ($p_{\text{corr}} < 0.001$, RMSE = 2.01, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)	0.3 ($p_{\text{corr}} < 0.001$, RMSE = 2.67, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)
DDT	-0.09 ($p_{\text{corr}} = \text{n.s.}$, RMSE = 2.39, $p_{\text{RMSE}} < \text{n.s.}$, $\text{BF}_{10} = 0.23$)	-0.02 ($p_{\text{corr}} = \text{n.s.}$, RMSE = 2.51, $p_{\text{RMSE}} < \text{n.s.}$, $\text{BF}_{10} = 0.15$)
DDT Feedback	0.47 ($p_{\text{corr}} < 0.001$, RMSE = 6.08, $p_{\text{RMSE}} < 0.001$, $\text{BF} > 1000$)	0.49 ($p_{\text{corr}} < 0.001$, RMSE = 6.17, $p_{\text{RMSE}} < 0.001$, $\text{BF} > 1000$)
MID _{val}	0.71 ($p_{\text{corr}} < 0.001$, RMSE = 2.85, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)	0.69 ($p_{\text{corr}} < 0.001$, RMSE = 2.89, $p_{\text{RMSE}} < 0.001$, $\text{BF}_{10} > 1000$)

CV = cross-validation; BF_{10} = Bayes Factor for evidence in favor of the alternative hypothesis; HCP = HCP gambling task; Percentage in the first row represents the number of voxels selected via feature selection for the different thresholds.

2017), making the signature approach more sensitive, generalizable and reproducible than traditional univariate approaches (Kragel et al., 2018). On the other hand, our results highlight that applying signatures outside of the bounds of the paradigms they have been trained on, entails the risk of ignoring the fact that the brain can utilize particular networks for rather different functions across different situations.

Credit Author Statement

Author Contributions: S.P.H.S., C.K. & V.G. conceived the BRS extraction. S.P.H.S., A.S. and M.A.S.B. conceived of the MID/DDT study. C.J.S.T., J.C.B. and S.P.H.S. performed the experiments; S.P.H.S. and J.C.B. analyzed data with input from C.K. and V.G.; S.P.H.S., C.K. & V.G. wrote the paper with comments from A.S., M.A.S.B., C.J.S.T. and T.D.W..

Data availability

The unthresholded BRS can be found on neurovault: <https://neurovault.org/images/775976/>. The scripts and links to the data needed to reproduce the analyses reported in the paper can be found on OSF: DOI 10.17605/OSF.IO/YDFRS (<https://osf.io/ydf/rs/>).

Our ethics approval however does not allow us to share the 7T fMRI data of the EVT task.

Funding

European Union's Horizon 2020 research and innovation programme grant, ERC-StG 'HelpUS' 758703 (VG)

Dutch Research Council (NWO) VICI grant 453-15-009 (CK)

Declaration of Competing Interest

No competing interests.

Data Availability

Links to some of the data has already been made available the rest of the data and scripts will be made available once the paper is at the review stage.

Acknowledgements

We also gratefully acknowledge financial support from the Erasmus Research Institute of Management (ERIM).

Supplementary materials

Supplementary material associated with this article can be found, in the online version, at [doi:10.1016/j.neuroimage.2023.119990](https://doi.org/10.1016/j.neuroimage.2023.119990).

