

Approach-avoidance reinforcement learning as a translational and computational model of anxiety-related avoidance

Reviewed Preprint

Revised by authors after peer review.

[About eLife's process](#)

Reviewed preprint version 2

August 24, 2023 (this version)

Reviewed preprint version 1

June 9, 2023

Posted to bioRxiv

April 4, 2023

Sent for peer review

April 4, 2023

Yumeya Yamamori, Oliver J Robinson, Jonathan P Roiser 

Institute of Cognitive Neuroscience, University College London, London, UK • Research Department of Clinical, Educational and Health Psychology, University College London, London, UK

 https://en.wikipedia.org/wiki/Open_access

 <https://creativecommons.org/licenses/by/4.0/>

Abstract

Although avoidance is a prevalent feature of anxiety-related psychopathology, differences in existing measures of avoidance between humans and non-human animals impede progress in its theoretical understanding and treatment. To address this, we developed a novel translational measure of anxiety-related avoidance in the form of an approach-avoidance reinforcement learning task, by adapting a paradigm from the non-human animal literature to study the same cognitive processes in human participants. We used computational modelling to probe the putative cognitive mechanisms underlying approach-avoidance behaviour in this task and investigated how they relate to subjective task-induced anxiety. In a large online study, participants ($n = 372$) who experienced greater task-induced anxiety avoided choices associated with punishment, even when this resulted in lower overall reward. Computational modelling revealed that this effect was explained by greater individual sensitivities to punishment relative to rewards. We replicated these findings in an independent sample ($n = 627$) and we also found fair-to-excellent reliability of measures of task performance in a sub-sample retested one week later ($n = 57$). Our findings demonstrate the potential of approach-avoidance reinforcement learning tasks as translational and computational models of anxiety-related avoidance. Future studies should assess the predictive validity of this approach in clinical samples and experimental manipulations of anxiety.

eLife assessment

This is a **valuable** paper demonstrating the validity of a novel task that could advance the field of reinforcement learning to better incorporate threat processing in approach-avoidance-conflict. A **compelling** methodology includes the use of online samples and computational modelling, psychometrics, discovery/replication and pre-registration. This work provides a foundation for future work, which is required to address potential confounds and establish this task as relevant to psychopathology and treatment.

1 Introduction

Avoiding harm or potential threats is essential for our well-being and survival, but it can sometimes lead to negative consequences in and of itself. This occurs when avoidance is costly, that is, when the act to avoid harm or threat requires the sacrifice of something positive. When offered the opportunity to attend a job interview, for example, the reward of landing a new job must be weighed against the risk of experiencing rejection and/or embarrassment due to fluffing the interview. On balance, declining a single interview would likely have little impact on one's life, but routinely avoiding job interviews would jeopardise one's long-term professional opportunities. More broadly, consistent avoidance in similar situations, which are referred to as involving *approach-avoidance conflict*, can have an increasingly negative impact as the sum of forgone potential reward accumulates, ultimately leading one to missing out on important things in life. Such excessive avoidance is a hallmark symptom of pathological anxiety (Barlow 2002 [↗](#)) and avoidance biases during approach-avoidance conflict are thought to drive the development and maintenance of anxiety-related (Hayes 1976 [↗](#), Stein and Paulus 2009 [↗](#), Aupperle and Paulus 2010 [↗](#), Mkrtchian, Aylward et al. 2017, Struijs, Lamers et al. 2018) and other psychiatric disorders (Aldao, Nolen-Hoeksema et al. 2010, Kakoschke, Albertella et al. 2019).

How is behaviour during approach-avoidance conflict measured quantitatively? In non-human animals, this can be achieved through behavioural paradigms that induce a conflict between natural approach and avoidance drives. For example, in the 'Geller-Seifter conflict test' (Geller, Kulak et al. 1962), rodents decide whether or not to press a lever that is associated with the delivery of both food pellets and shocks, thus inducing a conflict. Anxiolytic drugs typically increase lever-presses during the test (Treit, Engin et al. 2010), supporting its validity as a model of anxiety-related avoidance. Consequently, this test, along with others inducing approach-avoidance conflict are often used as non-human animal tests of anxiety (Griebel and Holmes 2013 [↗](#)). In humans, on the other hand, approach-avoidance conflict has historically been measured using questionnaires such as the Behavioural Inhibition/Activation Scale (Carver and White 1994 [↗](#)), or cognitive tasks that rely on motor biases, for example by using joysticks to approach/move towards positive stimuli and avoid/move away from negative stimuli, which have no direct non-human counterparts (Guitart-Masip, Huys et al. 2012, Phaf, Mohr et al. 2014, Kirlic, Young et al. 2017, Mkrtchian, Aylward et al. 2017).

Translational approaches, specifically when equivalent tasks are used to measure the same cognitive construct in humans and non-human animals, benefit the study of avoidance and its relevance to mental ill-health for two important reasons (Bach 2021 [↗](#), Pike, Lowther et al. 2021). First, precise causal manipulations of neural circuitry such as chemo/optogenetics are only feasible in non-human animals, whereas only humans can verbalise their subjective experiences – it is only by using translational measures that we can integrate data and theory across species to achieve a comprehensive mechanistic understanding. Second, translational paradigms can make the testing of novel anxiolytic drugs more efficient and cost-effective, since those that fail to reduce anxiety-related avoidance in humans could be identified earlier during the development pipeline.

Since the inception of non-human animal models of psychiatric disorders, researchers have recognised that it is infeasible to fully recapitulate a disorder like generalised anxiety disorder in non-human animals due to the complexity of human cognition and subjective experience. Instead, the main strategy has been to reproduce the physiological and behavioural responses elicited under certain scenarios across species, such as during fear conditioning (Craske, Hermans et al. 2006, Milad and Quirk 2012 [↗](#)). Computational modelling of behaviour (Kriegeskorte and Douglas 2018 [↗](#)), allows us to take a step further and probe the *mechanisms* underlying such physiological and behavioural responses, using a mathematically principled approach. A fundamental premise

of this approach is that the brain acts as an information-processing organ that performs computations responsible for observable behaviours, including approach and avoidance (for a recent review on the application of computational methods to approach-avoidance conflict, see Letkiewicz, Kottler et al. 2023). With respect to translational models, it follows that a valid translational measure of some behaviour not only matches the observable responses across species, but also their underlying computational mechanisms. This criterion, which has been termed *computational validity*, has been proposed as the primary basis for evaluating the validity of translational measures (Redish, Kepecs et al. 2022).

To exploit the advantages of translational and computational approaches, there has been a recent surge in efforts to create relevant approach-avoidance conflict tasks in humans based on non-human animal models (referred to as back-translation; Kirlic, Young et al. 2017). One stream of research has focused on exploration-based paradigms, which involve measuring how participants spend their time in environments that include potentially anxiogenic regions. Rodents typically avoid such regions, which can take the form of tall, exposed platforms as in the elevated plus maze (Pellow, Chopin et al. 1985), or large open spaces as in the open-field test (Hall 1934 [↗](#)). Similar behaviours have been observed in humans when virtual-reality was used to translate the elevated plus maze (Biedermann, Biedermann et al. 2017 [↗](#)), and when participants were free to walk around a field or around town (Walz, Muhlberger et al. 2016). These translational tasks have excellent face validity in that they involve almost identical behaviour compared to their non-human analogues, and their predictive validity is supported by findings that individual differences in self-reported anxiety are associated with greater avoidance of the exposed/open regions of space (Walz, Muhlberger et al. 2016, Biedermann, Biedermann et al. 2017 [↗](#)). But what of their computational validity? The computations underlying free movement are difficult to model as the action-space is large and complex (consisting of 2-dimensional directions, velocity, etc.), which makes it difficult to understand the precise cognitive mechanisms underlying avoidance in these tasks.

Other studies have designed tasks that emulate operant conflict paradigms in which non-human participants learn to associate certain actions with both reward and punishment, such as the Geller-Seifter and Vogel conflict tests (Geller, Kulak et al. 1962, Vogel, Beer et al. 1971). These studies used decision-making tasks in which participants choose between offers consisting of both reward and punishment outcomes (Aupperle, Sullivan et al. 2011, Sierra-Mercado, Deckersbach et al. 2015). The strengths of such decision tasks lie in the simplicity of their action-space (accept/reject an offer; choose offer A or B), which facilitates the computational modelling of decision-making, through for example value-based logistic regression models (Ironsides, Amemori et al. 2020) or drift-diffusion modelling (Pedersen, Ironsides et al. 2021). However, these tasks have less face validity compared to exploration-based tasks, as explicit offers ('Will you accept £1 to incur an electric shock?') are incomprehensible to the majority of non-human animals (although non-human primates can be trained to do so, see: Sierra-Mercado, Deckersbach et al. 2015, Ironsides, Amemori et al. 2020). Instead, non-human participants such as rodents typically have to learn the possible outcomes given their actions, for example that a lever-press is associated with both food pellets and electric shocks. Critically, this means that humans performing offer-based decision-making are doing something *computationally dissimilar* to non-human participants during operant conflict tasks, and the criterion of computational validity is unmet. Moreover, evidence from economic decision-making suggests that explicit offers of probabilistic outcomes can impact decision-making differently compared to when probabilistic contingencies need to be learned from experience (referred to as the 'description-experience gap'; Hertwig and Erev 2009 [↗](#)); this finding raises potential concerns regarding the use of offer-based tasks in humans as approximations of non-human tasks that do not involve explicit offers.

Therefore, in the current study, we developed a novel translational task that was specifically designed to mimic the computational processes involved in non-human operant conflict tasks, and was also amenable to computational modelling of behaviour. During the task, participants chose

between options that were asymmetrically associated with rewards and punishments, such that the more rewarding options were more likely to result in punishment. However, in contrast to previous translations of these tests using purely offer-based designs, participants had to learn the likelihoods of receiving rewards and/or punishments from experience, thus ensuring computational validity. Preserving the element of learning in our translational task did not necessarily risk confounding the study of approach-avoidance behaviour. Rather, a critical advantage of our approach was that we could model the cognitive processes underlying approach-avoidance behaviour using classical reinforcement learning algorithms (Sutton and Barto 2018 [↗](#)), and then use these algorithms to disentangle the relative contributions of learning and value-based decision-making to anxiety-related avoidance. Specifically, using the reinforcement learning framework; we accounted for individual differences in learning by estimating learning rates, which govern the degree to which recently observed outcomes affect subsequent actions; and we explained individual differences in value-based decision-making using outcome sensitivity parameters, which model the subjective values of rewards and punishments for each participant.

We therefore aimed to develop a new translational measure of approach avoidance conflict. Based on previous findings that anxiety predicts avoidance biases under conflict (Walz, Muhlberger et al. 2016, Biedermann, Biedermann et al. 2017 [↗](#)), we predicted that anxiety would increase avoidance responses (and conversely decrease approach responses) and that this would be reflected in changes in computational parameters. Having found evidence for this in a large initial online study, we retested the prediction that anxiety-related avoidance is explained by greater sensitivity to punishments relative to rewards in a pre-registered independent replication sample.

2 Results

2.1 The approach-avoidance reinforcement learning task

We developed a novel approach-avoidance conflict task, based on the “restless bandit” design (Daw, O’Doherty et al. 2006), in which participants’ goal was to accrue rewards whilst avoiding punishments (**Figure 1a** [↗](#)). On each of 200 trials, participants chose one of two options, depicted as distinct visual stimuli, to obtain certain outcomes. There were four possible outcomes: a monetary reward (represented by a coin); an aversive sound consisting of a combination of a female scream and a high-pitched sound; both the reward and aversive sound; or no reward and no sound (**Figure 1b** [↗](#)). The options were asymmetric in their associated outcomes, such that one option (which we refer to as the ‘conflict’ option) was associated with both the reward and aversive sound, whereas the other (the ‘safe’ option) was only associated with the reward (**Figure 1c** [↗](#)). In other words, choosing the conflict option could result in any of the four possible outcomes on any given trial, but the safe option could only lead to either the reward or no reward, and never the aversive sounds.

The probability of observing the reward and the aversive sound at each option fluctuated randomly and independently over trials, except the probability of observing the aversive sound at the safe option, which was set to 0 across all trials (**Figure 1c** [↗](#)). On average, the safe option was less likely to produce the reward compared to the conflict option across the task (**Figure 1d** [↗](#)). Therefore, we induced an approach-avoidance conflict between the options as the conflict option was generally more rewarding but also more likely to result in punishment. Participants were not informed about this asymmetric feature of the task, but they were told that the probabilities of observing the reward or the sound were independent and that they could change across trials.

We initially investigated behaviour in the task in a discovery sample ($n = 369$) after which the key effects-of-interest were retested in a pre-registered study (<https://osf.io/8dm95> [↗](#)) with an independent replication sample ($n = 629$). The discovery sample consisted of a significantly greater proportion of female participants than the replication sample (59% vs 52%, $\chi^2 = 4.64$, $p = 0.031$).

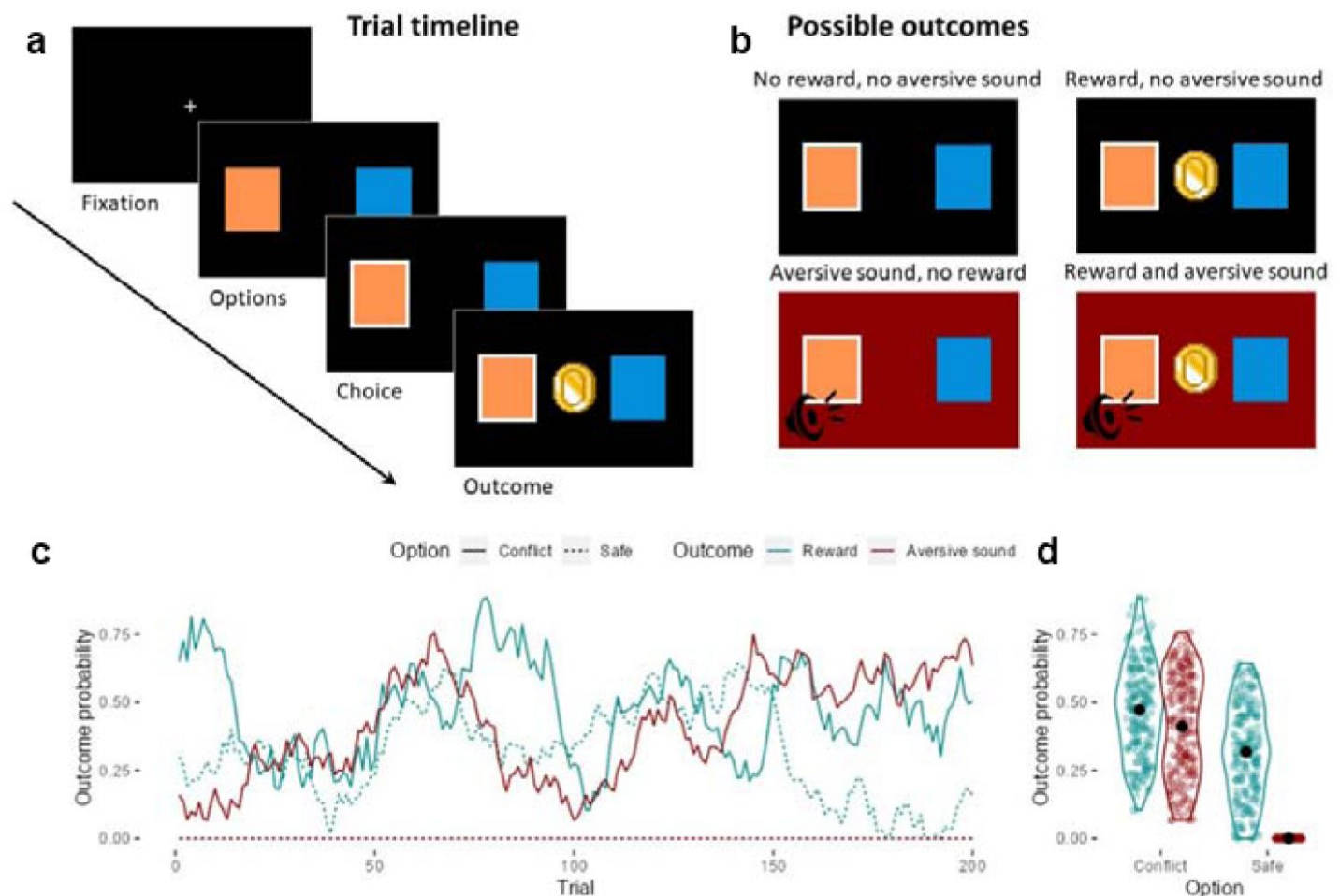


Figure 1.

The approach-avoidance reinforcement learning task.

a, Trial timeline. A fixation cross initiates the trial. Participants are presented with two options for up to 2 s, from which they choose one. The outcome is then presented for 1 s. **b**, Possible outcomes. There were four possible outcomes: 1) no reward and no aversive sound; 2) a reward and no aversive sound; 3) an aversive sound and no reward; or 4) both the reward and the aversive sound. **c**, Probabilities of observing each outcome given the choice of option. Unbeknownst to the participant, one of the options (which we refer to as the 'conflict' option - solid lines) was generally more rewarding compared to the other option (the 'safe' option - dashed line) across trials. However, the conflict option was also the only option of the two that was associated with a probability of producing an aversive sound (the probability that the safe option produced the aversive sound was 0 across all trials). The probabilities of observing each outcome given the choice of option fluctuated randomly and independently across trials. The correlations between these dynamic probabilities were negligible (mean Pearson's $r = 0.06$). **d**, Distribution of outcome probabilities by option and outcome. On average, the conflict option was more likely to produce a reward than the safe option. The conflict option had a variable probability of producing the aversive sound across trials, but this probability was always 0 for the safe option.

The average age was significantly different across samples (discovery sample mean = 37.7, SD = 10.3, replication sample mean = 34.3, SD = 10.4; $t_{785.5} = 5.06$, $p < 0.001$). The differences in self-reported psychiatric symptoms across samples did not reach significance ($p > 0.086$). Since our behavioural findings were broadly comparable across samples, we primarily report findings from the discovery sample and discuss deviations from these findings in the replication sample. A detailed comparison of statistical findings across samples is reported in the Supplement.

2.2 Choices reflect both approach and avoidance

First, we verified that participants demonstrated both approach and avoidance responses during the task through a hierarchical logistic regression model, in which we used the latent outcome probabilities to predict trial-by-trial choice. Overall, participants learned to optimise their choices to accrue rewards and avoid the aversive sounds. As expected, across trials participants chose (i.e. approached) the option with relatively higher reward probability ($\beta = 0.98 \pm 0.03$, $p < 0.001$; **Figure 2a, c**). At the same time, they avoided the conflict option when it was more likely to produce the punishment ($\beta = -0.33 \pm 0.04$, $p < 0.001$; **Figure 2a, d**). These effects were also clearly evident in our replication sample (effect of relative reward probability: $\beta = 0.95 \pm 0.02$, $p < 0.001$; effect of punishment probability: $\beta = -0.52 \pm 0.03$, $p < 0.001$; **Supplementary Table 1**). While sex was significantly associated with choice in the discovery sample ($\beta = 0.16 \pm 0.07$, $p = 0.028$) with males being more likely to choose the conflict option, this pattern was not evident in the replication sample ($\beta = 0.08 \pm 0.06$, $p = 0.173$), and age was not significantly associated with choice in either sample ($p > 0.2$). Across individuals, there was considerable variability in overall choice proportions (discovery sample: mean = 0.52, SD = 0.14, min/max = [0.03, 0.96]; replication sample: mean = 0.52, SD = 0.14, min/max = [0.01, 0.99]).

2.3 Task-induced anxiety is associated with greater avoidance

Immediately after the task, participants rated their task-induced anxiety in response to the question, ‘How anxious did you feel during the task?’, from ‘Not at all’ to ‘Extremely’ (scored from 0 to 50). On average, participants reported a moderate level of task-induced anxiety (mean rating of 21, SD = 14; **Figure 2b**).

Task-induced anxiety was significantly associated with multiple aspects of task performance. At the most basic level, task-induced anxiety was correlated with the proportion of conflict/safe option choices made by each participant across all trials, with greater anxiety corresponding to a preference for the safe option over the conflict option (permutation-based non-parametric correlation test: Kendall’s $\tau = -0.074$, $p = 0.033$), at the cost of fewer rewards on average.

To investigate this association further, we included task-induced anxiety into the hierarchical logistic regression model of trial-by-trial choice. Here, task-induced anxiety significantly interacted with punishment probability ($\beta = -0.10 \pm 0.04$, $p = 0.022$; **Figure 2a**), such that more anxious individuals were proportionately less willing to choose the conflict option during times it was likely to produce the punishment, relative to less-anxious individuals (**Figure 2e**). This explained the effect of anxiety on choices, as the main effect of anxiety was not significant in the model, although it was when the interaction effects were excluded. The analogous interaction between task-induced anxiety and relative reward probability was non-significant ($\beta = -0.06 \pm 0.03$, $p = 0.074$, **Figure 2a**).

These associations with task-induced anxiety were again clearly evident in the replication sample (correlation between task-induced anxiety and proportion of conflict/safe option choices: $\tau = -0.075$, $p = 0.005$; interaction between task-induced anxiety and punishment probability in the hierarchical model: $\beta = -0.23 \pm 0.03$, $p < 0.001$; **Supplementary Table 1**).

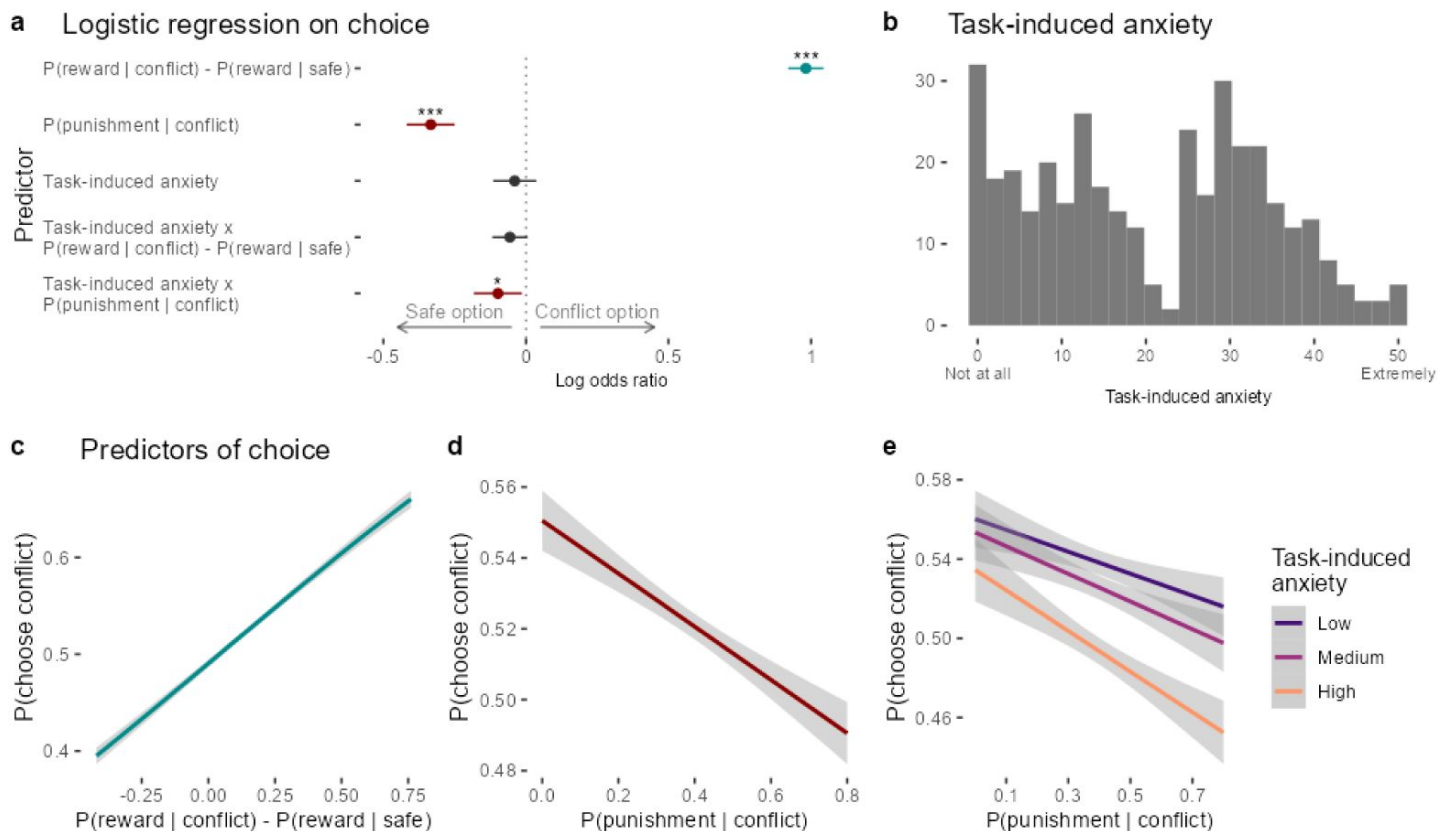


Figure 2.

