

1  
2  
3  
4  
5  
6  
7  
8  
9  
10  
11  
12  
13  
14  
15  
16  
17  
18  
19  
20  
21

# Supporting information for

## **A computational model for individual differences in non-reinforced learning**

Tom Salomon<sup>+</sup>, Alon Itzkovitch<sup>+</sup>, Nathaniel D. Daw, Tom Schonberg<sup>\*</sup>

<sup>+</sup> These authors contributed equally to this work.

<sup>\*</sup>Corresponding author. Email [Schonberg@tauex.tau.ac.il](mailto:Schonberg@tauex.tau.ac.il)

### **This PDF file includes:**

Supplementary Text  
Figs. S1 to S18  
Codes S1 to S5  
Table S1  
References

## Supplementary Text

### Unpublished data included in the CAT meta-analysis

The meta-analysis of CAT included data from 28 different CAT experiments; nine of these experiments include data, which were not published as they were collected for pilot testing before running another published study or would be published in the future. In this section we describe the unique features in these experiments and their results.

Overall, in most experiments we found that CAT resulted in enhanced preferences for the associated Go stimuli. The results of the studies are summarized in Supplementary Fig. S10 and its corresponding Supplementary Table 1. For each experiment below we describe its unique context and methods, as well as the results if those differed from the general trend of the enhanced preference effect for CAT stimuli over NoGo stimuli of similar initial value.

The proportion of trials in which participants chose Go over NoGo stimuli was analyzed using a mixed model logistic regression (with random intercept modeled per participant). For consistency with previous work with the CAT task (Bakkour et al., 2016; Schonberg et al., 2014), we analyzed separately choices between high-value stimuli, and choices between low-value stimuli. However, based on more recent finding with non-consumable stimuli as well as work that applied a slightly modified mixed-model analysis on the data (Botvinik-Nezer et al., 2021; Salomon et al., 2018, 2019), we did not expect to find a significant difference between the two value categories. Thus, we did not directly compare the effect of value category.

*Experiments 10-11: familiar faces.* In two experiments we performed the CAT study using stimuli of familiar faces of famous past and present Israeli parliament member. The stimuli set comprised of equal proportion of left-wing, right-wing, and central-trending political figures. The aim of these experiments was to examine the CAT effect with stimuli for which participants would have strong initial preference variability (e.g., a participant with right-wing tendencies should have strong preference for right-wing politicians over left-wing politicians). The procedure followed the protocol of similar CAT experiments with non-consumable stimuli – initial preferences evaluation using binary choices, training phase with 20 training runs, and probe phase in which Go stimuli were pitted against NoGo stimuli with similar initial value (both either high-value or low-value). Both experiments were identical in design, as the second experiment was designated

to be a direct replication of the first experiment ( $n = 25$ ) using larger sample size in the replication experiment ( $n = 39$ ).

*Experiments 13-14: emotional faces.* In two experiments, participants were exposed to stimuli set of 80 faces from the Karolinska face dataset (Lundqvist et al., 1998); the stimuli set comprised of 40 unique figures (21 male and 19 female), posing a neutral facial expression and a positive facial expression. Participants performed a CAT study design in which (unbeknownst to the participants) preferences were evaluated for each emotional category separately. During probe, participants chose from pairs of similar value and same facial expression (happy or neutral), in which one stimulus was a Go stimulus. One important goal of these experiments was to examine whether CAT had a stronger preference modification effect for faces of positive affect than neutral one.

In both experiments we found that CAT enhanced preferences for the associated Go stimuli, both for the face stimuli of positive affect and neutral affect). However, contrary to our initial hypothesis, we did not consistently observe a stronger preference modification effect for the positive affective face stimuli (see Supplementary Fig. S10 and Supplementary Table 1).

*Experiments 16: snacks adolescence.* In this CAT experiment with snack food stimuli, the effect of CAT was probed in a group of adolescent participants and young adults, aged 14 – 20 (mean age 18.2, SD = 1.90). The experiment protocol followed that of other CAT experiments with snacks (Bakkour et al., 2017; Schonberg et al., 2014), using a relatively short training of 12 training runs. The effect of CAT on preferences was not fully replicated in this sample. While participants demonstrated enhanced preferences for the low-value Go stimuli over low-value NoGo stimuli, for the high-value Go stimuli an enhanced preferences effect was not observed (see Supplementary Fig. S10 and Supplementary Table 1). This unexpected result could be due to less robust effects associated with the shorter training procedure or unique features of the population sampled this study.