References

- Arbabshirani, M.R., Plis, S., Sui, J., Calhoun, V.D., 2017. Single subject prediction of brain disorders in neuroimaging: promises and pitfalls. *Neuroimage* 145, 137–165. doi:10.1016/j.neuroimage.2016.02.079.
- Ashar, Y.K., Andrews-Hanna, J.R., Dimidjian, S., Wager, T.D., 2017. Empathic care and distress: predictive brain markers and dissociable brain systems. *Neuron* 94 (6), 1263–1273. doi:10.1016/j.neuron.2017.05.014, e4.
- Balodis, I.M., Potenza, M.N., 2015. Anticipatory reward processing in addicted populations: a focus on the monetary incentive delay task. *Biol. Psychiatry* 77 (5), 434–444. doi:10.1016/j.biopsych.2014.08.020.
- Bartra, O., McGuire, J.T., Kable, J.W., 2013. The valuation system: a coordinate-based meta-analysis of BOLD fMRI experiments examining neural correlates of subjective value. *Neuroimage* 76, 412–427. doi:10.1016/j.neuroimage.2013.02.063.
- Bateson, M., Healy, S.D., Hurly, T.A., 2003. Context-dependent foraging decisions in rufous hummingbirds. In: *Proceedings of the Royal Society of London. Series B: Biological Sciences*, 270, pp. 1271–1276. doi:10.1098/rspb.2003.2365.
- Bennett, D., Bode, S., Brydevall, M., Warren, H., Murawski, C., 2016. Intrinsic valuation of information in decision making under uncertainty. *PLoS Comput. Biol.* 12 (7), e1005020. doi:10.1371/journal.pcbi.1005020.
- Blaimer, M., Choli, M., Jakob, P.M., Griswold, M.A., Breuer, F.A., 2013. Multiband phase-constrained parallel MRI: multiband phase-constrained parallel MRI. *Magn. Reson. Med.* 69 (4), 974–980. doi:10.1002/mrm.24685.
- Caballero-Gaudes, C., Reynolds, R.C., 2017. Methods for cleaning the BOLD fMRI signal. *Neuroimage* 154, 128–149. doi:10.1016/j.neuroimage.2016.12.018.
- Carandini, M., Heeger, D.J., 2012. Normalization as a canonical neural computation. *Nat. Rev. Neurosci.* 13 (1), 51–62. doi:10.1038/nrn3136.
- Caseras, X., Lawrence, N.S., Murphy, K., Wise, R.G., Phillips, M.L., 2013. Ventral striatum activity in response to reward: differences between bipolar I and II disorders. *Am. J. Psychiatry* 170 (5), 533–541. doi:10.1176/appi.ajp.2012.12020169.
- Caspar, E.A., Ioumpa, K., Keysers, C., Gazzola, V., 2020. Obeying orders reduces vicarious brain activation towards victims' pain. *Neuroimage* 222, 117251. doi:10.1016/j.neuroimage.2020.117251.
- Chang, L.J., Gianaros, P.J., Manuck, S.B., Krishnan, A., Wager, T.D., 2015. A sensitive and specific neural signature for picture-induced negative affect. *PLoS Biol.* 13 (6), e1002180. doi:10.1371/journal.pbio.1002180.