Predictors of choice in the approach-avoidance reinforcement learning task.

a, Coefficients from the mixed-effects logistic regression of trial-by-trial choices in the task. On any given trial, participants chose the option that was more likely to produce a reward. They also avoided choosing the conflict option when it was more likely to produce the punishment. Task-induced anxiety significantly interacted with punishment probability. Significance levels are shown according to the following: $p < 0.05$ - *; $p < 0.05$ - **; $p < 0.001$ - ***. **b**, Subjective ratings of task-induced anxiety, given on a scale from 'Not at all' (0) to 'Extremely' (50). **c**, On each trial, participants were likely to choose the option with greater probability of producing the reward. **d**, Participants tended to avoid the conflict option when it was likely to produce a punishment. **e**, Compared to individuals reporting lower anxiety during the task, individuals experiencing greater anxiety showed greater avoidance of the conflict option, especially when it was more likely to produce the punishment. *Note*. Figures **c-e** show logistic curves fitted to the raw data using the 'glm' function in R. For visualisation purposes, we categorised continuous task-induced anxiety into tertiles.

2.4 A reinforcement learning model explains individual choices

Next, we investigated the putative cognitive mechanisms driving behaviour on the task by fitting standard reinforcement learning models to trial-by-trial choices. The winning model in both the discovery and replication samples included reward- and punishment-specific learning rates, and reward- and punishment-specific sensitivities (i.e. four parameters in total; **Figure 3a, b**). Briefly, it models participants as learning four values, which represent the probabilities of observing each outcome (reward/punishment) for each option (conflict/safe). At every trial, the probability estimates for the chosen option are updated after observing the outcome(s) according to the Rescorla-Wagner update rule. Individual differences in the rate of updating are captured by the learning rate parameters (α^r , α^p), whereas differences in the extent to which each outcome impacts choice are captured by the sensitivity parameters (β^r , β^p).

To assess how well this model captured our observed data, we simulated data from the model given each participant's parameter values, choices and observed outcomes. These synthetic choice data were qualitatively similar to the actual choices made by the participants (**Figure 3c**). Parameter recovery from the winning model was excellent (Pearson's r between data generating and recovered parameters ≈ 0.77 to 0.86 ; **Supplementary Figure 3**).

To quantify individual approach-avoidance bias computationally, we calculated the ratio between the reward and punishment sensitivity parameters (β^r/β^p). As the task requires simultaneously balancing reward pursuit (i.e., approach) and punishment avoidance, this composite measure provides an index of where an individual lies on the continuum of approach vs avoidance, with higher and lower values indicating approach and avoidance biases, respectively. We refer to this measure as the 'reward-punishment sensitivity index' (**Figure 3d**).

Comparing parameters across sexes via Welch's t-tests revealed significant differences in reward sensitivity ($t_{289} = -2.87$, $p = 0.004$, $d = 0.34$; lower in females) and consequently reward-punishment sensitivity index ($t_{336} = -2.03$, $p = 0.043$, $d = 0.22$; lower in females i.e. more avoidance-driven). In the replication sample, we observed the same effect on reward-punishment sensitivity index ($t_{626} = -2.79$, $p = 0.005$, $d = 0.22$; lower in females). However, the sex difference in reward sensitivity did not replicate ($p = 0.441$), although we did observe a significant sex difference in punishment sensitivity in the replication sample ($t_{626} = 2.26$, $p = 0.024$, $d = 0.18$).

2.5 Computational mechanisms of anxiety-related avoidance

We assessed the relationship between task-induced anxiety and individual model parameters through permutation testing of non-parametric correlations, since the distribution of task-induced anxiety was non-Gaussian. The punishment learning rate and reward-punishment sensitivity index were significantly and negatively correlated with task-induced anxiety in the discovery sample (punishment learning rate: Kendall's $\tau = -0.088$, $p = 0.015$; reward-punishment sensitivity index : $\tau = -0.099$, $p = 0.005$; **Figure 4a-e**). These associations were also significant in the replication sample (punishment learning rate: $\tau = -0.064$, $p = 0.020$; reward-punishment sensitivity index : $\tau = -0.096$, $p < 0.001$), with an additional positive correlation between punishment sensitivity and anxiety ($\tau = 0.076$, $p = 0.004$; **Supplementary Table 1**) which may have been detected due to greater statistical power. The learning rate equivalent of the reward-punishment sensitivity index (i.e., α^r/α^p) did not significantly correlate with task-induced anxiety (discovery sample: $\tau = 0.048$, $p = 0.165$; replication sample: $\tau = 0.029$, $p = 0.298$).

Given that both the punishment learning rate and the reward-punishment sensitivity index were significantly and reliably correlated with task-induced anxiety, we next asked whether these computational measures explained the effect of anxiety on avoidance behaviour. We tested this via structural equation modelling, including both the punishment learning rate and reward-punishment sensitivity index as parallel mediators for the effect of task-induced anxiety on proportion of choice for conflict/safe options. This suggested that the mediating effect of the

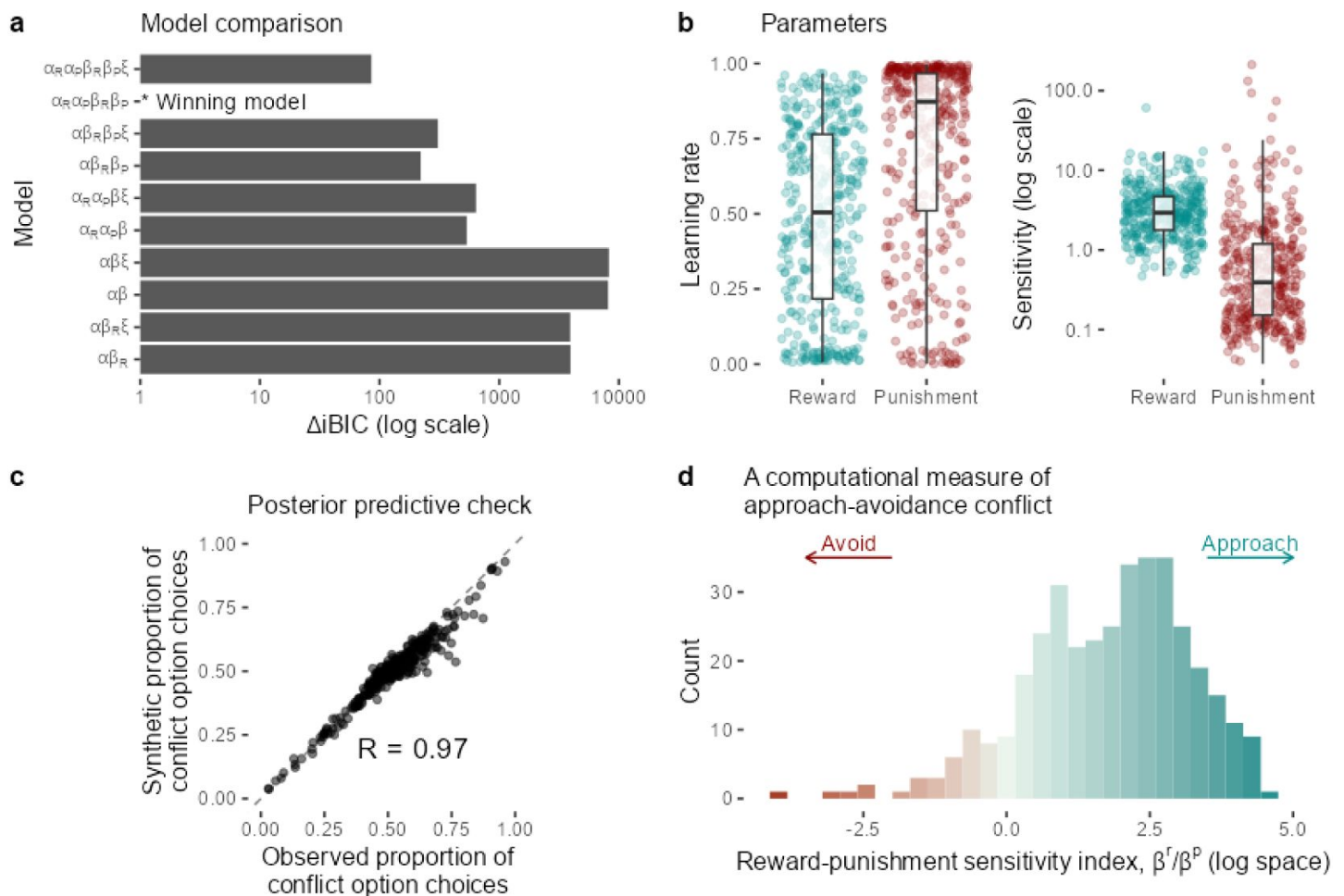


Figure 3.

Computational modelling of approach-avoidance reinforcement learning.

a, Model comparison results. The difference in integrated Bayesian Information Criterion scores from each model relative to the winning model is indicated on the x-axis. The winning model included specific learning rates for reward (α^R) and punishment learning (α^P), and specific outcome sensitivity parameters for reward (β^R) and punishment (β^P). Some models were tested with the inclusion of a lapse term (ξ). **b**, Distributions of individual parameter values from the winning model. **c**, The winning model was able to reproduce the proportion of conflict option choices over all trials in the observed data with high accuracy (observed vs predicted data $r = 0.97$). **d**, The distribution of the reward-punishment sensitivity index – the computational measure of approach-avoidance bias. Higher values indicate approach biases, whereas lower values indicate avoidance biases.

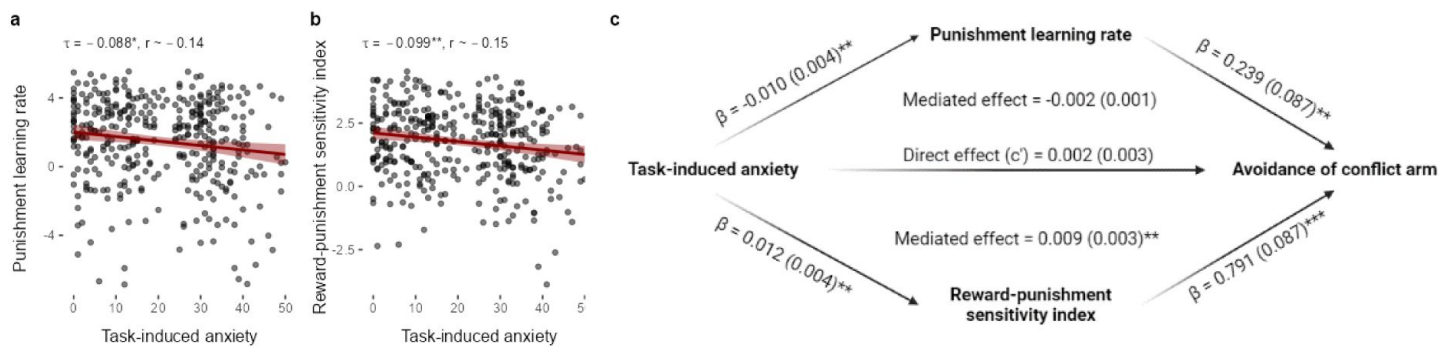


Figure 4.

Relationships between task-induced anxiety, model parameters and avoidance.

a, Task-induced anxiety was negatively correlated with the punishment learning rate. **b**, Task-induced anxiety was also negatively correlated with reward-punishment sensitivity index. Kendall's tau correlations and approximate Pearson's r equivalents are reported above each figure. **c**, The mediation model. Mediation effects were assessed using structural equation modelling. Bold terms represent variables and arrows depict regression paths in the model. The annotated values next to each arrow show the regression coefficient associated with that path, denoted as *coefficient (standard error)*. Only the reward-punishment sensitivity index significantly mediated the effect of task-induced anxiety on avoidance. Significance levels in all figures are shown according to the following: $p < 0.05$ - *; $p < 0.05$ - **; $p < 0.001$ - * * *, *

reward-punishment sensitivity index (standardised $\beta = -0.009 \pm 0.003$, $p = 0.003$) was over four times larger than the mediating effect of the punishment learning rate, which did not reach the threshold as a significant mediator (standardised $\beta = -0.002 \pm 0.001$, $p = 0.052$; **Figure 4f**). A similar pattern was observed in the replication sample (mediating effect of reward-punishment sensitivity index: standardised $\beta = -0.009 \pm 0.002$, $p < 0.001$; mediating effect of punishment learning rate: standardised $\beta = -0.001 \pm 0.0003$, $p = 0.132$; **Supplementary Table 1**). **Supplementary Figure 3**.

This analysis reveals two effects: first, that individuals who reported feeling more anxious during the task were slower to update their estimates of punishment probability; and second, more anxious individuals placed a greater weight on the aversive sound *relative to the reward* when making their decisions, meaning that the potential for a reward was more strongly offset by potential for punishment. Critically, the latter but not the former effect explained the effect of task-induced anxiety on avoidance behaviour, thus providing an important insight into the computational mechanisms driving anxiety-related avoidance.

2.6 Psychiatric symptoms predict task-induced anxiety, but not avoidance

We assessed the clinical potential of the measure by investigating how model-agnostic and computational indices of behaviour correlate with self-reported psychiatric symptoms of anxiety, depression and experiential avoidance; the latter indexes an unwillingness to experience negative emotions/experiences (Hayes, Wilson et al. 1996). We included experiential avoidance as a direct measure of avoidance symptoms/beliefs, given that our task involved avoiding stimuli inducing negative affect. Given that we recruited samples without any inclusion/exclusion criteria relating to mental health, the instances of clinically-relevant anxiety symptoms, defined as a score of 10 or greater in the Generalised Anxiety Disorder 7-item scale (Spitzer, Kroenke et al. 2006), were low in both the discovery (16%, 60 participants) and the replication sample (20%, 127 participants).

Anxiety, depression and experiential avoidance symptoms were all significantly and positively correlated with task-induced anxiety ($p < 0.001$; **Supplementary Table 1**). Only experiential avoidance symptoms were significantly correlated with the proportion of conflict/safe option choices ($\tau = -0.059$, $p = 0.010$; **Supplementary Table 1**), but this effect was not observed in the replication sample ($\tau = -0.029$, $p = 0.286$). Experiential avoidance was also significantly correlated with punishment learning rates (Kendall's $\tau = -0.09$, $p = 0.024$) and the reward-punishment sensitivity index ($\tau = -0.09$, $p = 0.034$) in the discovery sample. However, these associations were non-significant in the replication sample ($p > 0.3$), and no other significant correlations were detected between the psychiatric symptoms and model parameters (**Supplementary Table 1**).

2.7 Split-half reliability

We assessed the split-half reliability of the task by correlating the overall proportion of conflict option choices and model parameters from the winning model across the first and second half of trials. For overall choice proportion, reliability was simply calculated via Pearson's correlations. For the model parameters, we calculated model-derived estimates of Pearson's r values from the parameter covariance matrix when first- and second-half parameters were estimated within a single model, following a previous approach recently shown to accurately estimate parameter reliability (Waltmann, Schlagenhauf et al. 2022). We interpreted indices of reliability based on conventional values of < 0.40 as poor, $0.4 - 0.6$ as fair, $0.6 - 0.75$ as good, and > 0.75 as excellent reliability (Fleiss 1986). Overall choice proportion showed good reliability (discovery sample $r = 0.63$; replication sample $r = 0.63$; **Supplementary Figure 5**). The model parameters showed good-to-excellent reliability (model-derived r values ranging from 0.61 to 0.85 [0.76 to 0.92 after Spearman-Brown correction]; **Supplementary Figure 5**).

2.8 Test-retest reliability

We assessed the test-retest reliability of the task by retesting 57 participants from the replication sample at least 10 days after they initially completed the task. The retest version of the task was identical to the initial version, except for the use of new stimuli to represent the conflict and safe options and different latent outcome probabilities to limit practise effects. The test-retest reliability of model-agnostic effects was assessed using the intra-class correlation (ICC) coefficient, specifically ICC(3,1) (following the notation of Shrout and Fleiss 1979). Test-retest reliability of the model parameters was assessed by calculating Pearson correlations derived from the parameter covariance matrix when parameters from both sessions were estimated within a single model.

Task-induced anxiety, overall proportion of conflict option choices and the model parameters demonstrated fair to excellent reliability (reliability indices ranging from 0.4 to 0.77;

Supplementary Figure 6), except for the punishment learning rate which showed a model-derived reliability of $r = -0.45$. Practise effects analyses showed that task induced anxiety and the punishment learning rate were also significantly lower in the second session (anxiety: $t_{56} = 2.21$, $p = 0.031$; punishment learning rate: $t_{56} = 2.24$, $p = 0.029$), whilst the overall probability of choosing the conflict option and the other parameters did not significantly change over sessions (**Supplementary Figure 7**). As the dynamic (latent) outcome probabilities were different across the test vs. retest sessions, in other words participants were not performing the exact same task across sessions, we argue that even ‘fair’ reliability (i.e., estimates in range of 0.4 to 0.6) is encouraging, as this suggests that model parameters generalise across different environments if the required computations are preserved. However, potential within-subjects inferences based on the punishment learning rate may not be justified in this task.

2.9 Sensitivity analyses

As our procedure for estimating model parameters (the expectation-maximisation algorithm, see ‘Methods’) produced high inter-parameter correlations in our data (**Supplementary Figure 2**), we also re-estimated the parameters using Stan’s variational Bayesian inference algorithm (Stan Development Team 2023) – this resulted in lower inter-parameter correlations, but our primary computational finding, that the effect of anxiety on choice is mediated by relative sensitivity to reward/punishment was consistent across algorithms (see Supplement section 9.8 for details).

Since participants self-determined the volume of the punishments in the task, and therefore (at least in part) their aversiveness, we also conducted a sensitivity analysis by accounting for self-reported unpleasantness ratings of the punishment (see the Supplement). Our finding that anxiety impacts approach-avoidance behaviour was robust to this sensitivity analysis ($p < 0.001$), however the mediating effect of the reward-sensitivity sensitivity index was not ($p > 0.1$; see Supplement section 9.9 for details). Although sound volume and stimulus unpleasantness correlate (Seow and Hauser 2022), unpleasantness ratings are not a perfect indicator of volume, so this analysis should be interpreted with caution.

3 Discussion

To develop a translational and computational model of anxiety-related avoidance, we created a novel task that involves learning under approach-avoidance conflict and investigated the impact of anxiety on task performance. Importantly, we show that individuals who experienced more task-induced anxiety avoid actions that lead to aversive outcomes, at the cost of potential reward. Computational modelling of behaviour revealed that this was explained by relative sensitivity to punishment over reward. However, this effect was limited to task-induced anxiety, and was only evident on a self-report measure of *general* avoidance in our first sample. Finally, split-half and

test-retest analyses of the task indicated fair-to-excellent reliability indices of task behaviour and model parameters. These results demonstrate the potential of approach-avoidance reinforcement learning tasks as translational measures of anxiety-related avoidance.

We can evaluate the translational validity of our task based on the criteria of face, construct, predictive and computational validity (Willner 1984 [↗](#), Redish, Kepecs et al. 2022). The task shows sufficient face validity when compared to the Vogel/Geller-Seifter conflict tests; the design involves a series of choices between an action associated with high reward and punishment, or low reward and no punishment, and these contingencies must be learned. One potentially important difference is that our task involved a choice between two active options, referred to as a two-alternative forced choice (2AFC), whereas in the Vogel/Geller-Seifter tests, animals are free to perform an action (e.g. press the lever) or withhold the action (don't press the lever; i.e. a go/no-go design). The disadvantage of the latter approach is that human participants can and typically do show motor response biases (i.e. for either action or inaction), which can confound value-based response biases with approach or avoidance (Guitart-Masip, Huys et al. 2012). Fortunately, there is ample evidence that rodents can perform 2AFC learning (Metha, Brian et al. 2020) and conflict tasks (Simon, Gilbert et al. 2009, Oberrrauch, Sigrist et al. 2019, Glover, Postle et al. 2020). That avoidance responses were positively associated with task-induced anxiety also lends support for the construct validity of our measure, given the theoretical links between anxiety and avoidance under approach-avoidance conflict (Hayes 1976 [↗](#), Stein and Paulus 2009 [↗](#), Aupperle and Paulus 2010 [↗](#)). However, we note that the empirical effect size was small ($|r| \sim 0.13$) - we discuss potential reasons for this below.

The predictive and computational validity of the task needs to be determined in future work. Predictive validity can be assessed by investigating the effects of anxiolytic drugs on task behaviour. Reduced avoidance after anxiolytic drug administration, or indeed any manipulation designed to reduce anxiety, would support our measure's predictive validity. Similarly, computational validity would also need to be assessed directly in non-human animals by fitting models to their behavioural data. This should be possible even in the face of different procedures across species such as number of trials or outcomes used (shock or aversive sound). We are encouraged by our finding that the winning computational model in our study relies on a relatively simple classical reinforcement learning strategy. There exist many studies showing that non-human animals rely on similar strategies during reward and punishment learning (Schultz 2013 [↗](#), Mobbs, Headley et al. 2020); albeit to our knowledge this has never been modelled in non-human animals where rewards and punishment can occur simultaneously.

One benefit of designing a task amenable to the computational modelling of behaviour was that we could extract model parameters that represent individual differences in precise cognitive mechanisms. Our model recapitulated individual- and group-level behaviour with high fidelity, and mediation analyses suggested a potential mechanism of how anxiety may induce avoidance, namely through the relative sensitivity to punishment compared to reward. This demonstrates the potential of computational methods in allowing for mechanistic insights into anxiety-related avoidance. Further, we found satisfactory test-retest reliability of the reward-punishment sensitivity index in test-retest analyses, suggesting that the task can be used for within-subject experiments (e.g. on/off medication) to assess the impact of interventions on the mechanisms underlying avoidance.

The composite nature of the reward-punishment sensitivity index may be important for contextualising prior results of anxiety and outcome sensitivity. The traditional view has long been that anxiety is associated with greater automatic sensitivity to punishments and threats (i.e. that anxious individuals show a negative bias in information processing; Bishop 2007 [↗](#)), due to amygdala hypersensitivity (Mathews and Mackintosh 1998 [↗](#), Hartley and Phelps 2012 [↗](#)). However, subsequent studies have found inconsistent findings relative to this hypothesis (Bishop and Gagne 2018 [↗](#)). We found that, rather than a straightforward link between anxiety and

punishment sensitivity, the *relative* sensitivity to punishment over reward provided a better explanation of anxiety's effect on behaviour. This effect is more in keeping with cognitive accounts of anxiety's effect on information processing that involve higher-level cognitive functions such as attention and cognitive control, which propose that anxiety involves a global reallocation of resources away from processing positive and towards negative information (Beck and Clark 1997 [↗](#), Cisler and Koster 2010 [↗](#), Robinson, Vytal et al. 2013), possibly through prefrontal control mechanisms (Carlisi and Robinson 2018 [↗](#)). Future studies using experimental manipulations of anxiety will be needed to test the causal role of anxiety on relative sensitivities to punishments versus rewards.

Since we developed our task with the primary focus on translational validity, its design diverges from other reinforcement learning tasks that involve reward and punishment outcomes (Pike and Robinson 2022 [↗](#)). One important difference is that we used distinct reinforcers as our reward and punishment outcomes, compared to many studies which use monetary outcomes for both (e.g. earning and losing £1 constitute the reward and punishment, respectively; Pizzagalli, Jahn et al. 2005, Aylward, Valton et al. 2019, Jean-Richard-Dit-Bressel, Lee et al. 2021, Sharp, Russek et al. 2022). Other tasks have been used that induce a conflict between value and motor biases, relying on prepotent biases to approach/move towards rewards and withdraw from punishments, which makes it difficult to approach punishments and withdraw from rewards (Guitart-Masip, Huys et al. 2012, Mkrtchian, Aylward et al. 2017). However, since translational operant conflict tasks typically induce a conflict between different types of outcome (e.g. food and shocks/sugar and quinine pellets; van den Bos, Koot et al. 2014, Oberrauch, Sigrist et al. 2019), we felt it was important to implement this feature. One study used monetary rewards and shock-based punishments, but also included four options for participants to choose from on each trial, with rewards and punishments associated with all four options (Seymour, Daw et al. 2012). This effectively requires participants to maintain eight probability estimates (i.e. reward and punishment at each of the four options) to solve the task, which may be too difficult for non-human animals to learn efficiently.