*Experiments 19-20: CAT pilot experiments with snacks.* These two experiments were used as pilot experiments, before running Experiment 28 in fMRI setting (Botvinik-Nezer et al., 2020). In Experiment 19 a short training procedure with 12 training runs was used, which did not result in a significant preference modification effect (Supplementary Figure S9 and Supplementary Table S1). Following these results, in the second pilot experiments, which we refer to in this work as

Experiment 20, a longer training procedure with 20 training runs was used. This procedural design was eventually selected for the final imaging experiment (Experiment 28).

An additional sample ( $n = 24$ ) with a monetary reward cue was included in a previous draft of this manuscript (see [www.biorxiv.org/content/10.1101/2022.03.20.484477v2](http://www.biorxiv.org/content/10.1101/2022.03.20.484477v2); Exp. 17). This sample was excluded from this report to include only non-reinforced studies; however, its removal had no impact on any of the reported statistical results or conclusions.

### Using simple alternative marker of learning

In an additional analysis, conducted following all other analyses, we examined the robustness of  $\theta_{slope_i}$  as a prediction of subsequent probe choices above and beyond a simple alternative marker. As an alternative to  $\theta_{slope_i}$ , we calculated a simple anticipation score per participant. Training RTs were classified as ‘anticipatory’ based on the threshold of the bottom 1% of all RTs during the first training runs. In the meta-analysis this threshold was RT of 144ms or less, while in Study 2’s experiments this threshold was 210.2ms. We then calculated per participant the proportion of RTs below this threshold and used that proportion score as a simple alternative marker, which should model similar patterns as the  $\theta_{slope_i}$ . We also repeated the process for individual stimuli, to define a participant-stimulus alternative marker comparable with  $\theta_{slope_{i,s}}$ , and in Study 2, a proportion score was calculated for each condition (50% and 100% contingency).

We first examined the similarity of the proportion of anticipatory responses score with  $\theta_{slope_i}$ . In the meta-analysis, the proportion of anticipatory responses was found to be strongly correlated with  $\theta_{slope_i}$ ,  $r = 0.66$ ,  $t(838) = 25.64$ ,  $p < .001$ , 95%CI [0.62, 0.70]. However, in study 2 preliminary exp., the proportion of anticipatory responses was positively correlated with  $\theta_{slope_{50\%_i}}$  ( $r = 0.57$ ,  $t(18) = 1.93$ ,  $p = .008$ , 95%CI [0.17, 0.98]), but to a lesser degree with  $\theta_{slope_{100\%_i}}$  ( $r = 0.41$ ,  $t(18) = 2.15$ ,  $p = .08$ , 95%CI [-0.05, 0.86]). In Study 2 replication exp. there was no significant association between anticipatory responses and  $\theta_{slope_i}$  within each condition; replication exp. 50% condition:  $r = 0.11$ ,  $t(57) = 0.8$ ,  $p = .428$ , 95%CI [-0.159, 0.369]; replication exp. 100% condition:  $r = 0.10$ ,  $t(57) = 0.78$ ,  $p = .441$ , 95%CI [-0.161, 0.366]. Examining the correlation plots (see Supplementary Fig. S11), it seems that very low  $\theta_{slope_i}$  were commonly

associated with close to 0% anticipatory responses throughout the training, however the proportion of anticipatory responses did not increase monotonically with  $\theta_{slope_i}$ .

To evaluate the predictive power of the proportion of anticipatory responses as a simple alternative to  $\theta_{slope_i}$ , we repeated the same logistic regression models using the proportion of anticipatory responses as a sole independent variable instead of  $\theta_{slope_i}$ . We examined if the simple alternative marker could predict probe choices in the subsequent probe phase and compare the two models' AIC score. In Study 1, we found that the proportion of anticipatory responses were predictive of a preference modification effect (OR = 7.26, 95% CI = [3.45, 15.25],  $Z = 5.23$ ,  $p < .001$ ; one-sided mixed model logistic regression), and even provided better explanatory power compared with the original model using  $\theta_{slope_i}$  ( $\Delta AIC = -329.20$ ). This pattern was not found in Study 2, where the proportion of anticipatory responses were not found to be strongly indicative of a subsequent preference modification effect (preliminary exp. 50% contingency: OR = 2.07, 95% CI = [0.22, 19.28],  $Z = 0.64$ ,  $p = .26$ ; preliminary exp. 100% contingency: OR = 1.98, 95% CI = [0.24, 16.31],  $Z = 0.63$ ,  $p = .26$ ; replication exp. 50% contingency: OR = 0.94, 95% CI = [0.41, 2.15],  $Z = -0.64$ ,  $p = .55$ ; replication exp. 100% contingency: OR = 1.39, 95% CI = [0.54, 3.54],  $Z = 0.69$ ,  $p = .25$ ; one-sided mixed model logistic regression), and using the proportion of anticipatory responses provided worse explanatory power compared to  $\theta_{slope_i}$  (preliminary exp.  $\Delta AIC = 10.95$ ; replication exp.  $\Delta AIC = 11.01$ ).