- Chib, V.S., Rangel, A., Shimojo, S., O'Doherty, J.P., 2009. Evidence for a common representation of decision values for dissimilar goods in human ventromedial prefrontal cortex. *J. Neurosci.* 29 (39), 12315–12320. doi:10.1523/JNEUROSCI.2575-09.2009.
- Clithero, J.A., Rangel, A., 2014. Informatic parcellation of the network involved in the computation of subjective value. *Soc. Cogn. Affect. Neurosci.* 9 (9), 1289–1302. doi:10.1093/scan/nst106.
- Cox, R.W., Hyde, J.S., 1997. Software tools for analysis and visualization of fMRI data. *NMR Biomed.* 10 (4–5), 171–178. doi:10.1002/(SICI)1099-1492(199706/08)10:4/5<171::AID-NBM453>3.0.CO;2-L.
- Delgado, M.R., Nystrom, L.E., Fissell, C., Noll, D.C., Fiez, J.A., 2000. Tracking the hemodynamic responses to reward and punishment in the striatum. *J. Neurophysiol.* 84 (6), 3072–3077. doi:10.1152/jn.2000.84.6.3072.
- Diekhof, E.K., Kaps, L., Falkai, P., Gruber, O., 2012. The role of the human ventral striatum and the medial orbitofrontal cortex in the representation of reward magnitude – An activation likelihood estimation meta-analysis of neuroimaging studies of passive reward expectancy and outcome processing. *Neuropsychologia* 50 (7), 1252–1266. doi:10.1016/j.neuropsychologia.2012.02.007.
- Esteban, O., Markiewicz, C.J., Blair, R.W., Moodie, C.A., Isik, A.I., Erramuzpe, A., Kent, J.D., Gonçalves, M., DuPre, E., Snyder, M., Oya, H., Ghosh, S.S., Wright, J., Duriez, J., Poldrack, R.A., Gorgolewski, K.J., 2019. fmrip: a robust preprocessing pipeline for functional MRI. *Nat. Methods* 16 (1), 111–116. doi:10.1038/s41592-018-0235-4.
- Esteban, O., Zosso, D., Daducci, A., Bach-Cuadra, M., Ledesma-Carbayo, M.J., Thiran, J.-P., Santos, A., 2016. Surface-driven registration method for the structure-informed segmentation of diffusion MR images. *Neuroimage* 139, 450–461. doi:10.1016/j.neuroimage.2016.05.011.
- Etzel, J.A., Valchev, N., Keysers, C., 2011. The impact of certain methodological choices on multivariate analysis of fMRI data with support vector machines. *Neuroimage* 54 (2), 1159–1167. doi:10.1016/j.neuroimage.2010.08.050.
- Fornari, L., Ioumpa, K., Nostro, A.D., Evans, N., De Angelis, L., Speer, S.P.H., Paracampo, R., Gallo, S., Spezio, M., Keysers, C., Gazzola, V., 2022. Neuro-computational mechanisms and individual biases in action-outcome learning under moral conflict. *Manuscript under rev.*.