Notably, although a recent computational meta-analysis of reinforcement learning studies showed that symptoms of anxiety and depression are associated with elevated punishment learning rates (Pike and Robinson 2022 [↗](#)), we did not observe this pattern in our data. Indeed, we even found the contrary effect in relation to task-induced anxiety, specifically that anxiety was associated with lower rates of learning from punishment. However, other work has suggested that the direction of this effect can depend on the form of anxiety, where cognitive anxiety may be associated with elevated learning rates, but somatic anxiety may show the opposite pattern (Wise and Dolan 2020 [↗](#)) and this may explain the discrepancy in findings. Additionally, parameter values are highly dependent on task design (Eckstein, Master et al. 2022), and study designs to date may be more optimised in detecting differences in learning rate (Pike and Robinson 2022 [↗](#)) – future work is needed to better understand the potentially complex association between anxiety and punishment learning rate. Lastly, as punishment learning rate was severely unreliable in the test-retest analyses, and the associations between punishment learning rate and state anxiety were not robust to an alternative method of parameter estimation (variational Bayesian inference), the negative correlation observed in our study should be treated with caution.

One potentially important limitation of our findings is the small effect size observed in the correlation between task-induced anxiety and avoidance (Kendall's tau values < 0.1 , mediation betas < 0.01). This may be attributed to the simplicity of using overall choice proportion as a measure of approach/avoidance, as the effect of anxiety on choice was also influenced by punishment probability. This may have also been due to the aversiveness of the punishment outcomes (aversive sounds) used in our study, compared to other studies which included aversive images (Aupperle, Sullivan et al. 2011), shocks (Talmi, Dayan et al. 2009), or used virtual reality (Biedermann, Biedermann et al. 2017 [↗](#)) to provide negative reinforcement. Using more salient punishments would likely produce larger effects. Further, whilst we included tests to check

whether participants were wearing headphones for this online study, we could not be certain that the sounds were not played over speakers and/or played at a sufficiently high volume to be aversive. Finally, it may simply be that the relationships between psychiatric symptoms and behaviour are small due to the inherent variability of individuals, the lack of precision of self-reported symptoms or the unreliability of mental health diagnoses (Freedman, Lewis et al. 2013, Wise, Robinson et al. 2022). Despite these limitations, an advantage of using web-based auditory stimuli is their efficiency in scaling with sample size, as only a web-browser is needed to complete the task. This is especially relevant for psychiatric research, given the need for studies with larger sample sizes in the field (Rutledge, Chekroud et al. 2019).

Relatedly, participants had some control over the intensity at which the punishments were presented, which may have driven our findings relating to anxiety and putative mechanisms of anxiety-related avoidance. Sensitivity analyses showed that our finding that anxiety is positively associated with avoidance in the task was robust to individual differences in self-reported punishment unpleasantness, whilst the mediation effects were not. Future work imposing better control over the stimuli presented, and/or using within-subjects designs will be needed to validate the role of reward/punishment sensitivities in anxiety-related avoidance.

The punishment learning rate was severely unreliable given a model-derived correlation across sessions of $r = -0.45$, whereas the other indices of task performance showed fair-to-excellent reliability. This might be a result of practise effects, where participants learned about the punishments differently across sessions, but not about the rewards. Alternatively, as participants self-determined the loudness of the punishments, differences in volume settings across sessions may have impacted the reliability of this parameter (and indeed other measures). Further, the asymmetric nature of the task may have impacted our ability to estimate the punishment learning rate, as there were fewer occurrences of the punishment compared to the reward. However, given that the sensitivity parameters were more important for explaining how anxiety affected avoidance, the low observed reliability of the punishment learning rate may not present a barrier for future use of this task in repeated-measures designs.

We also did not find any significant correlations between task performance and clinical symptoms of anxiety, which raises questions around the predictive validity of the task. As participants were recruited without imposing specific inclusion/exclusion criteria related to mental health, the low level of clinically-relevant symptoms in the data may have contributed to these results. Further studies in samples with high levels of anxiety or examining individuals at the extreme ends of symptom distributions may offer increased sensitivity in assessing the predictive validity of the task. Further, it is worth noting that many animal paradigms were developed and widely adopted due to their sensitivity to anxiolytic medication (Cryan and Holmes 2005 [↗](#)). Given the lack of associations with clinical measures in our results, it is possible that current translational models of anxiety may not fully capture behaviours that are directly relevant to pathological anxiety. To develop translational paradigms of clinical utility, future research should place a stronger emphasis on assessing their clinical validity in humans. Using multi-level or continuous outcomes would also improve the ecological validity of the present approach and interpretation of the sensitivity indices.

Finally, whilst there is a broad literature on the roles of behavioural inhibition and avoidance tendency traits on decision-making and behaviour (Gray 1982 [↗](#), Carver and White 1994 [↗](#), Corr 2004 [↗](#)), we did not replicate the correlation of experiential avoidance and avoidance responses or the reward-punishment sensitivity index. Since there were also no significant correlations across task performance indices and clinical symptom measures, our findings suggest that the measure may be more sensitive to behaviours relating to state anxiety, rather more stable traits. Nevertheless, how performance in the present task relates to other traits such as behavioural

approach/inhibition tendencies (Carver and White 1994 [↗](#)), as has been found in previous studies on reward/punishment learning (Wise and Dolan 2020 [↗](#), Sharp, Russek et al. 2022) and approach-avoidance conflict (Aupperle, Sullivan et al. 2011), will be an important question for future work.

4 Conclusion

Avoidance is a hallmark symptom of anxiety-related disorders and there is a need for better translational models of avoidance. In our novel translational task designed to involve the same computational processes as those involved in non-human animal tests, anxiety was reliably induced, associated with avoidance behaviour, and linked to a computational model of behaviour. We hope that similar approaches will facilitate progress in the theoretical understanding of avoidance and ultimately aid in the development of novel anxiolytic treatments for anxiety-related disorders.

5 Methods

5.1 Participants

We recruited participants from Prolific (www.prolific.co (<http://www.prolific.co/> [↗](#))). Participants had to be aged 18-60, speak English as their first language, have no language-related disorders/literacy difficulties, have no vision/hearing difficulties/have no mild cognitive impairment or dementia, and be resident in the UK. Participants were reimbursed at a rate of £7.50 per hour, and could earn a bonus of up to an additional £7.50 per hour based on their performance on the task. The study had ethical approval from the University College London Research Ethics Committee (ID 15253/001).

We first recruited a discovery sample to investigate whether anxiety was associated with avoidance in the novel task. We conducted an a priori power analysis to detect a minimally-interesting correlation between anxiety and avoidance of $r = 0.15$ ($\alpha_{\text{two-tailed}} = 0.05$, power = 0.80), resulting in a required sample size of 346. We aimed to replicate our initial finding (that task-induced anxiety was associated with the model-derived measure of approach-avoidance bias) in a replication sample. In the initial sample, the size of this effect was Kendall's tau = 0.099, which approximates to a Pearson's r of 0.15 - this was the basis for the power analysis for the replication sample. Therefore, the power analysis was configured to detect a correlation of $r = 0.15$ ($\alpha_{\text{one-tailed}} = 0.05$, power = 0.95), resulting in a required sample size of $n = 476$. A total of 423 participants completed the initial study (mean age = 38, SD = 10, 59% female), and 726 participants completed the replication study (mean age = 34, SD = 10, 51% female).

5.2 Headphone screen and calibration

We used aversive auditory stimuli as our punishments in the reinforcement learning task, which consisted of combinations of female screams and high-frequency tones that were individually rated as highly aversive in a previous online study (Seow and Hauser 2022 [↗](#)). The stimuli can be found online at <https://osf.io/m3zev/> [↗](#). To maximise their aversiveness, we asked participants to use headphones whilst they completed the task. All participants completed an open-source, validated auditory discrimination task (Woods, Siegel et al. 2017) to screen out those who were unlikely to be using headphones. Participants could only continue to the main task if they achieved an accuracy threshold of five out of six trials in this screening task.

To further increase the aversiveness of the punishments, we wanted each participant to hear them at a loud volume. We initially used a calibration task to increase the system volume and implemented attention checks in the main task to ensure that participants had kept the volume at

a sufficiently high level during the task (see below). Before beginning the calibration task, participants were advised to set their system volume to around 10% of the maximum possible. Then, we presented two alternating auditory stimuli: a burst of white noise which was set to have the same volume as the aversive sounds in the main task, and a spoken number played at a lower volume to one ear ('one', 'two' or 'three'; these stimuli are also available online at <https://osf.io/m3zev/>). We instructed participants to adjust their system volume so that the white noise was loud but comfortable. At the same time, they also had to press the number key on their keyboard corresponding to the number spoken between each white noise presentation. The rationale for the difference in volume was to fix a lower boundary for each participant's system volume, so that the volume of the punishments presented in the main task was locked above this boundary. We used the white noise as a placeholder for the punishments to limit desensitisation. Since the volume of the punishments, and therefore at least in part their aversiveness, was determined by the participants themselves, we conducted sensitivity analyses on all reported findings after accounting for self-reported unpleasantness ratings for the punishment – this did not materially affect our main findings (see the Supplement).

5.3 The approach-avoidance reinforcement learning task

We developed a novel task on which participants learned to simultaneously pursue rewards and avoid punishments. On each trial, participants had 2 s to choose one of two option, depicted as distinct visual stimuli (**Figure 2a**), using the left/right arrow keys on their keyboard. The chosen option was then highlighted by a white border, and participants were presented with the outcome for that trial for 2 s. This could be one of 1) no outcome, 2) a reward (an animated coin image), 3) a punishment (an aversive sound), or 4) both the reward and punishment. We incentivised participants to collect the coins by informing them that they could earn a bonus payment based on the number of coins earned during the task. If no choice was made, the text 'Too slow!' was shown for 2s and the trial was skipped. Participants completed 200 trials which took 10 mins on average across the sample. Each option was associated with separate probabilities of producing the reward/aversive sound outcomes (**Figure 2b**), and these probabilities fluctuated randomly over trials. Both options were associated with the reward, but only one was associated with the punishment (i.e., the 'conflict' option could produce both rewards and punishments, but the 'safe' option could produce only rewards and never punishments). When asked to rate the unpleasantness of the punishments from 'Not at all' to 'Extremely' (scored from 0 to 50, respectively), participants gave a mean rating of 31.7 (SD = 12.8).

The latent outcome probabilities were generated as decaying Gaussian random walks in a similar fashion to previous restless bandit tasks (Daw, O'Doherty et al. 2006), where the probability, p , of observing outcome, a , at option, i , on trial, t , drifted according to the following:

$$p_{i,o}(t+1) = \lambda p_{i,o}(t) + (1 - \lambda)\theta + e$$

$$e \sim N(0, 0.04^2)$$

The decay parameter, λ , was set to 0.97, which pulled all values towards the decay centre θ , which was set to 0.5. The random noise, e , added on each trial was drawn from a zero-mean Gaussian with a standard deviation of 0.04. All probabilities of observing the reward at the safe option were scaled down by a factor of 0.7 so that it was less rewarding than the conflict option, on average. We initialised each walk at the following values: conflict option reward probability – 0.7, safe option reward probability – 0.35, conflict option punishment probability – 0.2. We only used sets of latent outcome probabilities where the cross-correlations across each probability walk was below 0.3. A total of four sets of outcome probabilities were used across our studies – they are available online at <https://osf.io/m3zev/>.

We implemented three auditory attention checks during the main task, to ensure that participants were still using an appropriate volume to listen to the auditory stimuli. As participants could potentially mute their device during the outcome phase of the trial only (i.e. only when the aversive sounds could occur), the timing of these checks had to coincide with moment of outcome presentation. Therefore, we implemented these checks at the outcome phase of two randomly selected trials where the participants' choices led to no outcome (no reward or punishment), and also once at the very end of the task. Specifically, we played one of the spoken numbers used during the calibration task (see above) and the participant had 2.5 s to respond using the corresponding number key on their keyboard.

Immediately after the task and the final attention check, we obtained subjective ratings about the task using the following questions: "How unpleasant did you find the sounds?" and "How anxious did you feel during the task?". Responses were measured on a continuous slider ranging from 'Not at all' (0) to 'Extremely' (50).

5.4 Data cleaning

As the task was implemented online where we could not ensure the same testing standards as we could in-person, we used three exclusion criteria to improve data quality. Firstly, we excluded participants who missed a response to more than one auditory attention check (see above; 8% in both discovery and replication samples) – as these occurred infrequently and the stimuli used for the checks were played at relatively low volume, we allowed for incorrect responses so long as a response was made. Secondly, we excluded those who responded with the same response key on 20 or more consecutive trials (> 10% of all trials; 4%/6% in discovery and replication samples, respectively) – note that as the options randomly switched sides on the screen across trials, this did not exclude participants who frequently and consecutively chose a certain option. Lastly, we excluded those who did not respond on 20 or more trials (1%/2% in discovery and replication samples, respectively). Overall, we excluded 51 out of 423 (12%) in the discovery sample, and 98 out of 725 (14%) in the replication sample. We conducted all the analyses with and without these exclusions (**Supplementary Table 1** [↗](#)). These exclusions had no effects on the main findings reported in the manuscript, and differences before and after excluding data are reported in the Supplement.

5.5 Mixed-effects regressions of trial-by-trial choice

We specified two hierarchical logistic regression models to test the effect of the dynamic outcome probabilities on choice. Both models included predictors aligned with approach and avoidance responses, based on the latent probabilities of observing an outcome given the chosen option at trial t , $P(\text{outcome} \mid \text{option})_t$. First, if individuals were learning to choose maximise their rewards earned, they should generally be choosing the option which was relatively more likely to produce the reward on each trial. This would result in a significant effect of the difference in reward probabilities across options, as given by:

$$\delta P(\text{reward})_t = P(\text{reward} \mid \text{conflict})_t - P(\text{reward} \mid \text{safe})_t.$$

Second, if individuals were learning to avoid punishments, they should generally be avoiding the conflict option on trials when it is likely to produce the punishment. This would result in a significant effect of the difference in punishment probabilities across options, as given by:

$$\delta P(\text{punishment})_t = P(\text{punishment} \mid \text{conflict})_t,$$

where $P(\text{punishment} \mid \text{safe})$ is omitted as it was always 0 across trials.

The first regression model simply included these two predictors and random intercepts varying by participant to predict choice:

$$\text{Choice} \sim \delta P(\text{reward}) + \delta P(\text{punishment}) + (1 \mid \text{participant}).$$

The second model tested the effect of task-induced anxiety on choice, by including anxiety as a main effect and its interactions with $\delta P(\text{reward})_t$ and $\delta P(\text{punishment})_t$:

$$\begin{aligned} \text{Choice} \sim & \delta P(\text{reward}) + \delta P(\text{punishment}) + \text{Anxiety} + \\ & \text{Anxiety} * \delta P(\text{reward}) + \text{Anxiety} * \delta P(\text{punishment}) + (1 \mid \text{participant}). \end{aligned}$$

5.6 Computational modelling using reinforcement learning

5.6.1 Model specification

We used variations of classical reinforcement learning algorithms (Sutton and Barto 2018 [↗](#)), specifically the ‘Q-learning’ algorithm (Watkins and Dayan 1992 [↗](#)), to model participants’ choices in the main task (**Table 1** [↗](#)). Briefly, our models assume that participants estimate the probabilities of observing a reward and punishment separately for each punishment, and continually update these estimates after observing the outcomes of their actions. Consequently, participants can then use these estimates to choose the option which will maximise subjective value, for example by choosing the option which is more likely to produce a reward, or avoiding an option if it is very likely to produce a punishment. By convention, we refer to the probability estimates as ‘Q-values’; for example, the estimated probability of observing a reward after choosing the conflict option is written as $Q^r(\text{conflict})$, whereas the probability of observing a punishment after choosing the safe option would be written as $Q^p(\text{safe})$.

In all models, the probability estimates of observing outcome, $a = \{r, p\}$, of the chosen option, $a = \{\text{conflict}, \text{safe}\}$, on trial, t , were updated via:

$$Q_{t+1}^o(a) = Q_t^o(a) + \alpha \cdot [o^t - Q_t^o(a)],$$

where the learning rate is controlled by $\alpha \in (0,1)$. In some models (**Table 1** [↗](#)), the learning rate was split into reward- and punishment-specific parameters, α^r and α^p .

On each trial, the probability estimates for observing rewards and aversive sounds are integrated into action weights, W , via:

$$W = \beta \cdot (Q^r - Q^p),$$

where individual differences in sensitivity to the outcomes (how much the outcomes affect value-based choice) are parameterised by the outcome sensitivity parameter, β . In models with a single sensitivity parameter, participants were assumed to value the reward and punishment equally, such that obtaining a reward had the same subjective value as avoiding a punishment. However, as individual differences in approach-avoidance conflict could manifest in asymmetries in the valuation of reward/punishment, we also specified models (**Table 1** [↗](#)) where β was split into reward- and punishment-specific parameters, β^r and β^p :

$$W = \beta^r \cdot Q^r - \beta^p \cdot Q^p.$$

Model	Parameters			
Reward learning only	α^r		β^r	
Punishment learning only	α^p		β^p	
Symmetrical learning	α		β	
Asymmetric learning rates	α^r	α^p	β	
Asymmetric sensitivities	α		β^r	β^p
Asymmetric learning rates and sensitivities	α^r	α^p	β^r	β^p

Table 1.

Model specification

The action weights are then the basis on which a choice was made between the options, which was modelled with the softmax function:

$$P(a) = \frac{W(a)}{\sum_i W(i)}.$$

We also tested all models with the inclusion of a lapse term, which allows for choice stochasticity to vary across participants, but since this did not improve model fit for any model, we do not discuss this further.

All custom model scripts are available online at <https://osf.io/m3zev/>.

5.6.2 Model fitting and comparison

We fit the models using a hierarchical expectation-maximisation algorithm (Huys, Cools et al. 2011) – the algorithm and supporting code is available online at www.quentinhuys.com/pub/emfit (<http://www.quentinhuys.com/pub/emfit>). All prior distributions of the untransformed parameter values were set as Gaussian distributions to regularise fitting and to limit extreme values. Parameters were transformed within the models, via the sigmoid function for parameters constrained to values between 0 and 1; or exponentiated for strictly positive parameters. We compared evidence across models using integrated Bayesian Information Criterion (iBIC; Huys, Cools et al. 2011) which integrates both model likelihood and model complexity into a single measure for model comparison.

5.7 Mediation analysis of anxiety-related avoidance

To investigate the cognitive processes underlying anxiety-related avoidance, we tested for potential mediators of the effect of task-induced anxiety on avoidance choices in the approach-avoidance reinforcement learning task. Mediation was assessed using structural equation modelling with the ‘lavaan’ package in R (Rosseel 2012). This allowed us to test multiple potential mediators within a single model. We used maximum-likelihood estimation with and without bootstrapped standard errors (using 2000 bootstrap samples) for estimation. The potential mediators were any parameters from our computational model that significantly correlated with task-induced anxiety. Therefore, the model included the following variables: task-induced anxiety, avoidance choices (computed as the proportion of safe option choices for each participant), and model parameters which correlated with task-induced anxiety. The model was constructed such that the model parameters were parallel mediators of the effect of task-induced anxiety on avoidance choices. The model parameters were allowed to covary across one another.

5.8 Symptom questionnaires

Psychiatric symptoms were measured using the Generalised Anxiety Disorder scale (GAD-7; Spitzer, Kroenke et al. 2006), Patient Health Questionnaire depression scale (PHQ-8; Kroenke, Strine et al. 2009) and the Brief Experiential Avoidance Questionnaire (BEAQ; Gamez, Chmielewski et al. 2014).

5.9 Test-retest reliability analysis

We retested a subset of our sample ($n = 57$) to assess the reliability of our task. We determined our sample size via *ana priori* power analysis using GPower (Faul, Erdfelder et al. 2007), which was set to detect the smallest meaningful effect size of $r = 0.4$, as any lower test-retest correlations are considered poor by convention (Fleiss 1981, Cicchetti 1994). The required sample size was thus 50 participants, based on a one-tailed significance threshold of 0.05 and 90% power. Participants were recruited from Prolific. We invited participants from the replication study (providing they had not excluded on the basis of data quality) to complete the study again, at least

10 days after their first session. The retest version of the approach-avoidance reinforcement learning task was identical to that in the first session, except for the use of new stimuli to represent the conflict and safe options and different latent outcome probabilities to limit practise effects.

We calculated intra-class correlations coefficients to assess the reliability of our model-agnostic measures. Specifically, we used consistency-based ICC(3,1) (following the notation of Shrout and Fleiss 1979 [\[4\]](#)). To estimate the reliability of the model-derived measures, we estimated individual- and session-level model parameters via a joint model in an approach recently shown to improve the accuracy of estimating reliability (Waltmann, Schlagenhauf et al. 2022). Briefly, this involved fitting behavioural data from both test and retest sessions per participant in a single model, but specifying unique parameters per session. Since our fitting algorithm assumes parameters are drawn from a multivariate Gaussian distribution, information about the reliability of the parameters can be derived from the covariance of each parameter across time. These covariance values can subsequently be converted to Pearson's r values, and thus indices of reliability.

6 Data and code availability

The data and code required to replicate the data analyses are available online at <https://osf.io/m3zev/> [\[4\]](#).

7 Competing interests

YY has no financial or non-financial interests to declare. OJR's MRC senior fellowship is partially in collaboration with Cambridge Cognition Ltd (who plan to provide in-kind contribution) and he is running an investigator-initiated trial with medication donated by Lundbeck (escitalopram and placebo, no financial contribution). He also holds an MRC-Proximity to discovery award with Roche (who provide in-kind contributions and have sponsored travel for ACP) regarding work on heart-rate variability and anxiety. He also has completed consultancy work for Peak, IESO digital health, Roche and BlackThorn therapeutics. OJR sat on the committee of the British Association of Psychopharmacology until 2022. In the past 3 years, JPR has held a PhD studentship co-funded by Cambridge Cognition and performed consultancy work for GH Research and GE Ltd. These disclosures are made in the interest of full transparency and do not constitute a conflict of interest with the current work.

9 Supplementary material

9.1 Comparison of results in the discovery and replication samples, with and without data cleaning exclusions

A sensitivity analyses of the effect of implementing or not implementing our data cleaning exclusions showed that the significant and negative correlation between experiential avoidance symptoms and overall proportion of conflict option choices ($\tau = -0.059$, $p = 0.010$) was not significant when the data cleaning exclusions were not implemented ($\tau = -0.029$, $p = 0.286$), however since this effect was not replicated in the independent sample (with and without data exclusions), we do not discuss this further. The mediating effect of the punishment learning rate was also significant in the replication sample only before excluding data, however the effect sizes were small and similar before and after excluding data, and so we also do not discuss this further.

Statistical test/model	Sub-test	Sample (without data cleaning)	
		Discovery	Discovery
Hierarchical logistic regression; no task-induced anxiety	Reward coefficient	$\beta = 0.98 \pm 0.03, p < 0.001$ $\beta = 0.95 \pm 0.03, p < 0.001$	$\beta = 0.95 \pm 0.02, p < 0.001$ $\beta = 0.90 \pm 0.02, p < 0.001$
	Punishment coefficient	$\beta = -0.33 \pm 0.04, p < 0.001$ $\beta = -0.29 \pm 0.04, p < 0.001$	$\beta = -0.52 \pm 0.03, p < 0.001$ $\beta = -0.50 \pm 0.03, p < 0.001$
Distribution; task-induced anxiety		Mean = 21, SD = 14	Mean = 22, SD = 13
		Mean = 21, SD = 14	mean = 22, SD = 14
Correlation; task-induced anxiety and choice		$\tau = -0.074, p = 0.033$ $\tau = -0.068, p = 0.036$	$\tau = -0.075, p = 0.005$ $\tau = -0.081, p = 0.002$
Hierarchical logistic regression; with task-induced anxiety	Reward coefficient	$\beta = 0.98 \pm 0.03, p < 0.001$ $\beta = 0.95 \pm 0.03, p < 0.001$	$\beta = 0.95 \pm 0.02, p < 0.001$ $\beta = 0.90 \pm 0.02, p < 0.001$
	Punishment coefficient	$\beta = -0.33 \pm 0.04, p < 0.001$ $\beta = -0.29 \pm 0.04, p < 0.001$	$\beta = -0.52 \pm 0.03, p < 0.001$ $\beta = -0.50 \pm 0.03, p < 0.001$
	Task-induced anxiety coefficient	$\beta = -0.04 \pm 0.04, p = 0.304$ $\beta = -0.03 \pm 0.04, p = 0.377$	$\beta = 0.02 \pm 0.03, p = 0.637$ $\beta = 0.003 \pm 0.03, p = 0.925$
	Interaction of task-induced anxiety and reward	$\beta = -0.06 \pm 0.03, p = 0.074$ $\beta = -0.05 \pm 0.03, p = 0.076$	$\beta = -0.003 \pm 0.02, p = 0.901$ $\beta = -0.006 \pm 0.02, p = 0.774$
	Interaction of task-induced anxiety and punishment	$\beta = -0.10 \pm 0.04, p = 0.022$ $\beta = -0.10 \pm 0.04, p = 0.011$	$\beta = -0.23 \pm 0.03, p < 0.001$ $\beta = -0.23 \pm 0.03, p < 0.001$
Hierarchical logistic regression; with task-induced anxiety, no interaction	Reward coefficient	$\beta = 0.99 \pm 0.03, p < 0.001$ $\beta = 0.95 \pm 0.03, p < 0.001$	$\beta = 0.95 \pm 0.02, p < 0.001$ $\beta = 0.89 \pm 0.02, p < 0.001$
	Punishment coefficient	$\beta = -0.33 \pm 0.04, p < 0.001$ $\beta = -0.29 \pm 0.04, p < 0.001$	$\beta = -0.52 \pm 0.03, p < 0.001$ $\beta = -0.50 \pm 0.03, p < 0.001$
	Task-induced anxiety coefficient	$\beta = -0.09 \pm 0.03, p = 0.012$ $\beta = -0.08 \pm 0.03, p = 0.016$	$\beta = -0.08 \pm 0.03, p = 0.005$ $\beta = -0.10 \pm 0.03, p < 0.001$
Computational model comparison	Winning model	2 learning rate, 2 sensitivity 2 learning rate, 2 sensitivity	2 learning rate, 2 sensitivity 2 learning rate, 2 sensitivity
Correlation; task-induced anxiety and model parameters	Reward learning rate	$\tau = -0.019, p = 0.596$ $\tau = -0.006, p = 0.847$	$\tau = -0.010, p = 0.749$ $\tau = -0.012, p = 0.637$
	Punishment learning rate	$\tau = -0.088, p = 0.015$ $\tau = -0.073, p = 0.026$	$\tau = -0.064, p = 0.019$ $\tau = -0.074, p = 0.003$
	Learning rate ratio	$\tau = 0.048, p = 0.175$ $\tau = -0.037, p = 0.276$	$\tau = 0.029, p = 0.282$ $\tau = 0.045, p = 0.072$
	Reward sensitivity	$\tau = -0.038, p = 0.286$ $\tau = -0.047, p = 0.149$	$\tau = -0.028, p = 0.285$ $\tau = -0.023, p = 0.345$
	Punishment sensitivity	$\tau = 0.068, p = 0.051$	$\tau = 0.076, p = 0.004^*$

Supplementary Table 1.

Statistical results across the discovery and replication samples, and the effect of data cleaning exclusions

		$\tau = 0.047, p = 0.153$	$\tau = 0.094, p < 0.001$
	Reward-punishment sensitivity index	$\tau = -0.099, p = 0.005$ $\tau = -0.084, p = 0.011$	$\tau = -0.096, p < 0.001$ $\tau = -0.103, p < 0.001$
Mediation	Punishment learning rate	$\beta = -0.002 \pm 0.001, p = 0.052$ $\beta = -0.001 \pm 0.001, p = 0.222$	$\beta = -0.001 \pm 0.003, p = 0.132$ $\beta = -0.001 \pm 0.003, p = 0.031^{\dagger}$
	Reward-punishment sensitivity index	$\beta = 0.009 \pm 0.003, p = 0.003$ $\beta = 0.006 \pm 0.002, p = 0.011$	$\beta = 0.009 \pm 0.002, p < 0.001$ $\beta = 0.009 \pm 0.002, p < 0.001$
Correlation; psychiatric symptoms and choice (overall proportion of conflict option choices)	GAD7	$\tau = -0.026, p = 0.458$ $\tau = 0.005, p = 0.894$	$\tau = -0.001, p = 0.988$ $\tau = -0.002, p = 0.919$
	PHQ8	$\tau = -0.02, p = 0.579$ $\tau = 0.014, p = 0.684$	$\tau = -0.013, p = 0.639$ $\tau = -0.012, p = 0.646$
	BEAQ	$\tau = -0.059, p = 0.010$ $\tau = -0.048, p = 0.151^{\dagger}$	$\tau = -0.029, p = 0.286^*$ $\tau = -0.020, p = 0.423$
Correlation; psychiatric symptoms and task-induced anxiety	GAD7	$\tau = 0.256, p < 0.001$ $\tau = 0.267, p < 0.001$	$\tau = 0.222, p < 0.001$ $\tau = 0.231, p < 0.001$
	PHQ8	$\tau = 0.233, p < 0.001$ $\tau = 0.244, p < 0.001$	$\tau = 0.184, p < 0.001$ $\tau = 0.194, p < 0.001$
	BEAQ	$\tau = 0.15, p < 0.001$ $\tau = 0.172, p < 0.001$	$\tau = 0.176, p < 0.001$ $\tau = 0.17, p < 0.001$
Correlation; GAD7 and model parameters	Reward learning rate	$\tau = -0.06, p = 0.076$ $\tau = -0.061, p = 0.074$	$\tau = -0.03, p = 0.304$ $\tau = -0.028, p = 0.264$
	Punishment learning rate	$\tau = -0.07, p = 0.077$ $\tau = -0.037, p = 0.265$	$\tau = -0.01, p = 0.621$ $\tau = -0.017, p = 0.534$
	Learning rate ratio	$\tau = -0.001, p = 0.999$ $\tau = -0.032, p = 0.35$	$\tau = -0.034, p = 0.229$ $\tau = -0.026, p = 0.301$
	Reward sensitivity	$\tau = -0.01, p = 0.802$ $\tau = -0.005, p = 0.877$	$\tau = -0.01, p = 0.745$ $\tau = -0.009, p = 0.718$
	Punishment sensitivity	$\tau = 0.05, p = 0.154$ $\tau = 0.022, p = 0.513$	$\tau = 0.003, p = 0.907$ $\tau = 0.012, p = 0.633$
	Reward-punishment sensitivity index	$\tau = -0.05, p = 0.149$ $\tau = -0.023, p = 0.496$	$\tau = 0.01, p = 0.746$ $\tau = 0.007, p = 0.797$
Correlation; PHQ8 and model parameters	Reward learning rate	$\tau = -0.03, p = 0.396$ $\tau = -0.04, p = 0.239$	$\tau = -0.03, p = 0.219$ $\tau = -0.035, p = 0.17$
	Punishment learning rate	$\tau = -0.06, p = 0.100$ $\tau = -0.022, p = 0.504$	$\tau = -0.01, p = 0.630$ $\tau = -0.025, p = 0.348$
	Learning rate ratio	$\tau = 0.008, p = 0.809$ $\tau = -0.038, p = 0.259$	$\tau = -0.033, p = 0.231$ $\tau = -0.016, p = 0.521$
	Reward sensitivity	$\tau = -0.012, p = 0.610$ $\tau = -0.009, p = 0.801$	$\tau = 0.01, p = 0.729$ $\tau = 0.004, p = 0.901$
	Punishment sensitivity	$\tau = 0.05, p = 0.179$ $\tau = 0.008, p = 0.802$	$\tau = -0.002, p = 0.948$ $\tau = 0.012, p = 0.619$
	Reward-punishment sensitivity index	$\tau = -0.06, p = 0.123$ $\tau = -0.01, p = 0.773$	$\tau = 0.02, p = 0.557$ $\tau = 0.005, p = 0.867$
	Reward learning rate	$\tau = -0.06, p = 0.085$	$\tau = -0.02, p = 0.394$

Supplementary Table 1. (continued)

Correlation; BEAQ and model parameters		$\tau = -0.047, p = 0.147$	$\tau = -0.017, p = 0.503$
	Punishment learning rate	$\tau = -0.08, p = 0.024$ $\tau = -0.071, p = 0.032$	$\tau = -0.03, p = 0.337^*$ $\tau = -0.031, p = 0.228$
	Learning rate ratio	$\tau = 0.018, p = 0.618$ $\tau = 0.002, p = 0.938$	$\tau = -0.005, p = 0.883$ $\tau = 0.007, p = 0.759$
	Reward sensitivity	$\tau = -0.01, p = 0.739$ $\tau = -0.036, p = 0.29$	$\tau = 0.01, p = 0.753$ $\tau = 0.002, p = 0.927$
	Punishment sensitivity	$\tau = 0.07, p = 0.061$ $\tau = 0.05, p = 0.127$	$\tau = 0.02, p = 0.477$ $\tau = 0.025, p = 0.308$
	Reward-punishment sensitivity index	$\tau = 0.02, p = 0.034$ $\tau = -0.073, p = 0.028$	$\tau = -0.01, p = 0.745^*$ $\tau = -0.015, p = 0.548$
Note. All correlations tests were conducted using permutation-based Kendall's tau correlations, using 10,000 permutations per test. Symbols represent the following: * - Inconsistent findings across discovery replication samples. † - Inconsistent findings before and after data cleaning exclusions. Abbreviations: GAD7 – Generalised Anxiety Disorder 7-item scale; PHQ8 – Patient Health Questionnaire 8-item depression scale; BEAQ – Brief Experiential Avoidance Questionnaire.			

Supplementary Table 1. (continued)

9.2 Computational modelling

9.3 Mediation analyses

9.4 Parameter recovery

To establish how well we could recover our parameters from the best fitting model, we simulated data from each study (discovery, replication) and refitted these new datasets with the identical model fitting procedure used for the empirical data. We matched the number of participants, number of trials per participant, and latent outcome probabilities in the simulated datasets to the empirical data, and used the best-fitting parameters for each participant to generate their choices. Trial-by-trial outcomes were generated stochastically according to the latent outcome probabilities, which meant that the simulated participants did not necessarily receive the same series of outcomes to our real participants. To average over this stochasticity in the outcomes, we repeated the simulations 100 times for each study. After fitting our best-fitting model to the simulated data, we computed Pearson's correlation coefficients across the data-generating parameters and the recovered parameters. Across all parameters, the Pearson's r values between data-generating and recovered parameters ranged from 0.77 to 0.86, indicating excellent recovery (**Supplementary Figure 4** [↗](#)).

9.5 Split-half reliability

9.6 Test-retest reliability

9.7 Practise effects

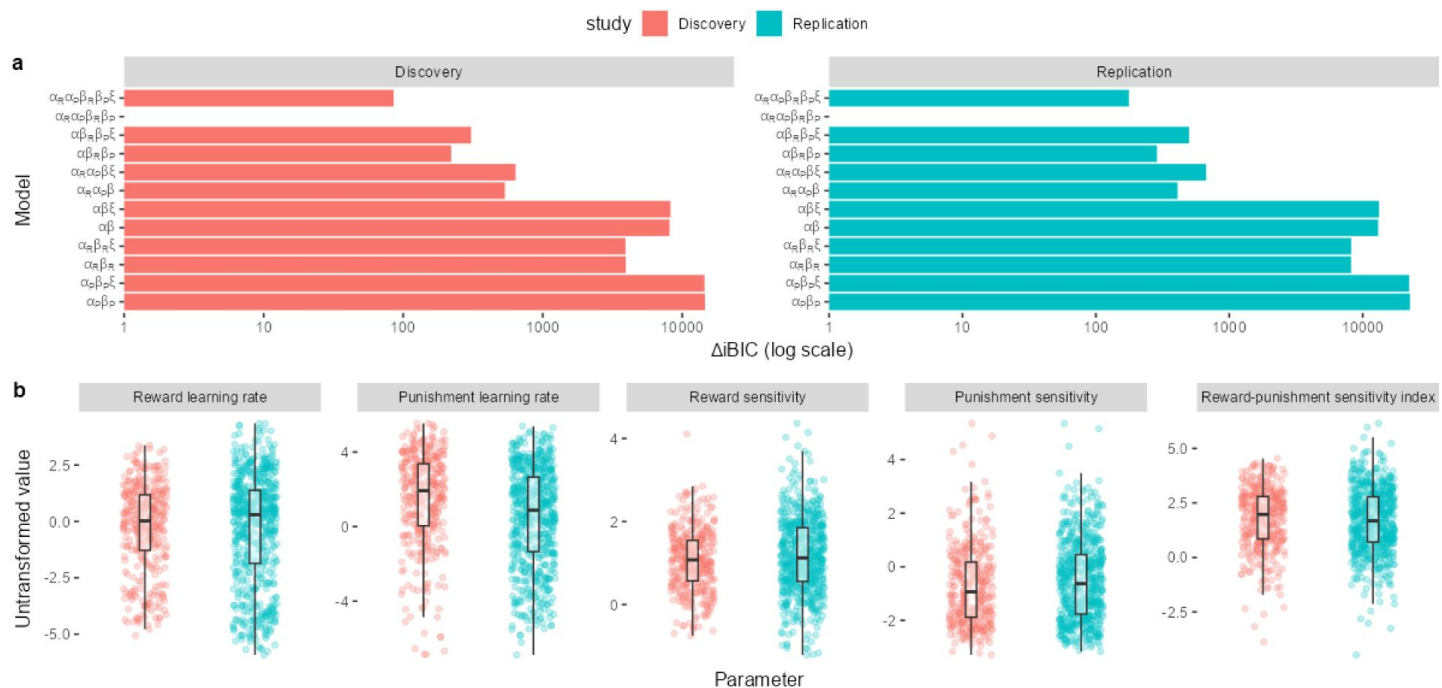
We checked for practise effects in the behavioural data by comparing the behavioural measures and computational model parameters across sessions via paired t-tests. These tests were all non-significant ($p > 0.07$) except for task-induced anxiety ($t_{56} = 2.21$, $p = 0.031$) and the punishment learning rate ($t_{56} = 2.24$, $p = 0.029$; **Supplementary Figure 7** [↗](#)).

9.8 Sensitivity analysis: estimating parameters via expectation maximisation and variational Bayesian inference algorithms

Given that the expectation maximisation (EM) algorithm produced high inter-parameter correlations, we ran a sensitivity analysis by assessing the robustness of our computational findings to an alternative method of parameter estimation – (mean-field) variational Bayesian inference (VBI) via Stan ([Stan Development Team 2023](#) [↗](#)). Since, unlike EM, the results of VBI are very sensitive to initial values, we fitted the data 10 times with different initial values.

Inter-parameter correlations

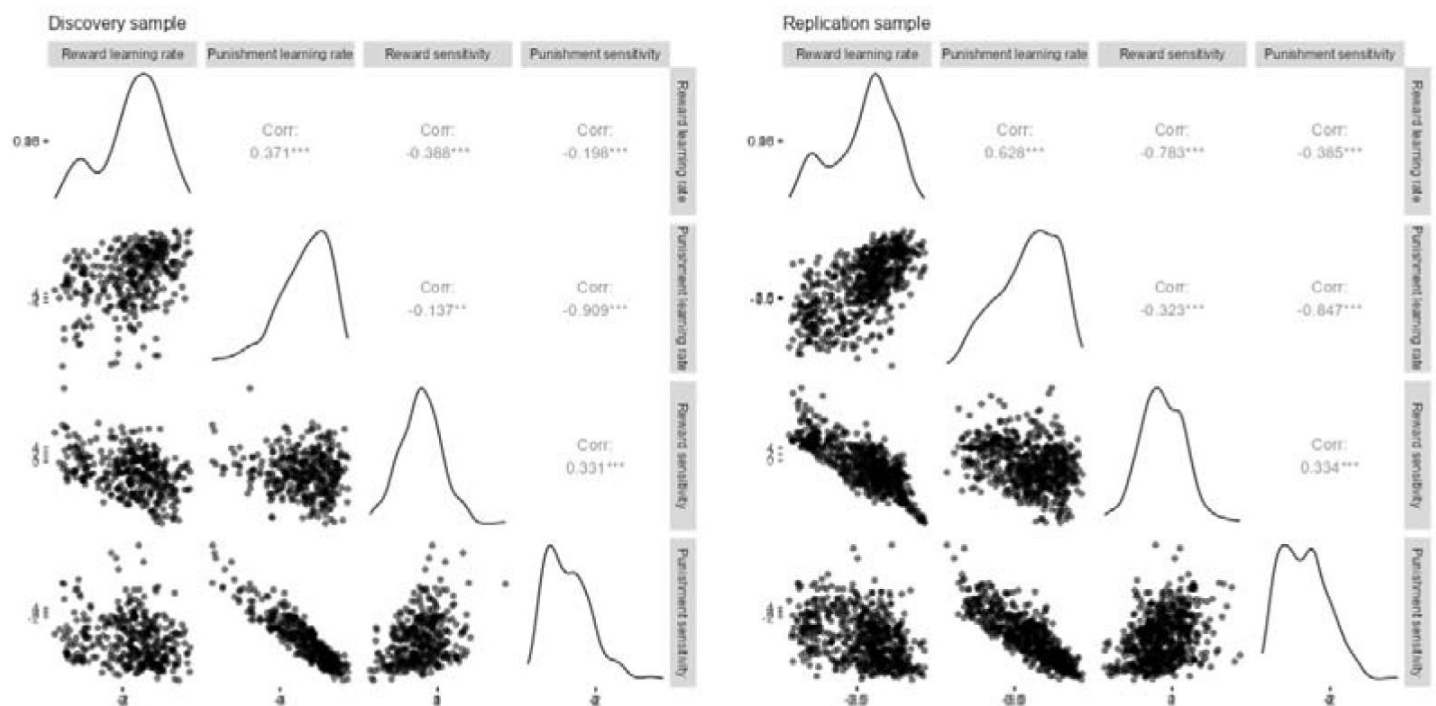
The VBI produced lower inter-parameter correlations than the EM algorithm (**Supplementary Figure 8** [↗](#)).



Supplementary Figure 1.

Model comparison and parameter distributions across studies.

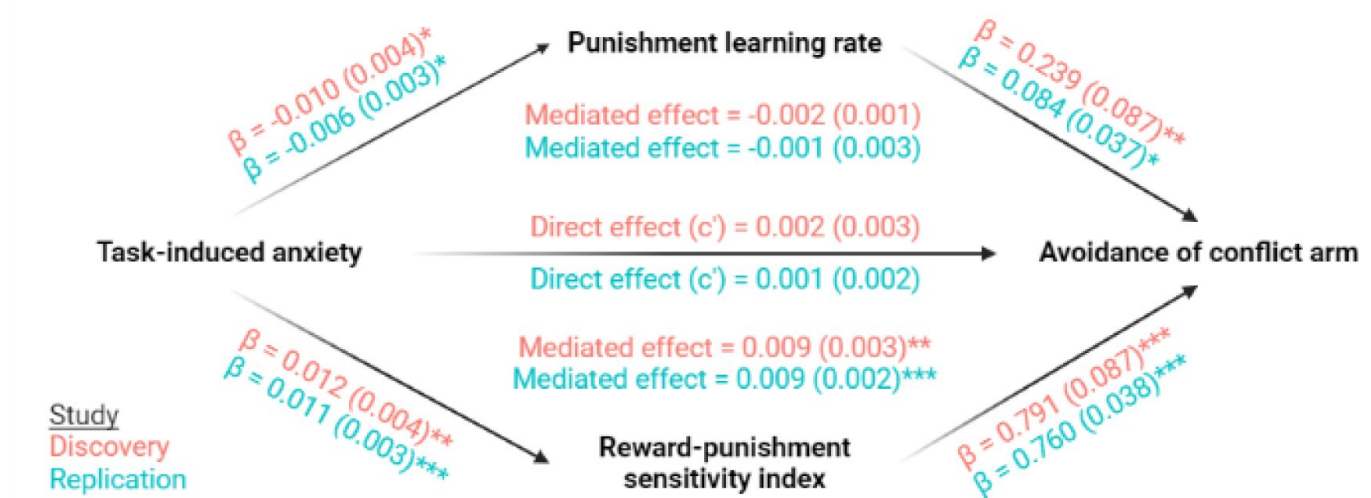
a, Model comparison results. The difference in integrated Bayesian Information Criterion scores from each model relative to the winning model is indicated on the x-axis. The winning model in both studies included specific learning rates for reward (α^R) and punishment learning (α^P), and specific outcome sensitivity parameters for reward (β^R) and punishment (β^P). Some models were tested with the inclusion of a lapse term (ξ). **b**, Distributions of individual parameter values from the winning model across studies. The reward-punishment sensitivity index constituted our computational measure of approach-avoidance bias, calculated by taking the ratio between the reward and punishment sensitivity parameters.



Supplementary Figure 2.

Correlation matrices for the estimated parameters across studies.

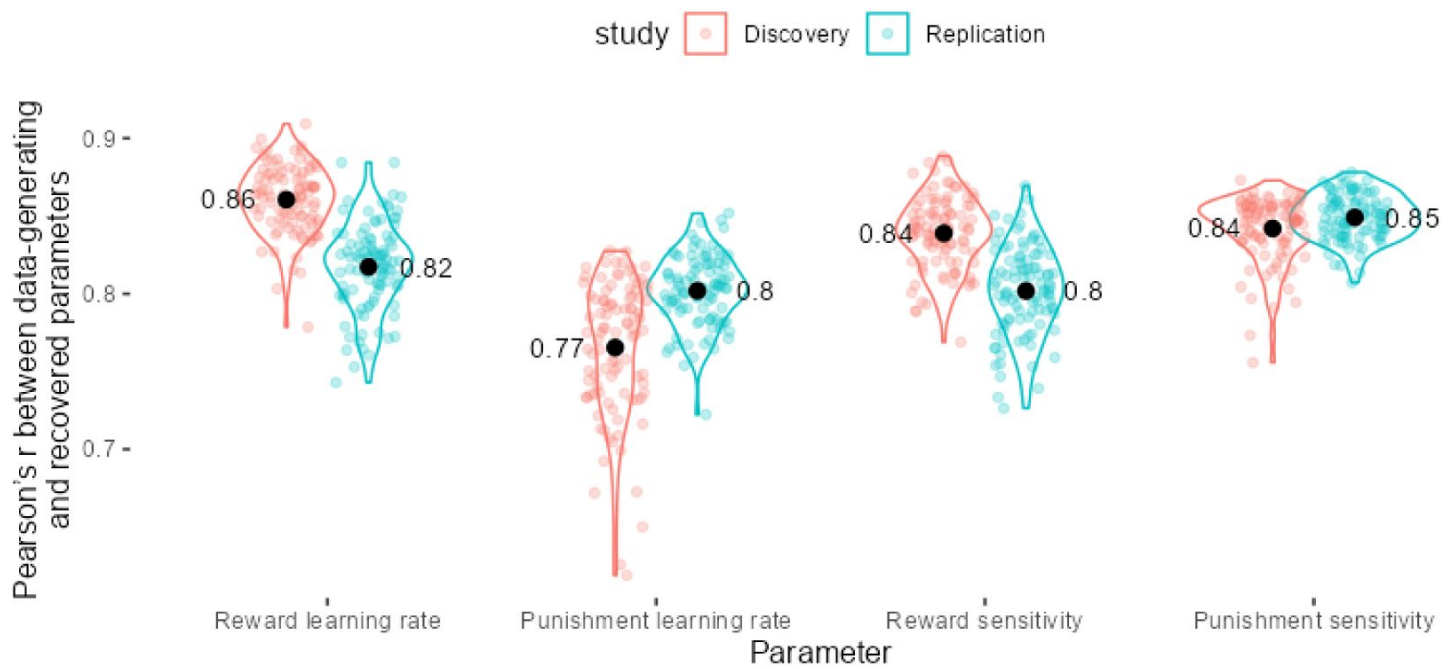
Lower-right diagonal of each matrix shows a scatterplot of cross-parameter correlations. Upper-right diagonal denotes the Pearson's r correlation coefficients for each pair of parameters, based on the untransformed parameter values.



Supplementary Figure 3.

Mediation analyses across studies.

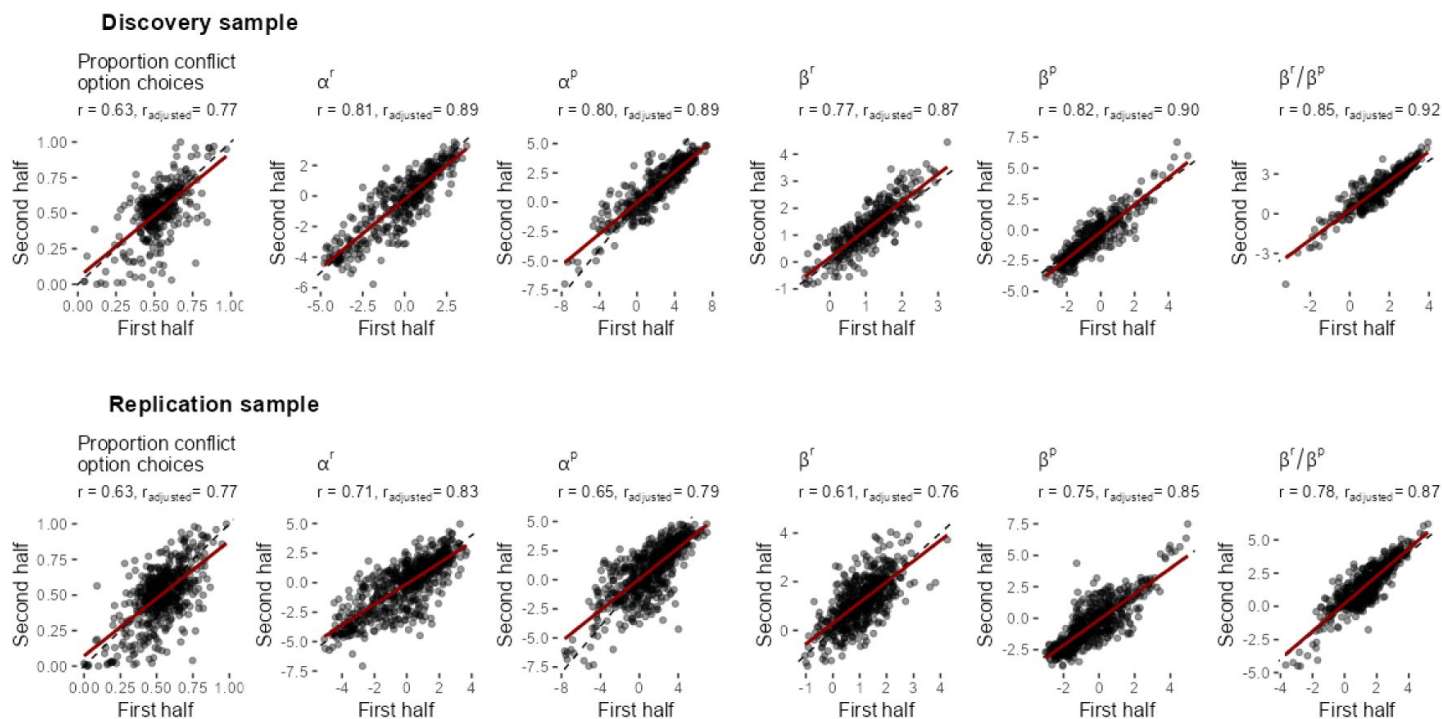
Mediation effects were assessed using structural equation modelling. Bold terms represent variables and arrows depict regression paths in the model. The annotated values next to each arrow show the regression coefficient associated with that path, denoted as *coefficient (standard error)*. Only the reward-punishment sensitivity index significantly mediated the effect of task-induced anxiety on avoidance. Significance levels in all figures are shown according to the following: $p < 0.05$ - *; $p < 0.05$ - **; $p < 0.001$ - ***.