- Galtress, T., Marshall, A.T., Kirkpatrick, K., 2012. Motivation and timing: clues for modelling the reward system. *Behav. Processes* 90 (1), 142–153. doi:10.1016/j.beproc.2012.02.014.
- Gazzola, V., Keysers, C., 2009. The observation and execution of actions share motor and somatosensory voxels in all tested subjects: single-subject analyses of unsmoothed fMRI data. *Cereb. Cortex* 19 (6), 1239–1255. doi:10.1093/cercor/bhn181.
- Glasser, M.F., Sotiropoulos, S.N., Wilson, J.A., Coalson, T.S., Fischl, B., Andersson, J.L., Xu, J., Jbabdi, S., Webster, M., Polimeni, J.R., Van Essen, D.C., Jenkinson, M., 2013. The minimal preprocessing pipelines for the human connectome project. *Neuroimage* 80, 105–124. doi:10.1016/j.neuroimage.2013.04.127.
- Gorgolewski, K., Burns, C.D., Madison, C., Clark, D., Halchenko, Y.O., Waskom, M.L., Ghosh, S.S., 2011. Nipype: a flexible, lightweight and extensible neuroimaging data processing framework in python. *Front. Neuroinform.* 5. doi:10.3389/fninf.2011.00013.
- Grosenick, L., Klingenberg, B., Katovich, K., Knutson, B., Taylor, J.E., 2013. Interpretable whole-brain prediction analysis with GraphNet. *Neuroimage* 72, 304–321. doi:10.1016/j.neuroimage.2012.12.062.
- Haber, S.N., Knutson, B., 2010. The reward circuit: linking primate anatomy and human imaging. *Neuropsychopharmacology* 35 (1), 4–26. doi:10.1038/npp.2009.129.
- Hare, T.A., O'Doherty, J., Camerer, C.F., Schultz, W., Rangel, A., 2008. Dissociating the role of the orbitofrontal cortex and the striatum in the computation of goal values and prediction errors. *J. Neurosci.* 28 (22), 5623–5630. doi:10.1523/JNEUROSCI.1309-08.2008.
- Haxby, J.V., Gobbini, M.I., Furey, M.L., Ishai, A., Schouten, J.L., Pietrini, P., 2001. Distributed and overlapping representations of faces and objects in ventral temporal cortex. *Science* 293 (5539), 2425–2430. doi:10.1126/science.1063736.
- Haynes, J.-D., Sakai, K., Rees, G., Gilbert, S., Frith, C., Passingham, R.E., 2007. Reading hidden intentions in the human brain. *Curr. Biol.* 17 (4), 323–328. doi:10.1016/j.cub.2006.11.072.
- Huber, J., Payne, J.W., Puto, C., 1982. Adding asymmetrically dominated alternatives: violations of regularity and the similarity hypothesis. *J. Consum. Res.* 9 (1), 90. doi:10.1086/208899.
- Huth, A.G., de Heer, W.A., Griffiths, T.L., Theunissen, F.E., Gallant, J.L., 2016. Natural speech reveals the semantic maps that tile human cerebral cortex. *Nature* 532 (7600), 453–458. doi:10.1038/nature17637.
- Huth, A.G., Nishimoto, S., Vu, A.T., Gallant, J.L., 2012. A continuous semantic space describes the representation of thousands of object and action categories across the human brain. *Neuron* 76 (6), 1210–1224. doi:10.1016/j.neuron.2012.10.014.
- Kahneman, D., 2011. *Thinking, Fast and Slow*. Allen Lane.
- Kashdan, T.B., Silvia, P.J., 2009. Curiosity and interest: the benefits of thriving on novelty and challenge. *Oxford handbook of positive psychol.* 2, 367–374 Apr 21.
- Keysers, C., Gazzola, V., Wagenmakers, E.-J., 2020. Using Bayes factor hypothesis testing in neuroscience to establish evidence of absence. *Nat. Neurosci.* 23 (7), 788–799. doi:10.1038/s41593-020-0660-4.