Supplementary Figure 4.

Parameter recovery.

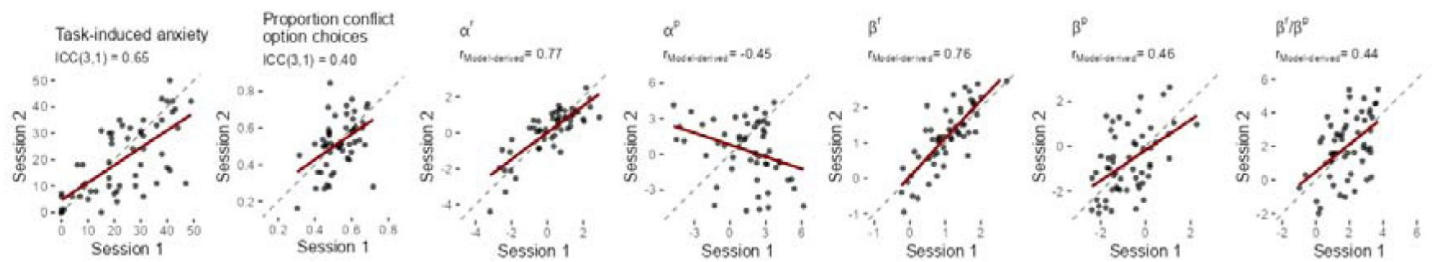
Pearson's r values across the data-generating and recovered parameters by parameter. Coloured points represent Pearson's r values for each of 100 simulation iterations, and black points represent the mean value across simulations.



Supplementary Figure 5.

Split-half reliability of the task.

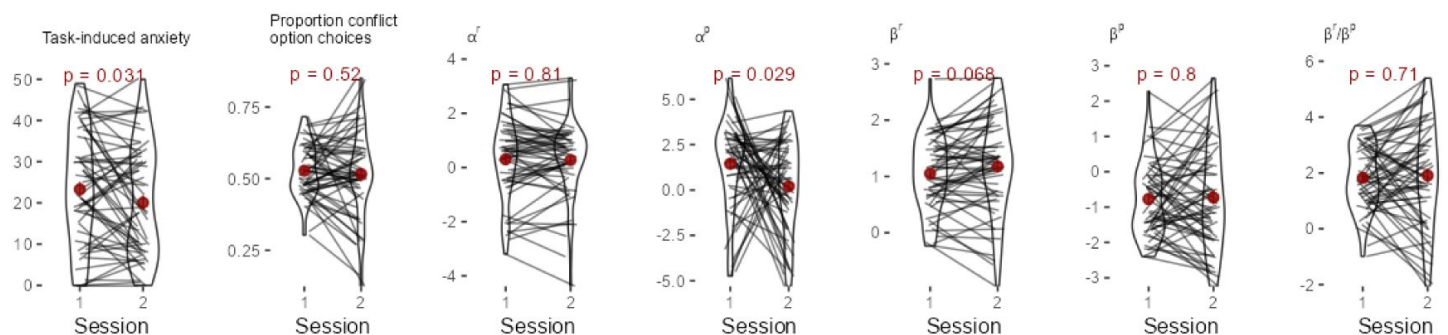
Reliability via Pearson's correlations of measures calculated from the first and second halves of the task are shown with their estimates of reliability. Reliability estimates for the computational measures from the winning computational model were computed by fitting split-half parameters within a single model, then using the parameter covariance matrix to derive Pearson's correlation coefficients for each parameter across halves. Reliability estimates are reported as unadjusted values (r) and after adjusting for reduced number of trials via Spearman-Brown correction (r_{adjusted}). Dotted lines represent the reference line, indicating perfect correlation. Red lines show lines-of-best-fit.



Supplementary Figure 6.

Test-retest reliability of the task.

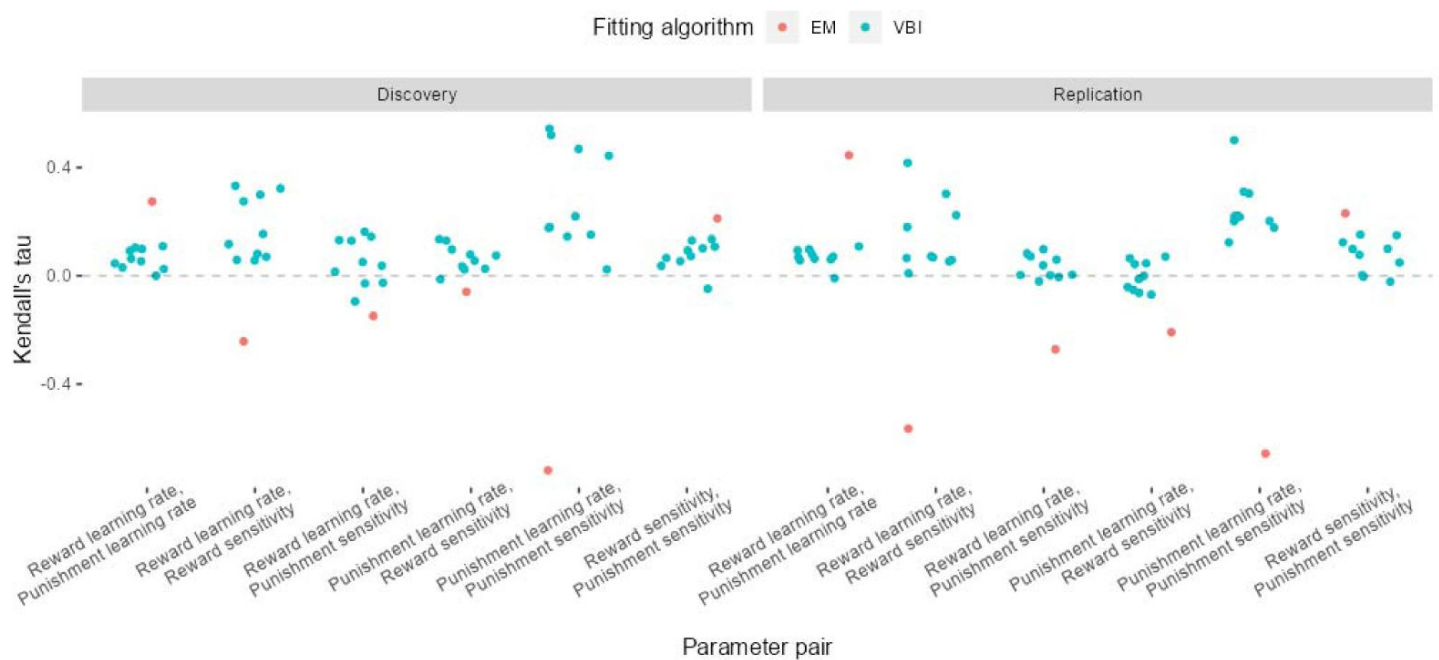
Correlations of measures across the test and retest sessions are shown with their estimates of reliability. Reliability estimates for the model-agnostic measures (task-induced anxiety, proportion of conflict option choices) were estimated using intra-class correlation coefficients. Reliability estimates for the computational measures from the winning computational model were computed by first fitting both sessions' parameters within a single model, then using the parameter covariance matrix to derive a Pearson's correlation coefficient ($r_{\text{Model-derived}}$) for each parameter across sessions to be calculated from their covariance. Dotted lines represent the reference line, indicating perfect correlation. Red lines show lines-of-best-fit.



Supplementary Figure 7.

Task practice effects.

Comparison of behavioural measures and model parameters across time. Lines represent individual data, red points represent mean values, and red lines represent standard error bars. P-values of paired t-tests are annotated above each plot. Task-induced anxiety and the punishment learning rate was significantly lower in the second session, whilst the other measures did not change significantly across sessions.



Supplementary Figure 8.

Inter-parameter correlations across the expectation-maximisation (EM, red) and variational Bayesian inference (VBI, blue) algorithms.

Overall, the VBI algorithm produced lower correlations compared to EM.

Sensitivity analysis

Since multicollinearity in the VBI-estimated parameters was lower than for EM, indicating less trade-off in the estimation, we re-tested our computational findings from the manuscript as part of a sensitivity analysis. We first assessed whether we observed the same correlations between task-induced anxiety and punishment learning, and reward-punishment sensitivity index (**Supplementary Figure 9a** [↗](#)). Punishment learning rate was not significantly associated with task-induced anxiety in any of the 10 VBI iterations in the discovery sample, although it was in 9/10 in the replication sample. On the other hand, the reward-punishment sensitivity index was significantly associated with task-induced anxiety in 9/10 VBI iterations in the discovery sample and all iterations in the replication sample. This suggests that the correlation of anxiety and sensitivity index is robust to these two fitting approaches.

We also re-estimated the mediation models, where in the EM-estimated parameters, we found that the reward-punishment sensitivity index mediated the relationship between task-induced anxiety and task choice proportions (**Supplementary Figure 9b** [↗](#)). Again, we found that the reward-punishment sensitivity index was a significant mediator in 9/10 VBI iterations in the discovery sample and all iterations in the replication sample. Punishment learning rate was also a significant mediator in 9/10 iterations in the replication sample, although it was not in the discovery sample for all iterations, and this was not observed for the EM-estimated parameters.

Overall, we found that our key results, that anxiety is associated with greater sensitivity to punishment over reward, and this mediates the relationship between anxiety and approach-avoidance behaviour, were robust across both fitting methods.

9.9 Sensitivity analysis: punishment unpleasantness

Distribution of unpleasantness

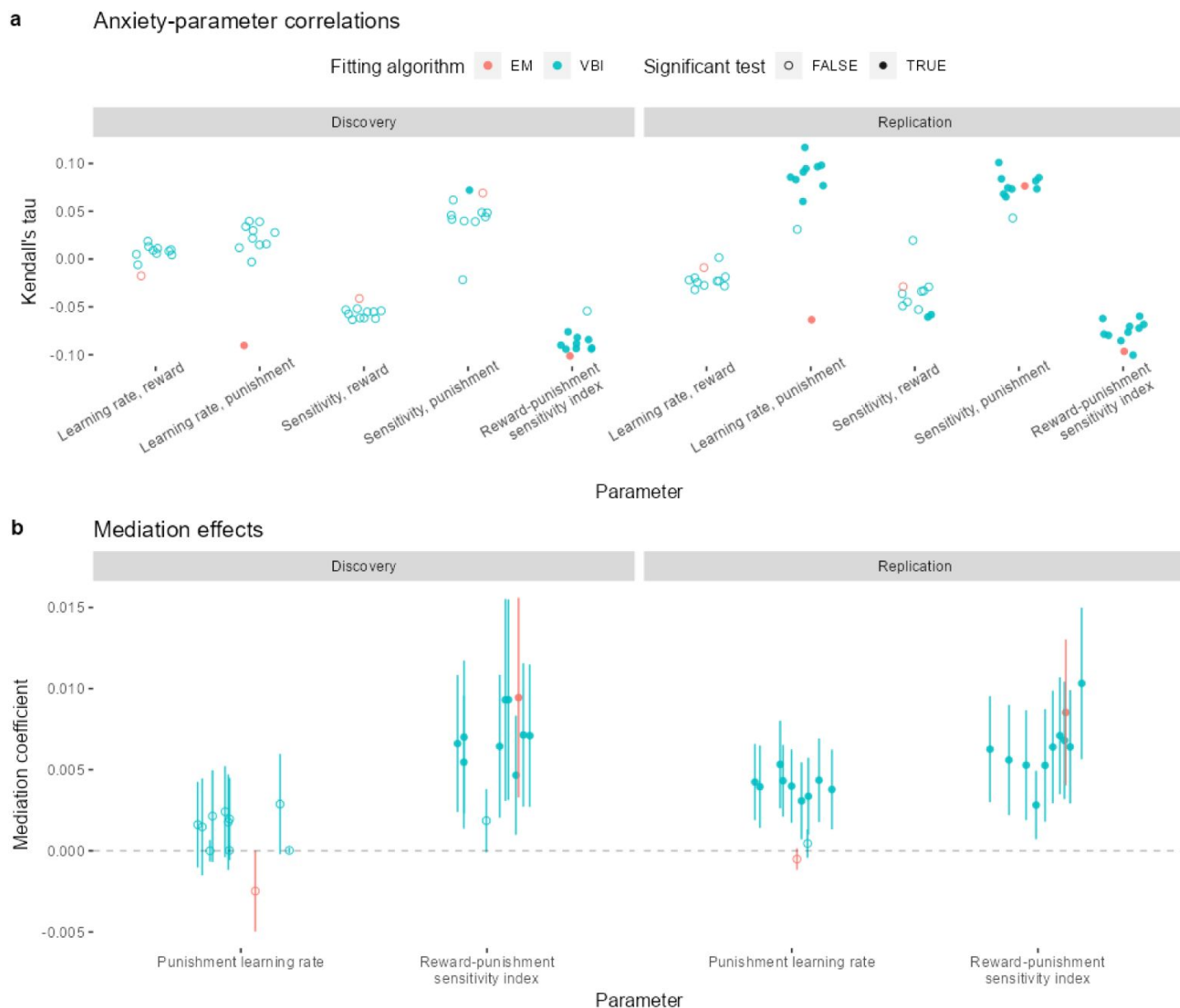
The punishments were rated as unpleasant by the participants, on average (discovery sample: mean rating = 31.1 [scored between 0 and 50], SD = 13.1; replication sample: mean rating = 32.1, SD = 12.7; **Supplementary Figure 10** [↗](#)).

Approach-avoidance hierarchical logistic regression model

We assessed whether approach and avoidance responses, and their relationships with state anxiety, were impacted by punishment unpleasantness, by including unpleasantness ratings as a covariate into the hierarchical logistic regression model. Whilst unpleasantness was a significant predictor of choice (positively predicting safe option choices), all significant predictors and interaction effects from the model without unpleasantness ratings survived (**Supplementary Figure 11** [↗](#)). Critically, this suggests that punishment unpleasantness does not account for all of the variance in the relationship between anxiety and avoidance.

Mediation model

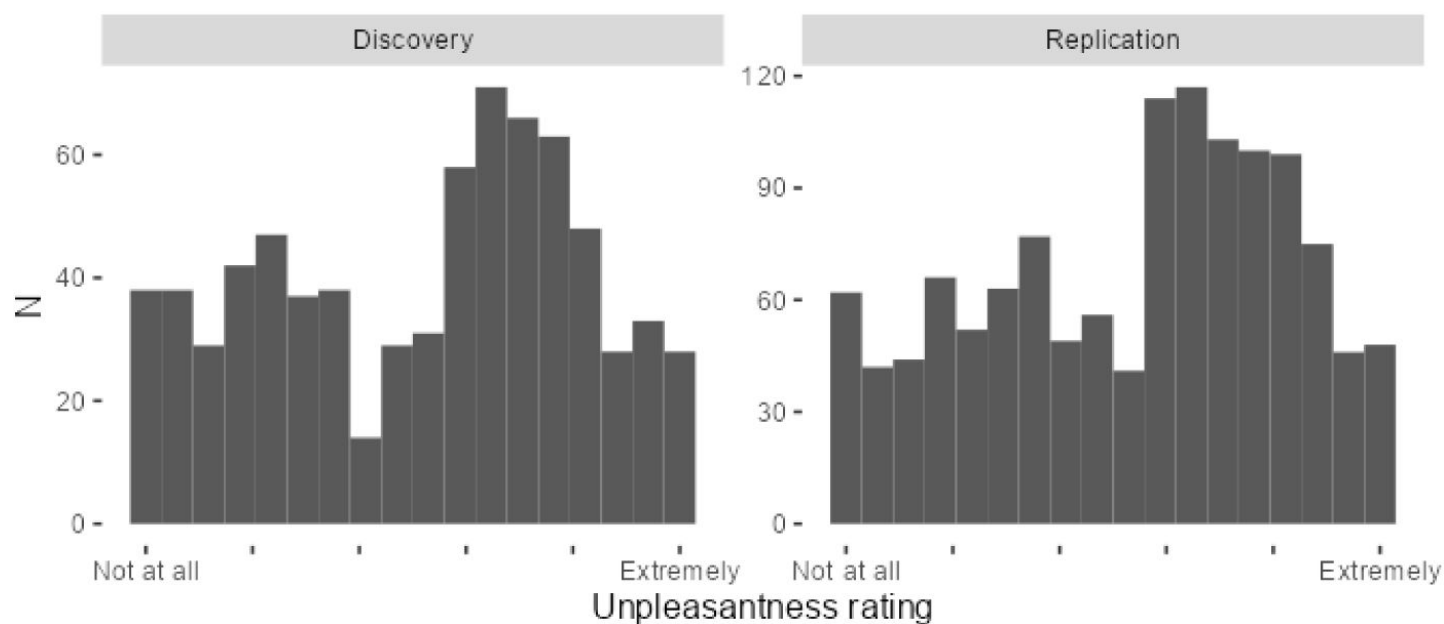
When unpleasantness ratings were included in the mediation models, the mediating effect of the reward-punishment sensitivity index did not survive (discovery sample: standardised $\beta = 0.003 \pm 0.003$, $p = 0.416$; replication sample: standardised $\beta = 0.004 \pm 0.003$, $p = 0.100$; **Supplementary Figure 12** [↗](#)). Pooling the samples resulted in an effect that narrowly missed the significance threshold (standardised $\beta = 0.004 \pm 0.002$, $p = 0.068$).



Supplementary Figure 9.

Sensitivity analysis of the computational findings relating to task-induced anxiety; comparing results when using parameters estimated via expectation maximisation (EM, red) and variational Bayesian inference (VBI, blue).

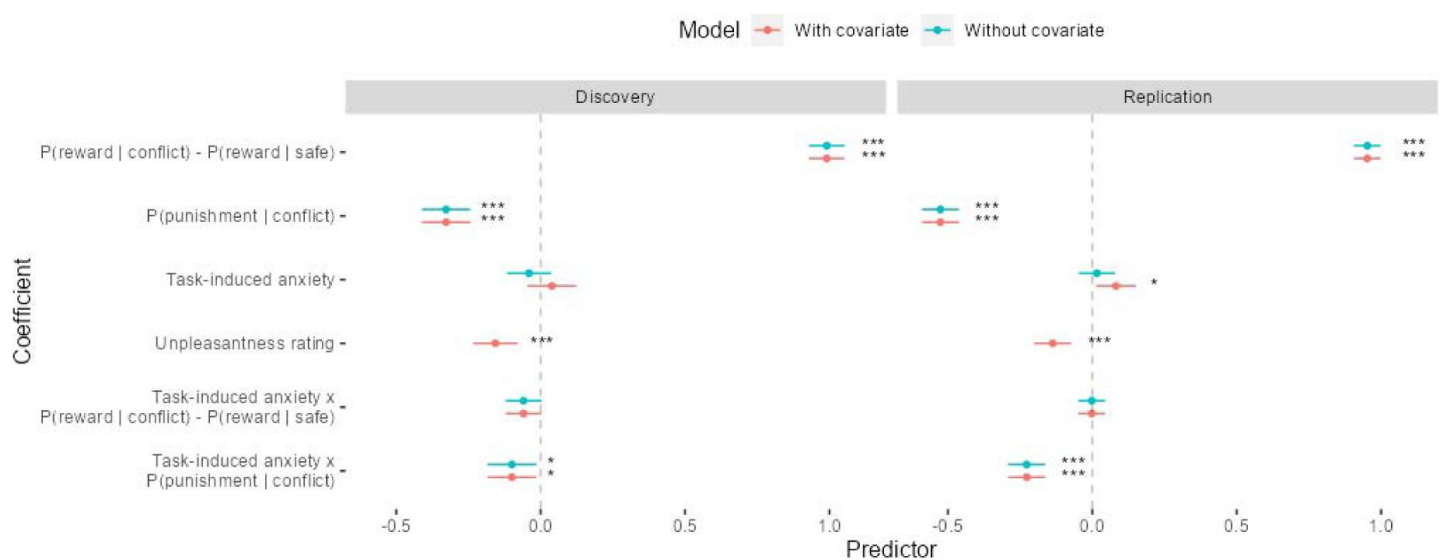
a, Kendall's tau correlations across each parameter and task-induced anxiety. **b**, Mediating effects of the punishment learning rate and reward-punishment sensitivity index.



Supplementary Figure 10.

Distribution of self-reported punishment unpleasantness ratings.

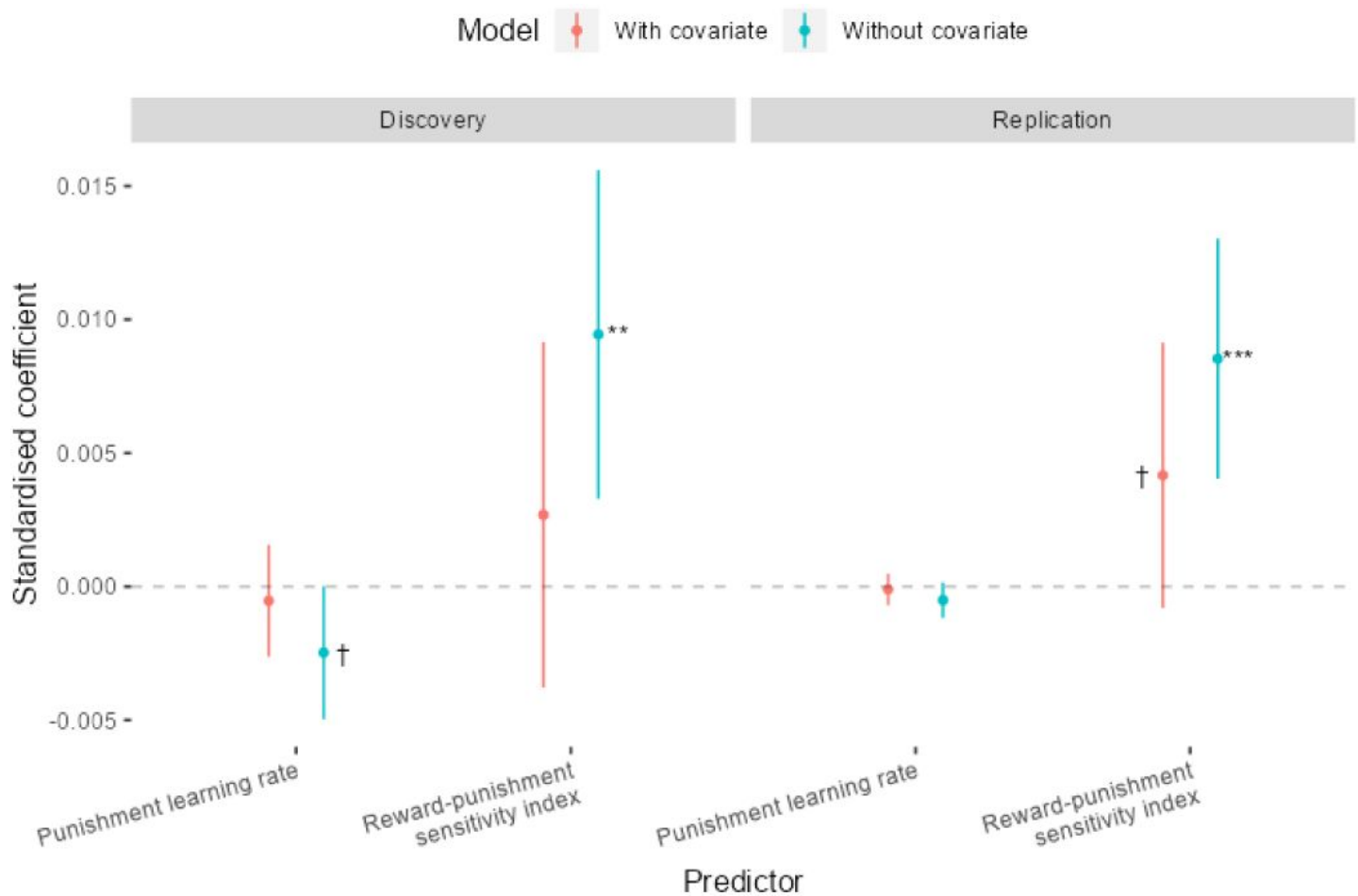
Ratings were scored from 'Not at all' to 'Extremely' (encoded as 0 and 50, respectively). Distributions are shown across the discovery and replication samples.