- Knutson, B., Fong, G.W., Adams, C.M., Varner, J.L., Hommer, D., 2001. Dissociation of reward anticipation and outcome with event-related fMRI. *Neuroreport* 12 (17), 3683–3687. doi:10.1097/00001756-200112040-00016.
- Knutson, B., Greer, S.M., 2008. Anticipatory affect: neural correlates and consequences for choice. *Philosophical Trans. Royal Society B: Biol. Sci.* 363 (1511), 3771–3786. doi:10.1098/rstb.2008.0155.
- Knutson, B., Taylor, J., Kaufman, M., Peterson, R., Glover, G., 2005. Distributed neural representation of expected value. *J. Neurosci.* 25, 4806–4812.
- Knutson, B., Westdorp, A., Kaiser, E., Hommer, D., 2000. fMRI visualization of brain activity during a monetary incentive delay task. *Neuroimage* 12 (1), 20–27. doi:10.1006/nimg.2000.0593.
- Kragel, P.A., Knodt, A.R., Hariri, A.R., LaBar, K.S., 2016. Decoding spontaneous emotional states in the human brain. *PLoS Biol.* 14 (9), e2000106. doi:10.1371/journal.pbio.2000106.
- Kragel, P.A., Koban, L., Barrett, L.F., Wager, T.D., 2018. Representation, pattern information, and brain signatures: from neurons to neuroimaging. *Neuron* 99 (2), 257–273. doi:10.1016/j.neuron.2018.06.009.
- Kragel, P.A., LaBar, K.S., 2015. Multivariate neural biomarkers of emotional states are categorically distinct. *Soc. Cogn. Affect. Neurosci.* 10 (11), 1437–1448. doi:10.1093/scan/nsv032.
- Kringelbach, M., 2004. The functional neuroanatomy of the human orbitofrontal cortex: evidence from neuroimaging and neuropsychology. *Prog. Neurobiol.* 72 (5), 341–372. doi:10.1016/j.pneurobio.2004.03.006.
- Krishnan, A., Woo, C.-W., Chang, L.J., Ruzic, L., Gu, X., López-Solà, M., Jackson, P.L., Pujol, J., Fan, J., Wager, T.D., 2016. Somatic and vicarious pain are represented by dissociable multivariate brain patterns. *Elife* 5. doi:10.7554/elife.15166.
- Levy, D.J., Glimcher, P.W., 2012. The root of all value: a neural common currency for choice. *Curr. Opin. Neurobiol.* 22 (6), 1027–1038. doi:10.1016/j.conb.2012.06.001.
- Lindquist, K.A., Barrett, L.F., 2012. A functional architecture of the human brain: emerging insights from the science of emotion. *Trends Cogn. Sci.* 16 (11), 533–540. doi:10.1016/j.tics.2012.09.005.
- Liu, X., Hairston, J., Schrier, M., Fan, J., 2011. Common and distinct networks underlying reward valence and processing stages: a meta-analysis of functional neuroimaging studies. *Neurosci. Biobehav. Rev.* 35 (5), 1219–1236. doi:10.1016/j.neubiorev.2010.12.012.
- Louie, K., Gratton, L.E., Glimcher, P.W., 2011. Reward value-based gain control: divisive normalization in parietal cortex. *J. Neurosci.* 31 (29), 10627–10639. doi:10.1523/JNEUROSCI.1237-11.2011.
- Louie, K., Khaw, M.W., Glimcher, P.W., 2013. Normalization is a general neural mechanism for context-dependent decision making. *Proc. Natl. Acad. Sci.* 110 (15), 6139–6144. doi:10.1073/pnas.1217854110.
- Louie, K., LoFaro, T., Webb, R., Glimcher, P.W., 2014. Dynamic Divisive Normalization Predicts Time-Varying Value Coding in Decision-Related Circuits. *J. Neurosci.* 34 (48), 16046–16057. doi:10.1523/JNEUROSCI.2851-14.2014.
- Lutz, K., Widmer, M., 2014. What can the monetary incentive delay task tell us about the neural processing of reward and punishment? *Neurosci. Neuroecon.* 33. doi:10.2147/NAN.S38864.
- Ly, A., Verhagen, J., Wagenmakers, E.-J., 2016. Harold Jeffreys's default Bayes factor hypothesis tests: explanation, extension, and application in psychology. *J. Math. Psychol.* 72, 19–32. doi:10.1016/j.jmp.2015.06.004.
- McClure, S.M., Laibson, D.I., Loewenstein, G., Cohen, J.D., 2004. Separate neural systems value immediate and delayed monetary rewards. *Science* 306 (5695), 503–507. doi:10.1126/science.1100907.
- McNamee, D., Rangel, A., O'Doherty, J.P., 2013. Category-dependent and category-independent goal-value codes in human ventromedial prefrontal cortex. *Nat. Neurosci.* 16 (4), 479–485. doi:10.1038/nn.3337.
- Nummenmaa, L., Hirvonen, J., Hannukainen, J.C., Immonen, H., Lindroos, M.M., Salminen, P., Nuutila, P., 2012. Dorsal striatum and its limbic connectivity mediate abnormal anticipatory reward processing in obesity. *PLoS One* 7 (2), e31089. doi:10.1371/journal.pone.0031089.
- O'Doherty, J., Dayan, P., Schultz, J., Deichmann, R., Friston, K., Dolan, R.J., 2004. Dissociable roles of ventral and dorsal striatum in instrumental conditioning. *Science* 304 (5669), 452–454. doi:10.1126/science.1094285.
- Oldham, S., Murawski, C., Fornito, A., Youssef, G., Yücel, M., Lorenzetti, V., 2018. The anticipation and outcome phases of reward and loss processing: A neuroimaging meta-analysis of the monetary incentive delay task. *Hum. Brain Mapp.* 39 (8), 3398–3418. doi:10.1002/hbm.24184.
- Oosterwijk, S., 2017. Choosing the negative: A behavioral demonstration of morbid curiosity. *PLoS One* 12 (7), e0178399. doi:10.1371/journal.pone.0178399.
- Oosterwijk, S., Snoek, L., Tekoppele, J., Engelbert, L.H., Scholte, H.S., 2020. Choosing to view morbid information involves reward circuitry. *Sci. Rep.* 10 (1), 15291. doi:10.1038/s41598-020-71662-y.
- Op de Beeck, H.P., 2010. Against hyperactivity in brain reading: Spatial smoothing does not hurt multivariate fMRI analyses? *Neuroimage* 49 (3), 1943–1948. doi:10.1016/j.neuroimage.2009.02.047.
- Peters, J., Büchel, C., 2010. Episodic future thinking reduces reward delay discounting through an enhancement of prefrontal-midtemporal interactions. *Neuron* 66 (1), 138–148. doi:10.1016/j.neuron.2010.03.026.
- Plassmann, H., O'Doherty, J., Rangel, A., 2007. Orbitofrontal cortex encodes willingness to pay in everyday economic transactions. *J. Neurosci.* 27 (37), 9984–9988. doi:10.1523/JNEUROSCI.2131-07.2007.
- Poldrack, R., 2006. Can cognitive processes be inferred from neuroimaging data? *Trends Cogn. Sci.* 10 (2), 59–63. doi:10.1016/j.tics.2005.12.004.
- Reddan, M.C., Lindquist, M.A., Wager, T.D., 2017. Effect size estimation in neuroimaging. *JAMA Psychiatry* 74 (3), 207. doi:10.1001/jamapsychiatry.2016.3356.
- Rutledge, R.B., Dean, M., Caplin, A., Glimcher, P.W., 2010. Testing the reward prediction error hypothesis with an axiomatic model. *J. Neurosci.* 30 (40), 13525–13536. doi:10.1523/JNEUROSCI.1747-10.2010.