Supplementary Figure 11.

The effect of including unpleasantness ratings as a covariate in the hierarchical logistic regression models of task choices.

Dots represent coefficient estimates from the model, with confidence intervals. Models are shown for both the discovery and replication samples. Significance levels are shown according to the following: $p < 0.05$ - *; $p < 0.05$ - **; $p < 0.001$ - ***.



Supplementary Figure 12.

The effect of including unpleasantness ratings as a covariate in the mediation models.

Dots represent coefficient estimates from the model, with confidence intervals. Models are shown for both the discovery and replication samples. Significance levels are shown according to the following: $p < 0.1$ - †; $p < 0.05$ - *; $p < 0.01$ - **; $p < 0.001$ - ***.

Test-retest reliability of unpleasantness

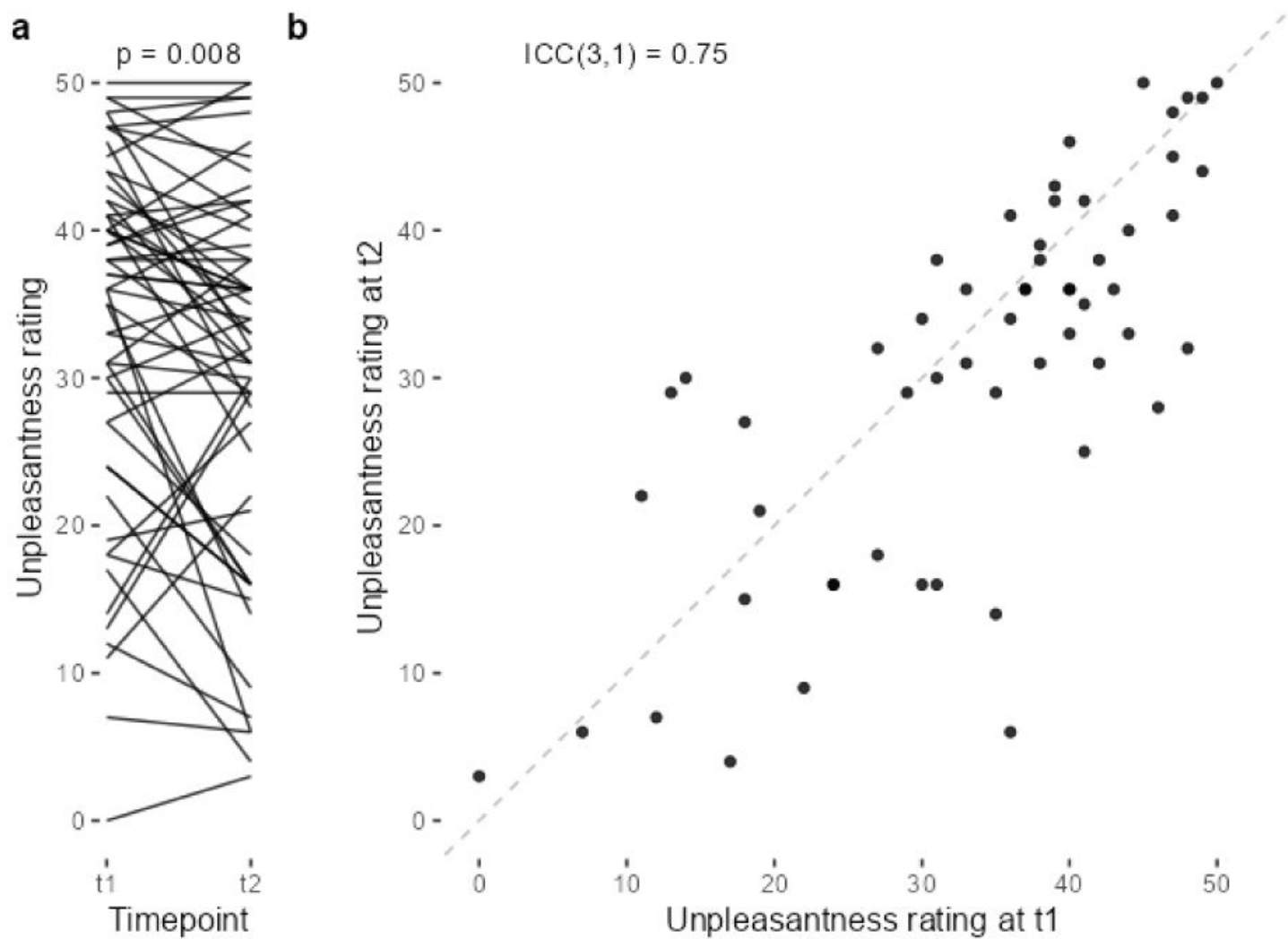
The test-retest reliability of unpleasantness ratings was excellent ($ICC(3,1) = 0.75$), although participants gave significantly lower ratings in the second session ($t_{56} = 2.7$, $p = 0.008$, $d = 0.37$; mean difference of 3.12, $SD = 8.63$).

Reliability of other measures with/out unpleasantness

To assess the effect of accounting for unpleasantness ratings on reliability estimates of task performance, we extracted variance components from linear mixed models, following a standard approach (Nakagawa, Johnson et al. 2017) – note that this was not the method used to estimate reliability values in the main analyses, but we used this specific approach to compare the reliability values with and without the covariate of unpleasantness ratings. The results indicated that unpleasantness ratings did not have a material effect on reliability (**Supplementary Figure 14** [↗](#)).

9.10 Sensitivity analysis: anxiety and inflexibility

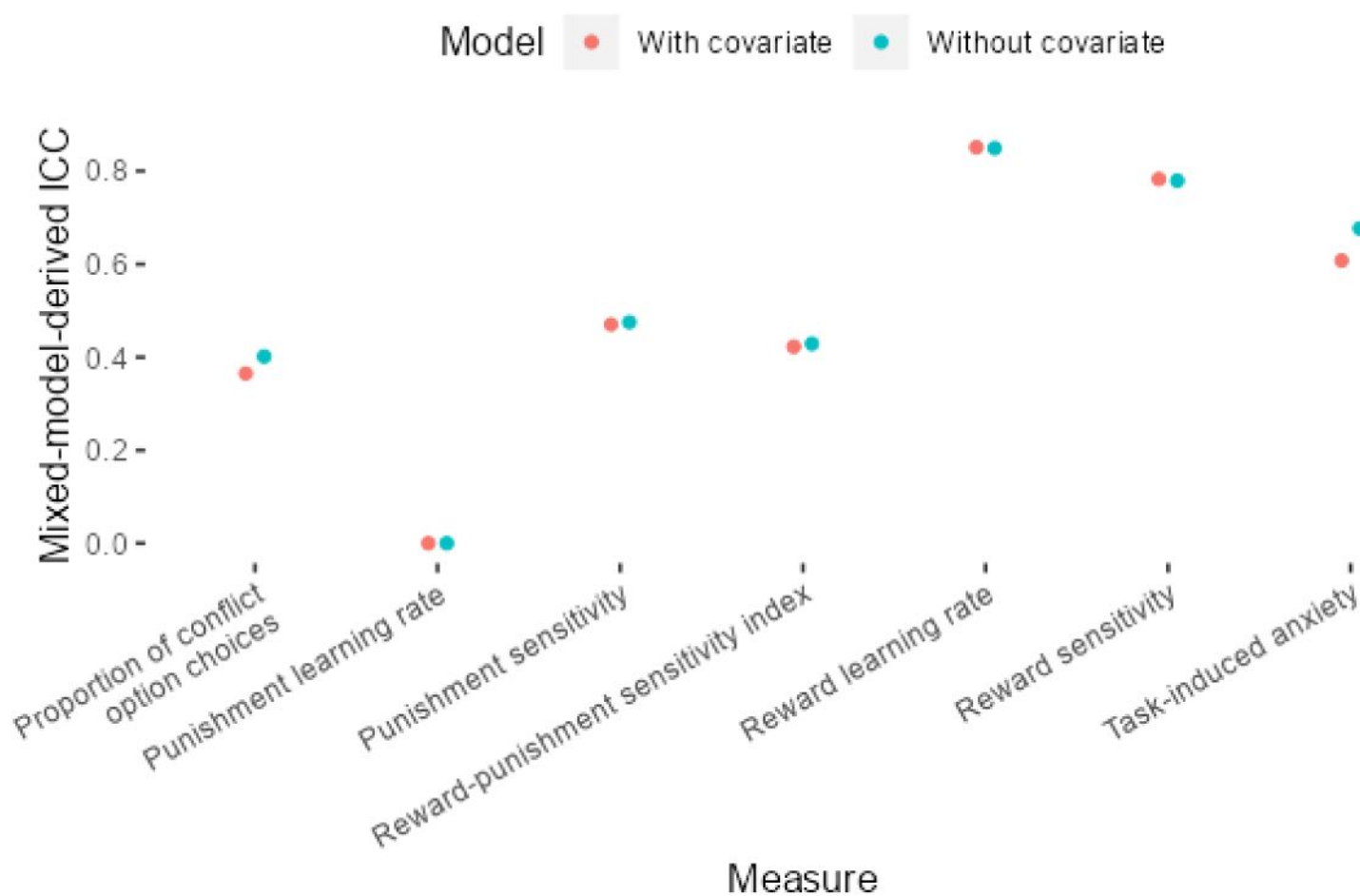
As anxious participants were slower to update their estimates of punishment probability, we determined whether this was due to greater general inflexibility by examining the model including two sensitivity parameters, but one general learning rate (i.e. not split by outcome). The correlation between this general learning rate and task-induced anxiety was not significant in either samples (discovery: $\tau = -0.02$, $p = 0.504$; replication: $\tau = -0.01$, $p = 0.625$), suggesting that the effect is specific to punishment.



Supplementary Figure 13.

Test-retest reliability of unpleasantness ratings.

a, Comparing unpleasantness ratings across timepoints, participants rated the punishments as significantly less unpleasant in the second session. **b**, Correlation of ratings across timepoints.



Supplementary Figure 14.

Mixed-model-derived intraclass correlation coefficients (ICCs) for measures of task performance, with and without accounting for unpleasantness.

Dots represent model-derived ICCs.

References

- Aldao A., Nolen-Hoeksema S., Schweizer S. (2010) **Emotion-regulation strategies across psychopathology: A meta-analytic review** *Clinical Psychology Review* **30**:217–237
- Aupperle R. L., Paulus M. P. (2010) **Neural systems underlying approach and avoidance in anxiety disorders** *Dialogues Clin Neurosci* **12**:517–531
- Aupperle R. L., Sullivan S., Melrose A. J., Paulus M. P., Stein M. B. (2011) **A reverse translational approach to quantify approach-avoidance conflict in humans** *Behavioural Brain Research* **225**:455–463
- Aylward J., Valton V., Ahn W. Y., Bond R. L., Dayan P., Roiser J. P., Robinson O. J. (2019) **Altered learning under uncertainty in unmedicated mood and anxiety disorders** *Nature Human Behaviour* **3**:1116–1123
- Bach D. R. (2021) **Cross-species anxiety tests in psychiatry: pitfalls and promises** *Molecular Psychiatry*
- Barlow D. H. (2002) **Anxiety and its disorders: The nature and treatment of anxiety and panic**
- Beck A. T., Clark D. A. (1997) **An information processing model of anxiety: Automatic and strategic processes** *Behaviour Research and Therapy* **35**:49–58
- Biedermann S. V., Biedermann D. G., Wenzlaff F., Kurjak T., Nouri S., Auer M. K., Wiedemann K., Briken P., Haaker J., Lonsdorf T. B., Fuss J. (2017) **An elevated plus-maze in mixed reality for studying human anxiety-related behavior** *Bmc Biology* **15**
- Bishop S. J. (2007) **Neurocognitive mechanisms of anxiety: an integrative account** *Trends in Cognitive Sciences* **11**:307–316
- Bishop S. J., Gagne C. (2018) **Anxiety, Depression, and Decision Making: A Computational Perspective** *Annual Review of Neuroscience* **41**:371–388
- Carlisi C. O., Robinson O. J. (2018) **The role of prefrontal-subcortical circuitry in negative bias in anxiety: Translational, developmental and treatment perspectives** *Brain Neurosci Adv* **2**
- Carver C. S., White T. L. (1994) **Behavioral-Inhibition, Behavioral Activation, and Affective Responses to Impending Reward and Punishment - the Bis Bas Scales** *Journal of Personality and Social Psychology* **67**:319–333
- Cicchetti D. V. (1994) **Guidelines, criteria, and rules of thumb for evaluating normed and standardized assessment instruments in psychology** *Psychological assessment* **6**
- Cisler J. M., Koster E. H. W. (2010) **Mechanisms of attentional biases towards threat in anxiety disorders: An integrative review** *Clinical Psychology Review* **30**:203–216
- Corr P. J. (2004) **Reinforcement sensitivity theory and personality** *Neuroscience and Biobehavioral Reviews* **28**:317–332

- Craske M. G., Hermans D. E., Vansteenwegen D. E. (2006) **Fear and learning: From basic processes to clinical implications** *American Psychological Association*
- Cryan J. F., Holmes A. (2005) **The ascent of mouse: Advances in modelling human depression and anxiety** *Nature Reviews Drug Discovery* **4**:775–790
- Daw N. D., O'Doherty J. P., Dayan P., Seymour B., Dolan R. J. (2006) **Cortical substrates for exploratory decisions in humans** *Nature* **441**:876–879
- Eckstein M. K., Master S. L., Xia L., Dahl R. E., Wilbrecht L., Collins A. G. E. (2022) **The interpretation of computational model parameters depends on the context** *Elife* **11**
- Faul F., Erdfelder E., Lang A.-G., Buchner A. (2007) **G* Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences** *Behavior research methods* **39**:175–191
- Fleiss J. L (1986) **The design and analysis of clinical experiments**
- Freedman R., Lewis D. A., Michels R., Pine D. S., Schultz S. K., Tamminga C. A., Gabbard G. O., Gau S. S. F., Javitt D. C., Oquendo M. A., Shrout P. E., Vieta E., Yager J. (2013) **The Initial Field Trials of DSM-5: New Blooms and Old Thorns** *American Journal of Psychiatry* **170**:1–5
- Gamez W., Chmielewski M., Kotov R., Ruggero C., Suzuki N., Watson D. (2014) **The Brief Experiential Avoidance Questionnaire: Development and Initial Validation** *Psychological Assessment* **26**:35–45
- Geller I., Kulak J. T., Seifter J. (1962) **The Effects of Chlordiazepoxide and Chlorpromazine on a Punishment Discrimination** *Psychopharmacologia* **3**:374–385
- Glover L. R., Postle A. F., Holmes A. (2020) **Touchscreen-based assessment of risky-choice in mice** *Behavioural Brain Research*
- Gray J. A (1982) **The Neuropsychology of Anxiety - an Inquiry into the Functions of the Septo- Hippocampal System** *Behavioral and Brain Sciences* **5**:469–484
- Griebel G., Holmes A. (2013) **50 years of hurdles and hope in anxiolytic drug discovery** *Nature Reviews Drug Discovery* **12**:667–687
- Guitart-Masip M., Huys Q. J., Fuentemilla L., Dayan P., Duzel E., Dolan R. J. (2012) **Go and no-go learning in reward and punishment: interactions between affect and effect** *Neuroimage* **62**:154–166
- Hall C. S (1934) **Emotional behavior in the rat I Defecation and urination as measures of individual differences in emotionality** *Journal of Comparative Psychology* **18**:385–403
- Hartley C. A., Phelps E. A. (2012) **Anxiety and Decision-Making** *Biological Psychiatry* **72**:113–118
- Hayes S. C (1976) **The role of approach contingencies in phobic behavior** *Behavior Therapy*.
- Hayes S. C., Wilson K. G., Gifford E. V., Follette V. M., Strosahl K. (1996) **Experiential avoidance and behavioral disorders: A functional dimensional approach to diagnosis and treatment** *Journal of Consulting and Clinical Psychology* **64**:1152–1168

- Hertwig R., Erev I. (2009) **The description-experience gap in risky choice** *Trends in Cognitive Sciences* **13**:517–523
- Huys Q. J. M., Cools R., Golzer M., Friedel E., Heinz A., Dolan R. J., Dayan P. (2011) **Disentangling the Roles of Approach, Activation and Valence in Instrumental and Pavlovian Responding** *Plos Computational Biology* **7**
- Ironside M., Amemori K., McGrath C. L., Pedersen M. L., Kang M. S., Amemori S., Frank M. J., Graybiel A. M., Pizzagalli D. A. (2020) **Approach-Avoidance Conflict in Major Depressive Disorder: Congruent Neural Findings in Humans and Nonhuman Primates** *Biological Psychiatry* **87**:399–408
- Jean-Richard-Dit-Bressel P., Lee J. C., Liew S. X., Weidemann G., Lovibond P. F., McNally G. P. (2021) **Punishment insensitivity in humans is due to failures in instrumental contingency learning** *Elife* **10**
- Kakoschke N., Albertella L., Lee R. S. C., Wiers R. W. (2019) **Assessment of Automatically Activated Approach-Avoidance Biases Across Appetitive Substances** *Current Addiction Reports* **6**:200–209
- Kirlic N., Young J., Aupperle R. L. (2017) **Animal to human translational paradigms relevant for approach avoidance conflict decision making** *Behaviour Research and Therapy* **96**:14–29
- Kriegeskorte N., Douglas P. K. (2018) **Cognitive computational neuroscience** *Nature Neuroscience* **21**:1148–1160
- Kroenke K., Strine T. W., Spitzer R. L., Williams J. B. W., Berry J. T., Mokdad A. H. (2009) **The PHQ-8 as a measure of current depression in the general population** *Journal of Affective Disorders* **114**:163–173
- Letkiewicz A. M., Kottler H. C., Shankman S. A., Cochran A. L. (2023) **Quantifying aberrant approach-avoidance conflict in psychopathology: A review of computational approaches** *Neuroscience and Biobehavioral Reviews* **147**
- Mathews A., Mackintosh B. (1998) **A cognitive model of selective processing in anxiety** *Cognitive Therapy and Research* **22**:539–560
- Metha J. A., Brian M. L., Oberrauch S., Barnes S. A., Featherby T. J., Bossaerts P., Murawski C., Hoyer D., Jacobson L. H. (2020) **Separating Probability and Reversal Learning in a Novel Probabilistic Reversal Learning Task for Mice** *Frontiers in Behavioral Neuroscience* **13**
- Milad M. R., Quirk G. J. (2012) **Fear Extinction as a Model for Translational Neuroscience: Ten Years of Progress** *Annual Review of Psychology* **63**:129–151
- Mkrtchian A., Aylward J., Dayan P., Roiser J. P., Robinson O. J. (2017) **Modeling Avoidance in Mood and Anxiety Disorders Using Reinforcement Learning** *Biological Psychiatry* **82**:532–539
- Mobbs D., Headley D. B., Ding W. L., Dayan P. (2020) **Space, Time, and Fear: Survival Computations along Defensive Circuits** *Trends in Cognitive Sciences* **24**:228–241
- Nakagawa S., Johnson P. C. D., Schielzeth H. (2017) **The coefficient of determination R² and intra-class correlation coefficient from generalized linear mixed-effects models revisited and expanded** *Journal of the Royal Society Interface* **14**

- Oberrauch S., Sigrist H., Sautter E., Gerster S., Bach D. R., Pryce C. R. (2019) **Establishing operant conflict tests for the translational study of anxiety in mice** *Psychopharmacology* **236**:2527–2541
- Pedersen M. L., Ironside M., Amemori K., McGrath C. L., Kang M. S., Graybiel A. M., Pizzagalli D. A., Frank M. J. (2021) **Computational phenotyping of brain-behavior dynamics underlying approach-avoidance conflict in major depressive disorder** *Plos Computational Biology* **17**
- Pellow S., Chopin P., File S. E., Briley M. (1985) **Validation of Open -Closed Arm Entries in an Elevated Plus-Maze as a Measure of Anxiety in the Rat** *Journal of Neuroscience Methods* **14**:149–167
- Phaf R. H., Mohr S. E., Rotteveel M., Wicherts J. M. (2014) **Approach, avoidance, and affect: a meta-analysis of approach-avoidance tendencies in manual reaction time tasks** *Frontiers in Psychology* **5**
- Pike A. C., Lowther M., Robinson O. J. (2021) **The Importance of Common Currency Tasks in Translational Psychiatry** *Current Behavioral Neuroscience Reports* **8**:1–10
- Pike A. C., Robinson O. J. (2022) **Reinforcement Learning in Patients With Mood and Anxiety Disorders vs Control Individuals A Systematic Review and Meta-analysis** *Jama Psychiatry*
- Pizzagalli D. A., Jahn A. L., O'Shea J. P. (2005) **Toward an objective characterization of an anhedonic phenotype: A signal detection approach** *Biological Psychiatry* **57**:319–327
- Redish A. D., Kepecs A., Anderson L. M., Calvin O. L., Grissom N. M., Haynos A. F., Heilbronner S. R., Herman A. B., Jacob S., Ma S. S., Vilares I., Vinogradov S., Walters C. J., Widge A. S., Zick J. L., Zilverstand A. (2022) **Computational validity: using computation to translate behaviours across species** *Philosophical Transactions of the Royal Society B-Biological Sciences* **377**
- Robinson O. J., Vytal K., Cornwell B. R., Grillon C. (2013) **The impact of anxiety upon cognition: perspectives from human threat of shock studies** *Frontiers in Human Neuroscience* **7**
- Rosseel Y (2012) **lavaan: An R package for structural equation modeling** *Journal of statistical software* **48**:1–36
- Rutledge R. B., Chekroud A. M., Huys Q. J. M. (2019) **Machine learning and big data in psychiatry: toward clinical applications** *Current Opinion in Neurobiology* **55**:152–159
- Schultz W (2013) **Updating dopamine reward signals** *Current Opinion in Neurobiology* **23**:229–238
- Seow T. X. F., Hauser T. U. (2022) **Reliability of web-based affective auditory stimulus presentation** *Behavior Research Methods* **54**:378–392
- Seymour B., Daw N. D., Roiser J. P., Dayan P., Dolan R. (2012) **Serotonin Selectively Modulates Reward Value in Human Decision-Making** *Journal of Neuroscience* **32**:5833–5842
- Sharp P. B., Russek E. M., Huys Q. J. M., Dolan R. J., Eldar E. (2022) **Humans persevere on punishment avoidance goals in multigoal reinforcement learning** *Elife* **11**
- Shrout P. E., Fleiss J. L. (1979) **Intraclass Correlations - Uses in Assessing Rater Reliability** *Psychological Bulletin* **86**:420–428