- Schultz, W., Dickinson, A., 2000. Neuronal coding of prediction errors. *Annu. Rev. Neurosci.* 23 (1), 473–500. doi:10.1146/annurev.neuro.23.1.473.
- Scrivner, C., 2021. The psychology of morbid curiosity: Development and initial validation of the morbid curiosity scale. *Pers. Individ. Dif.* 183, 111139. doi:10.1016/j.paid.2021.111139.
- Sescousse, G., Caldú, X., Segura, B., Dreher, J.-C., 2013. Processing of primary and secondary rewards: a quantitative meta-analysis and review of human functional neuroimaging studies. *Neurosci. Biobehav. Rev.* 37 (4), 681–696. doi:10.1016/j.neubiorev.2013.02.002.
- Shafir, S., Waite, T., Smith, B., 2002. Context-dependent violations of rational choice in honeybees (*Apis mellifera*) and gray jays (*Perisoreus canadensis*). *Behav. Ecol. Sociobiol.* 51 (2), 180–187. doi:10.1007/s00265-001-0420-8.
- Shmuel, A., Chaimow, D., Raddatz, G., Ugurbil, K., Yacoub, E., 2010. Mechanisms underlying decoding at 7 T: ocular dominance columns, broad structures, and macroscopic blood vessels in V1 convey information on the stimulated eye. *Neuroimage* 49 (3), 1957–1964. doi:10.1016/j.neuroimage.2009.08.040.
- Simonson, I., 1989. Choice based on reasons: the case of attraction and compromise effects. *J. Consum. Res.* 16 (2), 158. doi:10.1086/209205.
- Smith, S.M., Zhang, Y., Jenkinson, M., Chen, J., Matthews, P.M., Federico, A., De Stefano, N., 2002. Accurate, robust, and automated longitudinal and cross-sectional brain change analysis. *Neuroimage* 17 (1), 479–489. doi:10.1006/nimg.2002.1040.
- Soon, C.S., He, A.H., Bode, S., Haynes, J.-D., 2013. Predicting free choices for abstract intentions. *Proc. Natl. Acad. Sci.* 110 (15), 6217–6222. doi:10.1073/pnas.1212218110.
- Srirangarajan, T., Mortazavi, L., Bortoloni, T., Moll, J., Knutson, B., 2021. Multi-band fMRI compromises detection of mesolimbic reward responses. *Neuroimage* 244, 118617. doi:10.1016/j.neuroimage.2021.118617.
- Treiber, J.M., White, N.S., Steed, T.C., Bartsch, H., Holland, D., Farid, N., McDonald, C.R., Carter, B.S., Dale, A.M., Chen, C.C., 2016. Characterization and correction of geometric distortions in 814 diffusion weighted images. *PLoS One* 11 (3), e0152472. doi:10.1371/journal.pone.0152472.
- Tustison, N.J., Awate, S.P., Cai, J., Altes, T.A., Miller, G.W., de Lange, E.E., Mugler, J.P., Gee, J.C., 2010. Pulmonary kinematics from tagged hyperpolarized helium-3 MRI. *J. Magn. Reson. Imaging* 31 (5), 1236–1241. doi:10.1002/jmri.22137.
- Vallat, R., 2018. Pingouin: statistics in python. *J. Open Source Software* 3 (31), 1026. doi:10.21105/joss.01026.
- van Doorn, J., Ly, A., Marsman, M., Wagenmakers, E.-J., 2020. Bayesian rank-based hypothesis testing for the rank sum test, the signed rank test, and Spearman's ρ . *J. Appl. Statist.* 47 (16), 2984–3006. doi:10.1080/02664763.2019.1709053.
- Van Essen, D.C., Ugurbil, K., Auerbach, E., Barch, D., Behrens, T.E.J., Bucholz, R., Chang, A., Chen, L., Corbetta, M., Curtiss, S.W., Della Penna, S., Feinberg, D.,