- Sierra-Mercado D., Deckersbach T., Arulpragasam A. R., Chou T. N., Rodman A. M., Duffy A., McDonald E. J., Eckhardt C. A., Corse A. K., Kaur N., Eskandar E. N., Dougherty D. D. (2015) **Decision making in avoidance-reward conflict: a paradigm for non-human primates and humans** *Brain Structure & Function* **220**:2509–2517
- Simon N. W., Gilbert R. J., Mayse J. D., Bizon J. L., Setlow B. (2009) **Balancing Risk and Reward: A Rat Model of Risky Decision Making** *Neuropsychopharmacology* **34**:2208–2217
- Spitzer R. L., Kroenke K., Williams J. B. W., Lowe B. (2006) **A brief measure for assessing generalized anxiety disorder - The GAD-7** *Archives of Internal Medicine* **166**:1092–1097
- Stan Development Team (2023) **RStan: the R interface to Stan**
- Stein M. B., Paulus M. P. (2009) **Imbalance of Approach and Avoidance: The Yin and Yang of Anxiety Disorders** *Biological Psychiatry* **66**:1072–1074
- Struijs S. Y., Lamers F., Rinck M., Roelofs K., Spinhoven P., Penninx B. W. J. H. (2018) **The predictive value of Approach and Avoidance tendencies on the onset and course of depression and anxiety disorders** *Depression and Anxiety* **35**:551–559
- Sutton R. S., Barto A. G. (2018) **Reinforcement Learning: An Introduction, 2nd Edition** *Reinforcement Learning: An Introduction* :1–526
- Talmi D., Dayan P., Kiebel S. J., Frith C. D., Dolan R. J. (2009) **How Humans Integrate the Prospects of Pain and Reward during Choice** *Journal of Neuroscience* **29**:14617–14626
- Treit D., Engin E., McEown K. (2010) **Animal Models of Anxiety and Anxiolytic Drug Action** *Behavioral Neurobiology of Anxiety and Its Treatment* **2**:121–160
- van den Bos R., Koot S., de Visser L. (2014) **A rodent version of the Iowa Gambling Task: 7 years of progress** *Frontiers in Psychology* **5**
- Vogel J. R., Beer B., Clody D. E. (1971) **A simple and reliable conflict procedure for testing anti-anxiety agents** *Psychopharmacologia* **21**:1–7
- Waltmann M., Schlagenhauf F., Deserno L. (2022) **Sufficient reliability of the behavioral and computational readouts of a probabilistic reversal learning task** *Behavior Research Methods*
- Walz N., Muhlberger A., Pauli P. (2016) **A Human Open Field Test Reveals Thigmotaxis Related to Agoraphobic Fear** *Biological Psychiatry* **80**:390–397
- Watkins C. J. C. H., Dayan P. (1992) **Q-Learning** *Machine Learning* **8**:279–292
- Willner P (1984) **The Validity of Animal-Models of Depression** *Psychopharmacology* **83**:1–16
- Wise T., Dolan R. J. (2020) **Associations between aversive learning processes and transdiagnostic psychiatric symptoms in a general population sample** *Nature Communications* **11**
- Wise T., Robinson O. J., Gillan C. M. (2022) **Identifying Transdiagnostic Mechanisms in Mental Health Using Computational Factor Modeling** *Biol Psychiatry*
- Woods K. J. P., Siegel M. H., Traer J., McDermott J. H. (2017) **Headphone screening to facilitate web-based auditory experiments** *Attention Perception & Psychophysics* **79**:2064–2072

Author information

Yumeya Yamamori

Institute of Cognitive Neuroscience, University College London, London, UK

ORCID iD: [0000-0001-5508-7965](https://orcid.org/0000-0001-5508-7965)

Oliver J Robinson

Institute of Cognitive Neuroscience, University College London, London, UK, Research

Department of Clinical, Educational and Health Psychology, University College London, London, UK

ORCID iD: [0000-0002-3100-1132](https://orcid.org/0000-0002-3100-1132)

Jonathan P Roiser

Institute of Cognitive Neuroscience, University College London, London, UK

For correspondence: yumeya.yamamori@ucl.ac.uk

ORCID iD: [0000-0001-8269-1228](https://orcid.org/0000-0001-8269-1228)

Editors

Reviewing Editor

Claire Gillan

Trinity College Dublin, Ireland

Senior Editor

Michael Frank

Brown University, United States of America

Reviewer #1 (Public Review):

This paper describes the development and initial validation of an approach-avoidance task and its relationship to anxiety. The task is a two-armed bandit where one choice is 'safer' - has no probability of punishment, delivered as an aversive sound, but also lower probability of reward - and the other choice involves a reward-punishment conflict. The authors fit a computational model of reinforcement learning to this task and found that self-reported state anxiety during the task was related to a greater likelihood of choosing the safe stimulus when the other (conflict) stimulus had a higher likelihood of punishment. Computationally, this was represented by a smaller value for the ratio of reward to punishment sensitivity in people with higher task-induced anxiety. They replicated this finding, but not another finding that this behavior was related to a measure of psychopathology (experiential avoidance), in a second sample. They also tested test-retest reliability in a sub-sample tested twice, one week apart and found that some aspects of task behavior had acceptable levels of reliability. The introduction makes a strong appeal to back-translation and computational validity. The task design is clever and most methods are solid - it is encouraging to see attempts to validate tasks as they are developed. The lack of replicated effects with psychopathology may mean that this task is better suited to assess state anxiety, or to serve as a foundation for additional task development.

Reviewer #2 (Public Review):

Summary:

The authors develop a computational approach-avoidance-conflict (AAC) task, designed to overcome the limitations of existing offer based AAC tasks. The task incorporated likelihoods of receiving rewards/ punishments that would be learned by the participants to ensure computational validity and estimated model parameters related to reward/punishment and task induced anxiety. Two independent samples of online participants were tested. In both samples participants who experienced greater task induced anxiety avoided choices associated with greater probability of punishment. Computational modelling revealed that this effect was explained by greater individual sensitivities to punishment relative to rewards.

Strengths:

Large internet-based samples, with discovery sample ($n = 369$), pre-registered replication sample ($n = 629$) and test-retest sub group ($n = 57$). Extensive compliance measures (e.g. audio checks) seek to improve adherence.

There is a great need for RL tasks that model threatening outcomes rather than simply loss of reward. The main model parameters show strong effects and the additional indices with task based anxiety are a useful extension. Associations were broadly replicated across samples. Fair to excellent reliability of model parameters is encouraging and badly needed for behavioral tasks of threat sensitivity.

The task seems to have lower approach bias than some other AAC tasks in the literature.

Weaknesses:

The negative reliability of punishment learning rate is concerning as this is an important outcome.

The Kendall's tau values underlying task induced anxiety and safety reference/ various indices are very weak (all < 0.1), as are the mediation effects (all $\beta < 0.01$). The interaction with $P(\text{punishment} | \text{conflict})$ does explain some of this.

The inclusion of only one level of reward (and punishment) limits the ecological validity of the sensitivity indices.

Appraisal and impact:

Overall this is a very strong paper, describing a novel task that could help move the field of RL forward to take account of threat processing more fully. The large sample size with discovery, replication and test-retest gives confidence in the findings. The task has good ecological validity and associations with task-based anxiety and clinical self-report demonstrate clinical relevance. Test-retest of the punishment learning parameter is the only real concern. Overall this task provides an exciting new probe of reward/threat that could be used in mechanistic disease models.

Additional context:

The sex differences between the samples are interesting as effects of sex are commonly found in AAC tasks. It would be interesting to look at the main model comparison with sex included as a covariate.

Reviewer #3 (Public Review):

This study investigated cognitive mechanisms underlying approach-avoidance behavior using a novel reinforcement learning task and computational modelling. Participants could select a risky "conflict" option (latent, fluctuating probabilities of monetary reward and/or unpleasant sound [punishment]) or a safe option (separate, generally lower probability of reward). Overall, participant choices were skewed towards more rewarded options, but were also repelled by increasing probability of punishment. Individual patterns of behavior were well-captured by a reinforcement learning model that included parameters for reward and punishment sensitivity, and learning rates for reward and punishment. This is a nice replication of existing findings suggesting reward and punishment have opposing effects on behavior through dissociated sensitivity to reward versus punishment.

Interestingly, avoidance of the conflict option was predicted by self-reported task-induced anxiety. Importantly, when a subset of participants were retested over 1 week later, most behavioral tendencies and model parameters were recapitulated, suggesting the task may capture stable traits relevant to approach-avoidance decision-making.

The revised paper commendably adds important additional information and analyses to support these claims. The initial concern that not accounting for participant control over punisher intensity confounded interpretation of effects has been largely addressed in follow-up analyses and discussion.

This study complements and sits within a broad translational literature investigating interactions between reward/punishers and psychological processes in approach-avoidance decisions.

Author Response

The following is the authors' response to the original reviews.

Reviewer #1 (Public Review):

This paper describes the development and initial validation of an approach-avoidance task and its relationship to anxiety. The task is a two-armed bandit where one choice is 'safer' - has no probability of punishment, delivered as an aversive sound, but also lower probability of reward - and the other choice involves a reward-punishment conflict. The authors fit a computational model of reinforcement learning to this task and found that self-reported state anxiety during the task was related to a greater likelihood of choosing the safe stimulus when the other (conflict) stimulus had a higher likelihood of punishment. Computationally, this was represented by a smaller value for the ratio of reward to punishment sensitivity in people with higher task-induced anxiety. They replicated this finding, but not another finding that this behavior was related to a measure of psychopathology (experiential avoidance), in a second sample. They also tested test-retest reliability in a sub-sample tested twice, one week apart and found that some aspects of task behavior had acceptable levels of reliability. The introduction makes a strong appeal to back-translation and computational validity, but many aspects of the rationale for this task need to be strengthened or better explained. The task design is clever and most methods are solid - it is encouraging to see attempts to validate tasks as they are developed. There are a few methodological questions and interpretation issues, but they do not affect the overall findings. The lack of replicated effects with psychopathology may mean that this task is better suited to assess state anxiety, or to serve as a foundation for additional task development.

We thank the reviewer for their kind comments and constructive feedback. We agree that the approach taken in this paper appears better suited to state anxiety, and further work is needed to assess/improve its clinical relevance.

Reviewer #1 (Recommendations For The Authors):

1. For the introduction, the authors communicate well the appeal of tasks with translational potential, and setting up this translation through computational validity is a strong approach. However, I had some concerns about how the task was motivated in the introduction:

a) The authors state that current approach-avoidance tasks used in humans do not resemble those used in the non-human literature, but do not provide details on what exactly is missing from these tasks that makes translation difficult.

Our intention for the section that the reviewer refers to was to briefly convey that historically, approach-avoidance conflict would have been measured either using questionnaires or joystick-based tasks which have no direct non-human counterpart. However, we note that the phrasing was perhaps unfair to recent tasks that were explicitly designed to be translatable across species. Therefore, we have amended the text to the following:

In humans, on the other hand, approach-avoidance conflict has historically been measured using questionnaires such as the Behavioural Inhibition/Activation Scale (Carver & White, 1994), or cognitive tasks that rely on motor biases, for example by using joysticks to approach/move towards positive stimuli and avoid/move away from negative stimuli, which have no direct non-human counterparts (Guitart-Masip et al., 2012; Kirlic et al., 2017; Mkrtchian et al., 2017; Phaf et al., 2014).

b) Although back-translation to 'match' human paradigms to non-animal paradigms is useful for research, this isn't the end goal of task development. What really matters is how well these tasks, whether in humans or not, capture psychopathology-relevant behavior. Many animal paradigms were developed and brought into extensive use because they showed sensitivity to pharmacological compounds (e.g., benzodiazepines). The introduction accepts the validity of these paradigms at face value, and doesn't address whether developing human tests of psychopathology based on sensitivity to existing medication classes is the best way to generate new insights about psychopathology.

We agree that whilst paradigms with translational and computational validity have merits of their own for neuroscientific theory, clinical validity (i.e. how well the paradigm reflects a phenomenon relevant to psychopathology) is key in the context of clinical applications. While our findings of associations between task performance and self-reported (state) anxiety suggest that our approach is a step in the right direction, the lack of associations with clinical measures was disappointing. Although future work is needed to more directly test the sensitivity of the current approach to psychopathology, this may mean that it, and its non-human counterparts, do not measure behaviours relevant to pathological anxiety. Since our primary focus in this paper was on translational and computational validity, we have opted to discuss the author's suggestion in the 'Discussion' section, as follows:

Further, it is worth noting that many animal paradigms were developed and widely adopted due to their sensitivity to anxiolytic medication (Cryan & Holmes, 2005). Given the lack of associations with clinical measures in our results, it is possible that current translational models of anxiety may not fully capture behaviours that are directly relevant to pathological

anxiety. To develop translational paradigms of clinical utility, future research should place a stronger emphasis on assessing their clinical validity in humans.

c) The authors may want to bring in the literature on the description-experience gap (e.g., PMID: 19836292) when discussing existing decision tasks and their computational dissimilarity to non-human operant conditioning tasks.

We thank the reviewer for this useful addition to the introduction. We have now added the following to the 'Introduction' section:

Moreover, evidence from economic decision-making suggests that explicit offers of probabilistic outcomes can impact decision-making differently compared to when probabilistic contingencies need to be learned from experience (referred to as the 'description-experience gap'; Hertwig & Erev, 2009); this finding raises potential concerns regarding the use of offer-based tasks in humans as approximations of non-human tasks that do not involve explicit offers.

d) How does one evaluate how computationally similar human vs. non-human tasks are? What are the criteria for making this judgement? Specific to the current tasks, many animal learning tasks are not learning tasks in the same sense that human learning tasks are, in terms of the number of trials used and if the animals are choosing from a learned set of contingencies versus learning the contingencies during the testing.

The computational similarity of human and non-human strategies in a given translational task can be tested empirically. This can be done by fitting models to the data and assessing whether similar models explain choices, even if parameter distributions might vary across species due to, for example, physiological differences. Indeed, non-human animals require much more training to perform even uni-dimensional reinforcement learning, but once they are trained, it should be possible to model their responses. In fact, it should even be possible to take training data into account in some cases. For example, the training phase of the Vogel/Geller-Seifter preclinical tests require an animal to learn to emit a certain action (e.g. lever press) simply to obtain some reward. In the next phase, an aversive outcome is introduced as an additional outcome, but one could model both the training and test phase together – the winning model in our studies would be a suitable candidate to model behaviour here. As we also discuss predictive validity in the 'Discussion' section, we opted to add the following text there too:

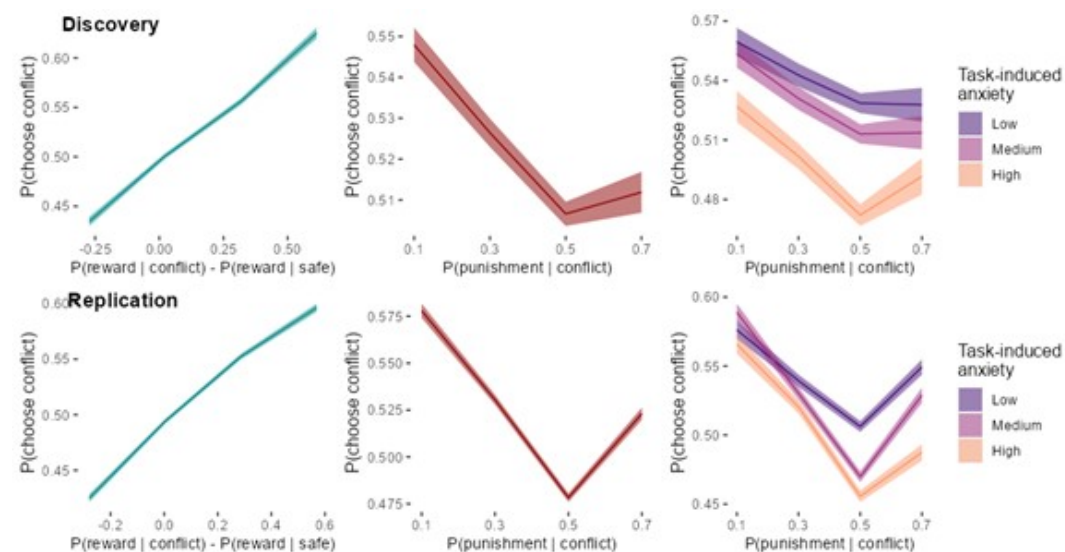
... computational validity would also need to be assessed directly in non-human animals by fitting models to their behavioural data. This should be possible even in the face of different procedures across species such as number of trials or outcomes used (shock or aversive sound). We are encouraged by our finding that the winning computational model in our study relies on a relatively simple classical reinforcement learning strategy. There exist many studies showing that non-human animals rely on similar strategies during reward and punishment learning (Mobbs et al., 2020; Schultz, 2013); albeit to our knowledge this has never been modelled in non-human animals where rewards and punishment can occur simultaneously.

1. What do the authors make of the non-linear relationship between probability of punishment and probability of choosing the conflict stimulus (Fig 2d), especially in the high task-induced anxiety participants? Did this effect show up in the replication sample as well?

Figures 2c-e were created by binning the continuous predictors of outcome probabilities into discrete bins of equal interval. Since punishment probability varied according to Gaussian random walks, it was also distributed with more of its mass in the central region (~ 0.4), and

so values at the extreme bins were estimated on fewer data and with greater variance. The non-linear relationships are likely thus an artefact of our task design and plotting procedure. The pattern was also evident in the replication sample, see Author response image 1:

Author response image 1.



However, since these effects were estimated as linear effects in the logistic regression models, and to avoid overfitting/interpretations of noise arising from our task design, we now plot logistic curves fitted to the raw data instead.

1. How correlated were learning rate and sensitivity parameters? The EM algorithm used here can sometimes result in high correlations among these sets of parameters.

As the reviewer suspects the parameters were strongly correlated, especially across the punishment-specific parameters. The Pearson's r estimates for the untransformed parameter values were as follows:

Reward parameters: discovery sample $r = -0.39$; replication sample $r = -0.78$

Punishment parameters: discovery sample $r = -0.91$; replication sample $r = -0.85$

We have included the correlation matrices of the estimated parameters as Supplementary Figure 2 in the 'Computational modelling' section of the Supplement.

We have now also re-fitted the winning model using variational Bayesian inference (VBI) via Stan, and found that the cross-parameter correlations were much lower than when the data were fitted using EM. We also ran a sensitivity analysis assessing whether using VBI changed the main findings of our studies. This showed that the correlation between task-induced anxiety and the reward-punishment sensitivity index was robust to fitting method, as was the mediating effect of reward-punishment sensitivity index on anxiety's effect on choice. This indicates that overall our key findings are robust to different methods of parameter-fitting.

We now direct readers to these analyses from the new 'Sensitivity analyses' section in the manuscript, as follows:

As our procedure for estimating model parameters (the expectation-maximisation algorithm, see 'Methods') produced high inter-parameter correlations in our data (Supplementary Figure 2), we also re-estimated the parameters using Stan's variational Bayesian inference

algorithm (Stan Development Team, 2023) – this resulted in lower inter-parameter correlations, but our primary computational finding, that the effect of anxiety on choice is mediated by relative sensitivity to reward/punishment was consistent across algorithms (see Supplement section 9.8 for details).

We have included the relevant analyses comparing EM and VBI in the Supplement, as follows:

[9.8 Sensitivity analysis: estimating parameters via expectation maximisation and variational Bayesian inference algorithms]

Given that the expectation maximisation (EM) algorithm produced high inter-parameter correlations, we ran a sensitivity analysis by assessing the robustness of our computational findings to an alternative method of parameter estimation – (mean-field) variational Bayesian inference (VBI) via Stan (Stan Development Team, 2023). Since, unlike EM, the results of VBI are very sensitive to initial values, we fitted the data 10 times with different initial values.

Inter-parameter correlations

The VBI produced lower inter-parameter correlations than the EM algorithm (Supplementary Figure 8).

Sensitivity analysis

Since multicollinearity in the VBI-estimated parameters was lower than for EM, indicating less trade-off in the estimation, we re-tested our computational findings from the manuscript as part of a sensitivity analysis. We first assessed whether we observed the same correlations between task-induced anxiety and punishment learning, and reward-punishment sensitivity index (Supplementary Figure 9a). Punishment learning rate was not significantly associated with task-induced anxiety in any of the 10 VBI iterations in the discovery sample, although it was in 9/10 in the replication sample. On the other hand, the reward-punishment sensitivity index was significantly associated with task-induced anxiety in 9/10 VBI iterations in the discovery sample and all iterations in the replication sample. This suggests that the correlation of anxiety and sensitivity index is robust to these two fitting approaches.

We also re-estimated the mediation models, where in the EM-estimated parameters, we found that the reward-punishment sensitivity index mediated the relationship between task-induced anxiety and task choice proportions (Supplementary Figure 9b). Again, we found that the reward-punishment sensitivity index was a significant mediator in 9/10 VBI iterations in the discovery sample and all iterations in the replication sample. Punishment learning rate was also a significant mediator in 9/10 iterations in the replication sample, although it was not in the discovery sample for all iterations, and this was not observed for the EM-estimated parameters.

Overall, we found that our key results, that anxiety is associated with greater sensitivity to punishment over reward, and this mediates the relationship between anxiety and approach-avoidance behaviour, were robust across both fitting methods.

As an aside, we were unable to run the model fitting using Markov chain Monte Carlo sampling approaches due to the computational power and time required for a sample of this size (Pike & Robinson, 2022, JAMA Psychiatry).

| 1. What is the split-half reliability of the task parameters?

We thank the reviewer for this query. We have now included a brief section on the (good-to-excellent) split-half reliability of the task in the manuscript:

We assessed the split-half reliability of the task by correlating the overall proportion of conflict option choices and model parameters from the winning model across the first and second half of trials. For overall choice proportion, reliability was simply calculated via Pearson's correlations. For the model parameters, we calculated model-derived estimates of Pearson's r values from the parameter covariance matrix when first- and second-half parameters were estimated within a single model, following a previous approach recently shown to accurately estimate parameter reliability (Waltmann et al., 2022). We interpreted indices of reliability based on conventional values of < 0.40 as poor, $0.4 - 0.6$ as fair, $0.6 - 0.75$ as good, and > 0.75 as excellent reliability (Fleiss, 1986). Overall choice proportion showed good reliability (discovery sample $r = 0.63$; replication sample $r = 0.63$; Supplementary Figure 5). The model parameters showed good-to-excellent reliability (model-derived r values ranging from 0.61 to 0.85 [0.76 to 0.92 after Spearman-Brown correction]; Supplementary Figure 5).

1. The authors do a good job of avoiding causal language when setting up the cross-sectional mediation analysis, but depart from this in the discussion (line 335). Without longitudinal data, they cannot claim that "mediation analyses revealed a mechanism of how anxiety induces avoidance".

Thank you for spotting this, we have now amended the text to:

... mediation analyses suggested a potential mechanism of how anxiety may induce avoidance.

Reviewer #2 (Public Review):

Summary:

The authors develop a computational approach-avoidance-conflict (AAC) task, designed to overcome limitations of existing offer based AAC tasks. The task incorporated likelihoods of receiving rewards/ punishments that would be learned by the participants to ensure computational validity and estimated model parameters related to reward/punishment and task induced anxiety. Two independent samples of online participants were tested. In both samples participants who experienced greater task induced anxiety avoided choices associated with greater probability of punishment. Computational modelling revealed that this effect was explained by greater individual sensitivities to punishment relative to rewards.

Strengths:

Large internet-based samples, with discovery sample ($n = 369$), pre-registered replication sample ($n = 629$) and test-retest sub group ($n = 57$). Extensive compliance measures (e.g. audio checks) seek to improve adherence.

There is a great need for RL tasks that model threatening outcomes rather than simply loss of reward. The main model parameters show strong effects and the additional indices with task based anxiety are a useful extension. Associations were broadly replicated across samples. Fair to excellent reliability of model parameters is encouraging and badly needed for behavioral tasks of threat sensitivity.

We thank the reviewer for their comments and constructive feedback.

The task seems to have lower approach bias than some other AAC tasks in the literature. Although this was inferred by looking at Fig 2 (it doesn't seem to drop below 46%) and Fig 3d seems to show quite a strong approach bias when using a reward/punishment

sensitivity index. It would be good to confirm some overall stats on % of trials approached/avoided overall.

The range of choice proportions is indeed an interesting statistic that we have now included in the manuscript:

Across individuals, there was considerable variability in overall choice proportions (discovery sample: mean = 0.52, SD = 0.14, min/max = [0.03, 0.96]; replication sample: mean = 0.52, SD = 0.14, min/max = [0.01, 0.99]).

Weaknesses:

The negative reliability of punishment learning rate is concerning as this is an important outcome.

We agree that this is a concerning finding. As reviewer 3 notes, this may have been due to participants having control over the volume used to play the aversive sounds in the task (see below for our response to this point). Future work with better controlled experimental settings will be needed to determine the reliability of this parameter more accurately.

This may also have been due to the asymmetric nature of the task, as only one option could produce the punishment. This means that there were fewer trials on which to estimate learning about the occurrence of a punishment. Future work using continuous outcomes, as the reviewer suggests below, whilst keeping the asymmetric relationship between the options, could help in this regard.

We have included the following comment on this issue in the manuscript:

Alternatively, as participants self-determined the loudness of the punishments, differences in volume settings across sessions may have impacted the reliability of this parameter (and indeed punishment sensitivity). Further, the asymmetric nature of the task may have impacted our ability to estimate the punishment learning rate, as there were fewer occurrences of the punishment compared to the reward.

The Kendall's tau values underlying task induced anxiety and safety reference/ various indices are very weak (all < 0.1), as are the mediation effects (all beta < 0.01). This should be highlighted as a limitation, although the interaction with P(punishment|conflict) does explain some of this.

We now include references to the effect sizes to emphasise this limitation. We also note, as the reviewer suggests, that this may be due to crudeness of overall choice proportion as a measure of approach/avoidance, as it is contaminated with variables such as P(punishment|conflict).

One potentially important limitation of our findings is the small effect size observed in the correlation between task-induced anxiety and avoidance (Kendall's tau values < 0.1, mediation betas < 0.01). This may be attributed to the simplicity of using overall choice proportion as a measure of approach/avoidance, as the effect of anxiety on choice was also influenced by punishment probability.