- Glasser, M.F., Harel, N., Heath, A.C., Larson-Prior, L., Marcus, D., Michalareas, G., Moeller, S., Yacoub, E., 2012. The human connectome project: a data acquisition perspective. *Neuroimage* 62 (4), 2222–2231. doi:[10.1016/j.neuroimage.2012.02.018](https://doi.org/10.1016/j.neuroimage.2012.02.018).
- Värbu, K., Muhammad, N., Muhammad, Y., 2022. Past, present, and future of EEG-based BCI applications. *Sensors* 22 (9), 3331. doi:[10.3390/s22093331](https://doi.org/10.3390/s22093331).
- Wager, T.D., Atlas, L.Y., Botvinick, M.M., Chang, L.J., Coghill, R.C., Davis, K.D., Iannetti, G.D., Poldrack, R.A., Shackman, A.J., Yarkoni, T., 2016. Pain in the ACC? *Proc. Nat. Acad. Sci. U.S.A.* 113 (18), E2474–E2475. doi:[10.1073/pnas.1600282113](https://doi.org/10.1073/pnas.1600282113).
- Wager, T.D., Atlas, L.Y., Leotti, L.A., Rilling, J.K., 2011. Predicting individual differences in placebo analgesia: contributions of brain activity during anticipation and pain experience. *J. Neurosci.* 31 (2), 439–452. doi:[10.1523/JNEUROSCI.3420-10.2011](https://doi.org/10.1523/JNEUROSCI.3420-10.2011).
- Wager, T.D., Atlas, L.Y., Lindquist, M.A., Roy, M., Woo, C.-W., Kross, E., 2013. An fMRI-based neurologic signature of physical pain. *N. Engl. J. Med.* 368 (15), 1388–1397. doi:[10.1056/NEJMoa1204471](https://doi.org/10.1056/NEJMoa1204471).
- Wager, T.D., Kang, J., Johnson, T.D., Nichols, T.E., Satpute, A.B., Barrett, L.F., 2015. A Bayesian model of category-specific emotional brain responses. *PLoS Comput. Biol.* 11 (4), e1004066. doi:[10.1371/journal.pcbi.1004066](https://doi.org/10.1371/journal.pcbi.1004066).
- Woo, C.-W., Chang, L.J., Lindquist, M.A., Wager, T.D., 2017. Building better biomarkers: brain models in translational neuroimaging. *Nat. Neurosci.* 20 (3), 365–377. doi:[10.1038/nn.4478](https://doi.org/10.1038/nn.4478).
- Woo, C.-W., Koban, L., Kross, E., Lindquist, M.A., Banich, M.T., Ruzic, L., Andrews-Hanna, J.R., Wager, T.D., 2014. Separate neural representations for physical pain and social rejection. *Nat. Commun.* 5 (1), 5380. doi:[10.1038/ncomms6380](https://doi.org/10.1038/ncomms6380).
- Yacubian, J., 2006. Dissociable systems for gain- and loss-related value predictions and errors of prediction in the human brain. *J. Neurosci.* 26 (37), 9530–9537. doi:[10.1523/JNEUROSCI.2915-06.2006](https://doi.org/10.1523/JNEUROSCI.2915-06.2006).
- Yarkoni, T., Poldrack, R.A., Nichols, T.E., Van Essen, D.C., Wager, T.D., 2011. Large-scale automated synthesis of human functional neuroimaging data. *Nat. Methods* 8 (8), 665–670. doi:[10.1038/nmeth.1635](https://doi.org/10.1038/nmeth.1635).
- Yip, S.W., Worhunsky, P.D., Rogers, R.D., Goodwin, G.M., 2015. Hypoactivation of the ventral and dorsal striatum during reward and loss anticipation in antipsychotic and mood stabilizer-naïve bipolar disorder. *Neuropsychopharmacology* 40 (3), 658–666. doi:[10.1038/npp.2014.215](https://doi.org/10.1038/npp.2014.215).
- Yu, H., Koban, L., Chang, L.J., Wagner, U., Krishnan, A., Vuilleumier, P., Zhou, X., Wager, T.D., 2020. A generalizable multivariate brain pattern for interpersonal guilt. *Cereb. Cortex* 30 (6), 3558–3572. doi:[10.1093/cercor/bhz326](https://doi.org/10.1093/cercor/bhz326).
- Zhou, F., Li, J., Zhao, W., Xu, L., Zheng, X., Fu, M., Yao, S., Kendrick, K.M., Wager, T.D., Becker, B., 2020. Empathic pain evoked by sensory and emotional-communicative cues share common and process-specific neural representations. *Elife* 9, e56929. doi:[10.7554/eLife.56929](https://doi.org/10.7554/eLife.56929).

Further reading

- Kurdi, B., Lozano, S., Banaji, M.R., 2017. Introducing the open affective standardized image set (OASIS). *Behav. Res. Methods* 49 (2), 457–470. doi:[10.3758/s13428-016-0715-3](https://doi.org/10.3758/s13428-016-0715-3).
- Saarimäki, H., Ejtehadian, L.F., Gleason, E., Jääskeläinen, I.P., Vuilleumier, P., Sams, M., Nummenmaa, L., 2018. Distributed affective space represents multiple emotion categories across the human brain. *Soc. Cogn. Affect. Neurosci.* 13 (5), 471–482. doi:[10.1093/scan/nsy018](https://doi.org/10.1093/scan/nsy018).
- Zhao, S., Gao, Y., Jiang, X., Yao, H., Chua, T.-S., Sun, X., 2014. Exploring principles-of-art features for image emotion recognition. In: *Proceedings of the 22nd ACM International Conference on Multimedia*, pp. 47–56. doi:[10.1145/2647868.2654930](https://doi.org/10.1145/2647868.2654930).