The inclusion of only one level of reward (and punishment) limits the ecological validity of the sensitivity indices.

We agree that using multi-level outcomes will be an important question for future work and now explicitly note this in the manuscript, as below:

Using multi-level or continuous outcomes would also improve the ecological validity of the present approach and interpretation of the sensitivity parameters.

Appraisal and impact:

Overall this is a very strong paper, describing a novel task that could help move the field of RL forward to take account of threat processing more fully. The large sample size with discovery, replication and test-retest gives confidence in the findings. The task has good ecological validity and associations with task-based anxiety and clinical self-report demonstrate clinical relevance. The authors could give further context but test-retest of the punishment learning parameter is the only real concern. Overall this task provides an exciting new probe of reward/threat that could be used in mechanistic disease models.

We thank the reviewer again for helping us to improve our analyses and manuscript.

Reviewer #2 (Recommendations For The Authors):

Additional context:

In the introduction "cognitive tasks that bear little semblance to those used in the non-human literature" seems a little unfair. One study that is already cited (Ironside et al, 2020) used a task that was adapted from non-human primates for use in humans. It has almost identical visual stimuli (different levels of simultaneous reward and aversive outcome/punishment) and response selection processes (joystick) between species and some overlapping brain regions were activated across species for conflict and aversiveness. The later point that non-human animals must be trained on the association between action and outcome is well taken from the point of view of computational validity but perhaps not sufficient to justify the previous statement.

Our intention for this section was to briefly convey that historically, approach-avoidance conflict would have been measured either using questionnaires or joystick-based tasks which have no direct non-human counterpart. However, we agree that this phrasing is unfair to recent studies such as those by Ironside and colleagues. Therefore, we have amended the text to the following:

In humans, on the other hand, approach-avoidance conflict has historically been measured using questionnaires such as the Behavioural Inhibition/Activation Scale (Carver & White, 1994), or cognitive tasks that rely on motor biases to approach/move towards positive stimuli and avoid/move away from negative stimuli which have no direct non-human counterparts (Guitart-Masip et al., 2012; Kirlic et al., 2017; Mkrtchian et al., 2017; Phaf et al., 2014).

It would be good to speculate on why task induced anxiety made participants slower to update their estimates of punishment probability.

Although a meta-analysis of reinforcement learning studies using reward and punishment outcomes suggests a positive association between punishment learning rate and anxiety symptoms (and depressed mood), we paradoxically found the opposite effect. However, previous work has suggested that distinct forms of anxiety associate differently with anxiety (Wise & Dolan, 2020, Nat. Commun.), where somatic anxiety was negatively correlated with punishment learning rate whereas cognitive anxiety showed the opposite effect. We have now added the following to the manuscript, and noted that future work is needed to understand the potentially complex relationship between anxiety and learning from punishments:

Notably, although a recent computational meta-analysis of reinforcement learning studies showed that symptoms of anxiety and depression are associated with elevated punishment learning rates (Pike & Robinson, 2022), we did not observe this pattern in our data. Indeed, we even found the contrary effect in relation to task-induced anxiety, specifically that anxiety was associated with lower rates of learning from punishment. However, other work has suggested that the direction of this effect can depend on the form of anxiety, where cognitive anxiety may be associated with elevated learning rates, but somatic anxiety may show the opposite pattern (Wise & Dolan, 2020) and this may explain the discrepancy in findings. Additionally, parameter values are highly dependent on task design (Eckstein et al., 2022), and study designs to date may be more optimised in detecting differences in learning rate (Pike & Robinson, 2022) – future work is needed to better understand the potentially complex association between anxiety and punishment learning rate. Lastly, as punishment learning rate was severely unreliable in the test-retest analyses, and the associations between punishment learning rate and state anxiety were not robust to an alternative method of parameter estimation (variational Bayesian inference), the negative correlation observed in our study should be treated with caution.

Were those with more task-based anxiety more inflexible in general?

The lack of associations across reward learning rate and task-induced anxiety suggest that this was not a general inflexibility effect. To test the reviewer's hypothesis more directly, we conducted a sensitivity analysis by examining the model with a general learning rate – this did not support a general inflexibility effect. Please see the new section in the Supplement below:

[9.10 Sensitivity analysis: anxiety and inflexibility]

As anxious participants were slower to update their estimates of punishment probability, we determined whether this was due to greater general inflexibility by examining the model including two sensitivity parameters, but one general learning rate (i.e. not split by outcome). The correlation between this general learning rate and task-induced anxiety was not significant in either samples (discovery: $\tau = -0.02$, $p = 0.504$; replication: $\tau = -0.01$, $p = 0.625$), suggesting that the effect is specific to punishment.

Was the 16% versus 20% of the two samples with clinically relevant anxiety symptoms significantly different? What about other demographics in the two samples?

The difference in proportions were not significantly different ($\chi^2 = 2.33$, $p = 0.127$). The discovery sample included more females and was older on average compared to the replication sample – information which we now report in the manuscript:

The discovery sample consisted of a significantly greater proportion of female participants than the replication sample (59% vs 52%, $\chi^2 = 4.64$, $p = 0.031$). The average age was significantly different across samples (discovery sample mean = 37.7, SD = 10.3, replication sample mean = 34.3, SD = 10.4; $t_{785.5} = 5.06$, $p < 0.001$). The differences in self-reported psychiatric symptoms across samples did not reach significance ($p > 0.086$).

It would be interesting to know how many participants failed the audio attention checks.

We have now included information about what proportion of participants fail each of the task exclusion criteria in the manuscript:

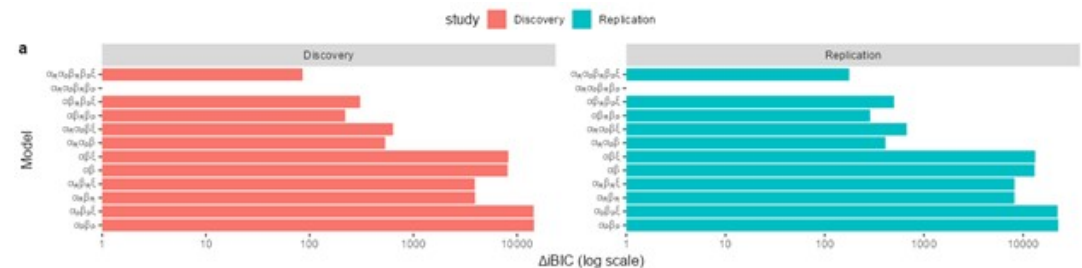
Firstly, we excluded participants who missed a response to more than one auditory attention check (see above; 8% in both discovery and replication samples) – as these occurred infrequently and the stimuli used for the checks were played at relatively low volume, we

allowed for incorrect responses so long as a response was made. Secondly, we excluded those who responded with the same response key on 20 or more consecutive trials ($> 10\%$ of all trials; 4/6% in discovery and replication samples, respectively). Lastly, we excluded those who did not respond on 20 or more trials (1/2% in discovery and replication samples, respectively). Overall, we excluded 51 out of 423 (12%) in the discovery sample, and 98 out of 725 (14%) in the replication sample.

There doesn't appear to be a model with only learning from punishment (i.e. no reward learning) included in the model comparison. It would be interesting to see how it compared.

We have fitted the suggested model and found that it is the least parsimonious of the models. Since participants were monetarily incentivised based on the rewards only, this was to be expected. We have now added this 'punishment learning only' model and its variant including a lapse term into the model comparison. The two lowest bars on the y-axis in Author response image 2 represent these models.

Author response image 2.



Were sex effects examined as these have been commonly found in AAC tasks. How about other covariates such as age?

We have now tested the effects of sex and age on behaviour and on parameter values. There were indeed some significant effects, albeit with some inconsistencies across the two samples, which for completeness we have included in the manuscript, as follows:

While sex was significantly associated with choice in the discovery sample ($\beta = 0.16 \pm 0.07$, $p = 0.028$) with males being more likely to choose the conflict option, this pattern was not evident in the replication sample ($\beta = 0.08 \pm 0.06$, $p = 0.173$), and age was not associated with choice in either sample ($p > 0.2$).

Comparing parameters across sexes via Welch's t-tests revealed significant differences in reward sensitivity ($t_{289} = -2.87$, $p = 0.004$, $d = 0.34$; lower in females) and consequently reward-punishment sensitivity index ($t_{336} = -2.03$, $p = 0.043$, $d = 0.22$; lower in females i.e. more avoidance-driven). In the replication sample, we observed the same effect on reward-punishment sensitivity index ($t_{626} = -2.79$, $p = 0.005$, $d = 0.22$; lower in females). However, the sex difference in reward sensitivity did not replicate ($p = 0.441$), although we did observe a significant sex difference in punishment sensitivity in the replication sample ($t_{626} = 2.26$, $p = 0.024$, $d = 0.18$).

Minor: Still a few placeholders (Supplementary Table X/ Table X) in the methods

We thank the reviewer for spotting these errors. We have now corrected these references.

Reviewer #3 (Public Review):

This study investigated cognitive mechanisms underlying approach-avoidance behavior using a novel reinforcement learning task and computational modelling. Participants could select a risky "conflict" option (latent, fluctuating probabilities of monetary reward and/or unpleasant sound [punishment]) or a safe option (separate, generally lower probability of reward). Overall, participant choices were skewed towards more rewarded options, but were also repelled by increasing probability of punishment. Individual patterns of behavior were well-captured by a reinforcement learning model that included parameters for reward and punishment sensitivity, and learning rates for reward and punishment. This is a nice replication of existing findings suggesting reward and punishment have opposing effects on behavior through dissociated sensitivity to reward versus punishment.

Interestingly, avoidance of the conflict option was predicted by self-reported task-induced anxiety. This effect of anxiety was mediated by the difference in modelled sensitivity to reward versus punishment (relative sensitivity). Importantly, when a subset of participants were retested over 1 week later, most behavioral tendencies and model parameters were recapitulated, suggesting the task may capture stable traits relevant to approach-avoidance decision-making.

We thank the reviewer for their useful analysis of our study. Indeed, it was reassuring to see that performance indices were reliable across time.

However, interpretation of these findings are severely undermined by the fact that the aversiveness of the auditory punisher was largely determined by participants, with the far-reaching impacts of this not being accounted for in any of the analyses. The manipulation check to confirm participants did not mute their sound is highly commendable, but the thresholding of punisher volume to "loud but comfortable" at the outset of the task leaves substantial scope for variability in the punisher delivered to participants. Indeed, participants' ratings of the unpleasantness of the punishment was moderate and highly variable ($M = 31.7$ out of 50, $SD = 12.8$ [distribution unreported]). Despite having this rating, it is not incorporated into analyses. It is possible that the key finding of relationships between task-induced anxiety, reward-punishment sensitivity and avoidance are driven by differences in the punisher experienced; a louder punisher is more unpleasant, driving greater task-induced anxiety, model-derived punishment sensitivity, and avoidance (and vice versa). This issue can also explain the counterintuitive findings from re-tested participants; lower/negatively correlated task-induced anxiety and punishment-related cognitive parameters may have been due to participants adjusting their sound settings to make the task less aversive (retest punisher rating not reported). It can therefore be argued that the task may not actually capture meaningful cognitive/motivational traits and their effects on decision-making, but instead spurious differences in punisher intensity.

We thank the reviewer for raising this important potential limitation of our study. We agree that how participants self-adjusted their sound volume may have important consequences for our interpretations of the data. Unfortunately, despite the scalability of online data collection, this highlights one of its major weaknesses in the lack of controllability over experimental parameters. The previous paper from which we obtained our aversive sounds (Seow & Hauser, 2021, Behav Res, doi.org/10.3758/s13428-021-01643-0) contains useful analyses with regards to this discussion. When comparing the unpleasantness of the sounds played at 50% vs 100% volume, the authors indeed found that the lower volumes lead to lower unpleasantness ratings. However, the magnitude of this effect did not appear to be substantial (Fig. 4 from the paper), and even at 50% volume, the scream sounds we used were

rated in the top quartile for unpleasantness, on average. This implies that the sounds have sufficient inherent unpleasantness, even when played at half intensity. We find this reassuring, in the sense that any self-imposed volume effects may not be large. Of note, our instructions to participants to adjust the volume to a ‘loud but comfortable’ level was based on the same phrasing used in this study.

To the reviewers point on how this might affect the reliability of the task, we have included the following in the ‘Discussion’ section:

Alternatively, as participants self-determined the loudness of the punishments, differences in volume settings across sessions may have impacted the reliability of this parameter (and indeed other measures).

Please see below for analyses accounting for punishment unpleasantness ratings.

This undercuts the proposed significance of this task as a translational tool for understanding anxiety and avoidance. More information about ratings of punisher unpleasantness and its relationship to task behavior, anxiety and cognitive parameters would be valuable for interpreting findings. It would also be of interest whether the same results were observed if the aversiveness of the punisher was titrated prior to the task.

As suggested, we have now included sensitivity analyses using the unpleasantness ratings that show their effect is minimal on our primary inference. We report relevant results below in the ‘Recommendations For The Authors’ section. At the same time, we think it is important to acknowledge that unpleasantness is a combination of both the inherent unpleasantness of the sound and the volume it is presented at, where only the latter is controlled by the participant. Therefore, these analyses are not a perfect indicator of the effect of participant control. For convenience, we reproduce the key findings from this sensitivity analysis here:

Approach-avoidance hierarchical logistic regression model

We assessed whether approach and avoidance responses, and their relationships with state anxiety, were impacted by punishment unpleasantness, by including unpleasantness ratings as a covariate into the hierarchical logistic regression model. Whilst unpleasantness was a significant predictor of choice (positively predicting safe option choices), all significant predictors and interaction effects from the model without unpleasantness survived (Supplementary Figure 11). Critically, this suggests that punishment unpleasantness does not account for all of the variance in the relationship between anxiety and avoidance.

Mediation model

When unpleasantness ratings were included in the mediation models, the mediating effect of the reward-punishment sensitivity index did not survive (discovery sample: standardised $\beta = 0.003 \pm 0.003$, $p = 0.416$; replication sample: standardised $\beta = 0.004 \pm 0.003$, $p = 0.100$; Supplementary Figure 12). Pooling the samples resulted in an effect that narrowly missed the significance threshold (standardised $\beta = 0.004 \pm 0.002$, $p = 0.068$).

More generally, whether or not to titrate the punishments (and indeed the rewards) is an interesting experimental decision, which we think should be guided by the research question. In our case, we were interested in individual differences in reward/punishment learning and sensitivity and their relation to anxiety, so variation in how aversive the sounds affected approach-avoidance decisions was an important aspect of our design. In studies where the aim is to understand more general processes of how humans act under approach-avoidance conflict, it may be better to tightly control the salience of reinforcers.

Ultimately, the best test of the causal role of anxiety on avoidance, and against the hypothesis that our results were driven by spurious volume control effects, would be to run within-

subjects anxiety interventions, where these volume effects are naturally accounted for. This will be an important direction for future studies using similar measures. We have added a paragraph in the ‘Discussion’ section on this point:

Relatedly, participants had some control over the intensity at which the punishments were presented, which may have driven our findings relating to anxiety and putative mechanisms of anxiety-related avoidance. Sensitivity analyses showed that our finding that anxiety is positively associated with avoidance in the task was robust to individual differences in self-reported punishment unpleasantness, whilst the mediation effects were not. Future work imposing better control over the stimuli presented, and/or using within-subjects designs will be needed to validate the role of reward/punishment sensitivities in anxiety-related avoidance.

Although the procedure and findings reported here remain valuable to the field, claims of novelty including its translational potential are perhaps overstated. This study complements and sits within a much broader literature that investigates roles for aversion and cognitive traits in approach-avoidance decisions. This includes numerous studies that apply reinforcement learning models to behavior in two-choice tasks with latent probabilities of reward and punishment (e.g., see doi: 10.1001/jamapsychiatry.2022.0051), as well as other translationally-relevant paradigms (e.g., doi: 10.3389/fpsyg.2014.00203, 10.7554/eLife.69594, etc).

We agree with the reviewer that our approach builds on previous work in reinforcement learning, approach-avoidance conflict and translational measures of anxiety. Whilst there are by now many studies using two-choice learning tasks with latent reward and punishment probabilities, our main, and which we refer to as ‘novel’, aim was to bring these fields together in such a way so as to model anxiety-related behaviour.

We note that we do not make strong statements about whether these effects speak to traits per se, and as Reviewer 1 notes, the evidence from our study suggests that the present measure may be better suited to assessing state anxiety. While computational model parameters can and are certainly often interpreted as constituting stable individual traits, a more simple interpretation of our findings may be that state anxiety is associated with a momentary preference for punishment avoidance over reward pursuit. This can still be informative for the study of anxiety, especially given the notion of a continuous relationship between adaptive/state anxiety and maladaptive/persistent anxiety.

Having said that, we agree with the underlying premise of the reviewer’s point that how the measure relates to trait-level avoidance/inhibition measures will be an interesting question for future work. We appreciate the importance of using tasks such as ours and those highlighted by the reviewer as trait-level measures, especially in computational psychiatry. We have now included a discussion on the potential roles of cognitive/motivational traits, in line with the reviewer’s recommendation – briefly, we have included the suggested references by the reviewer, discussed the measure’s potential relevance to cognitive/motivational traits, and direct interested readers to the broader literature. Please see below for details.

Reviewer #3 (Recommendations For The Authors):

As stated in the public review, punisher unpleasantness and its relationship to key findings (including for retest) should be reported and discussed.

We signpost readers to our new analyses, incorporating unpleasantness ratings into the statistical models, from the main manuscript as follows:

Since participants self-determined the volume of the punishments in the task, and therefore (at least in part) their aversiveness, we conducted sensitivity analyses by accounting for self-reported unpleasantness ratings of the punishment (see the Supplement). Our finding that anxiety impacts approach-avoidance behaviour was robust to this sensitivity analysis ($p < 0.001$), however the mediating effect of the reward-sensitivity sensitivity index was not ($p > 0.1$; see Supplement section 9.9 for details).

We reproduce the relevant section from the Supplement below. Overall, we found that the effect of anxiety on choices (via its interaction with punishment probability) remained significant after accounting for unpleasantness, however the mediating effect of reward-punishment sensitivity was no longer significant when unpleasantness ratings were included in the model. As noted above, unpleasantness ratings are not a perfect measure of self-imposed sound volume, and indeed punishment sensitivity is essentially a computationally-derived measure of unpleasantness, which makes it difficult to interpret the mediation model which contains both of these measures. However, since we found that anxiety affected choice over and above and effects of self-imposed sound volume (using unpleasantness ratings as a proxy measure), we argue that the task still holds value as a model of anxiety-related avoidance.

[Supplement Section 9.9: Sensitivity analyses of punishment unpleasantness]

Distribution of unpleasantness

The punishments were rated as unpleasant by the participants, on average (discovery sample: mean rating = 31.1 [scored between 0 and 50], SD = 13.1; replication sample: mean rating = 32.1, SD = 12.7; Supplementary Figure 10).

Approach-avoidance hierarchical logistic regression model

We assessed whether approach and avoidance responses, and their relationships with state anxiety, were impacted by punishment unpleasantness, by including unpleasantness ratings as a covariate into the hierarchical logistic regression model. Whilst unpleasantness was a significant predictor of choice (positively predicting safe option choices), all significant predictors and interaction effects from the model without unpleasantness ratings survived (Supplementary Figure 11). Critically, this suggests that punishment unpleasantness does not account for all of the variance in the relationship between anxiety and avoidance.

Mediation model

When unpleasantness ratings were included in the mediation models, the mediating effect of the reward-punishment sensitivity index did not survive (discovery sample: standardised $\beta = 0.003 \pm 0.003$, $p = 0.416$; replication sample: standardised $\beta = 0.004 \pm 0.003$, $p = 0.100$; Supplementary Figure 12). Pooling the samples resulted in an effect that narrowly missed the significance threshold (standardised $\beta = 0.004 \pm 0.002$, $p = 0.068$).

Test-retest reliability of unpleasantness

The test-retest reliability of unpleasantness ratings was excellent (ICC(3,1) = 0.75), although participants gave significantly lower ratings in the second session ($t_{56} = 2.7$, $p = 0.008$, $d = 0.37$; mean difference of 3.12, SD = 8.63).

Reliability of other measures with/out unpleasantness

To assess the effect of accounting for unpleasantness ratings on reliability estimates of task performance, we extracted variance components from linear mixed models, following a standard approach (Nakagawa et al., 2017) – note that this was not the method used to estimate reliability values in the main analyses, but we used this specific approach to

compare the reliability values with and without the covariate of unpleasantness ratings. The results indicated that unpleasantness ratings did not have a material effect on reliability (Supplementary Figure 14).

We discuss the findings of these sensitivity analyses in the ‘Discussion’ section, as follows:

Relatedly, participants had some control over the intensity at which the punishments were presented, which may have driven our findings relating to anxiety and putative mechanisms of anxiety-related avoidance. Sensitivity analyses showed that our finding that anxiety is positively associated with avoidance in the task was robust to individual differences in self-reported punishment unpleasantness, whilst the mediation effects were not. Future work imposing better control over the stimuli presented, and/or using within-subjects designs will be needed to validate the role of reward/punishment sensitivities in anxiety-related avoidance.

Introduction and discussion should spend more time relating the task and current findings to existing procedures and findings examining individual differences in avoidance and cognitive/motivational correlates.

We thank the reviewer for the opportunity to expand on the literature. Whilst there are numerous behavioural paradigms in both the human and non-human literature that involve learning about rewards and punishments, our starting point for the introduction was the state-of-the-art in translational models of approach-avoidance conflict models of anxiety. Therefore, for the sake of brevity and logical flow of our introduction, we have opted to bring in the discussion on other procedures primarily in the ‘Discussion’ section of the manuscript.

We have now included the reviewer’s suggested citations from their ‘Public Review’ as follows:

Since we developed our task with the primary focus on translational validity, its design diverges from other reinforcement learning tasks that involve reward and punishment outcomes (Pike & Robinson, 2022). One important difference is that we used distinct reinforcers as our reward and punishment outcomes, compared to many studies which use monetary outcomes for both (e.g. earning and losing £1 constitute the reward and punishment, respectively; Aylward et al., 2019; Jean-Richard-Dit-Bressel et al., 2021; Pizzagalli et al., 2005; Sharp et al., 2022). Other tasks have been used that induce a conflict between value and motor biases, relying on prepotent biases to approach/move towards rewards and withdraw from punishments, which makes it difficult to approach punishments and withdraw from rewards (Guitart-Masip et al., 2012; Mkrtchian et al., 2017). However, since translational operant conflict tasks typically induce a conflict between different types of outcome (e.g. food and shocks/sugar and quinine pellets; Oberrauch et al., 2019; van den Bos et al., 2014), we felt it was important to implement this feature. One study used monetary rewards and shock-based punishments, but also included four options for participants to choose from on each trial, with rewards and punishments associated with all four options (Seymour et al., 2012). This effectively requires participants to maintain eight probability estimates (i.e. reward and punishment at each of the four options) to solve the task, which may be too difficult for non-human animals to learn efficiently.

We have also included a discussion on the measure’s potential relevance to cognitive/motivational traits as follows:

Finally, whilst there is a broad literature on the roles of behavioural inhibition and avoidance tendency traits on decision-making and behaviour (Carver & White, 1994; Corr, 2004; Gray, 1982), we did not replicate the correlation of experiential avoidance and avoidance responses or the reward-punishment sensitivity index. Since there were also no significant correlations across task performance indices and clinical symptom measures, our findings suggest that

the measure may be more sensitive to behaviours relating to state anxiety, rather more stable traits. Nevertheless, how performance in the present task relates to other traits such as behavioural approach/inhibition tendencies (Carver & White, 1994), as has been found in previous studies on reward/punishment learning (Sharp et al., 2022; Wise & Dolan, 2020) and approach-avoidance conflict (Aupperle et al., 2011), will be an important question for future work.

We also now direct readers to a recent, comprehensive review on applying computational methods to approach-avoidance behaviours in the ‘Introduction’ section:

A fundamental premise of this approach is that the brain acts as an information-processing organ that performs computations responsible for observable behaviours, including approach and avoidance (for a recent review on the application of computational methods to approach-avoidance conflict, see Letkiewicz et al., 2023).

I am curious why participants were excluded if they made the same response on 20+ consecutive trials. How does this represent a cut-off between valid versus invalid behavioral profiles?

We apologise for the lack of clarity on this point in our original submission – this exclusion criterion was specifically if participants used the same response key (e.g. the left arrow button) on 20 or more consecutive trials, indicating inattention. Since the left-right positions of the stimuli were randomised across trials, this did not exclude participants who repeatedly chose the same option frequently. However, as we show in the Supplement, this, along with the other exclusion criteria, did not affect our main findings.

We have now clarified this as follows:

... we excluded those who responded with the same response key on 20 or more consecutive trials (> 10% of all trials; 4%/6% in discovery and replication samples, respectively) – note that as the options randomly switched sides on the screen across trials, this did not exclude participants who frequently and consecutively chose a certain option.