To examine the added value of  $\theta_{slope_i}$  above and beyond the simple proportion of the anticipatory response marker, we examined the significance of  $\theta_{slope_i}$  in a nested model using both indicators to predict probe choices. In Study 1,  $\theta_{slope_i}$  provided no additional explanatory power above and beyond the simple proportion of anticipatory responses indicators ( $\Delta AIC = 1.81$ ,  $\chi^2_{(1)} = 0.19$ ,  $p = .662$ ; likelihood ratio test). In both experiments of Study 2,  $\theta_{slope_i}$  significantly improved a predictive model using only anticipatory responses (preliminary exp.  $\Delta AIC = 7.56$ ,  $\chi^2_{(2)} = 11.56$ ,  $p = .003$ ; replication exp.  $\Delta AIC = 7.83$ ,  $\chi^2_{(2)} = 11.83$ ,  $p = .003$ ; likelihood ratio test).

#### Using alternative log-normal Bayesian model

In a post-hoc analysis, conducted following all other analyses, we examined the impact of changing the Gaussian Bayesian model with a log-normal modeling of RT distribution. The model had the same restrictions as the gaussian one, meaning the mean of anticipatory RTs must be lower than that of the cue-dependent RTs.

A few adaptations were required to help the log-normal model to converge: (i) We added a fixed value of 1,000ms to the effective RT to avoid negative values. (ii) We changed the RT scale to seconds (divided by 1000, while keeping milliseconds precision). (iii) We had to exclude out of this analysis participants which had a large number of more than 30% non-responses. This resulted in exclusion of 26 participants from Study 1 meta-analyses and one participant from Study 2 replication experiment). The model is fully detailed in Supplementary Code S5.

#### *Meta analysis results:*

After running four independent chains, the Bayesian model converged to a stable solution (see Supplementary Fig. S12). The converged model estimated effective RTs as a time-dependent mixture of two RT distributions: one of early anticipatory responses ( $\mu_1 = 21.22\text{ms}$ , 95%CI [18.26, 26.03],  $\sigma_{\epsilon_1} = 298.23\text{ms}$ , 95%CI [295.37, 301.87]) and late cue-dependent responses ( $\mu_2 = 284.54\text{ms}$ , 95%CI [284.03, 284.92],  $\sigma_{\epsilon_2} = 60.24\text{ms}$ , 95%CI [59.93, 60.56]). Overall, participants demonstrated an increase in proportion of anticipatory responses as training progressed, as manifested in the positive group-level parameter ( $\theta_{\text{slope}} = 3.42$ , 95%CI [3.14, 3.72]), with variation between participants ( $\sigma_{\theta_{\text{slope}}} = 4.11$ , 95% CI [3.87, 4.35]).

Posterior predictive checks (simulated distributions of RTs based on estimated model parameters) revealed a good fit of the model to the actual data. Simulated posterior distributions recreated the patterns observed empirically, showing more rapid transition to anticipatory responses in participants with higher  $\theta_{\text{slope}_i}$  parameter estimates as well as capturing the skewed right tail shape of RTs (Supplementary Fig. S13).

Similarly to the two-gaussian model, we examined whether variation in the slope parameter fit to RTs correlated with probe performance. The meta-analysis showed  $\theta_{\text{slope}_i}$  was positively associated with the preference modification effect – i.e., participants with higher  $\theta_{\text{slope}_i}$  parameter estimates, also demonstrated greater odds of choosing Go stimuli (OR = 1.06, 95% CI = [1.04, 1.07],  $Z = 4.92$ ,  $p < .001$ ; one-sided mixed model logistic regression; Supplementary Fig. S14). The model's intercept was significantly greater than zero, i.e., even when extrapolated to very low  $\theta_{\text{slope}_i}$  values, the model forecasted enhanced preference for Go stimuli (intercept odds = 1.22, 95% CI = [1.15, 1.29],  $Z = 8.36$ ,  $p < .001$ ).

#### *Study 2 experiments results:*

After running four independent chains, the Bayesian model converged to a stable solution (see Supplementary Fig. S15). The converged model estimated effective RTs as a time-dependent mixture of two RT distributions: one of the early anticipatory responses (preliminary experiment:  $\mu_1 = -72.2ms$ , 95%CI[95.1,39.2],  $\sigma_{\epsilon_1} = 263.6ms$ , 95%CI[246.0,284.0]; replication experiment:  $\mu_1 = -90.6ms$ , 95%CI[-104.2, -29.5],  $\sigma_{\epsilon_1} = 264.9ms$ , 95%CI[258.6,277.6]), and late cue-dependent responses (preliminary experiment:  $\mu_2 = 339.1ms$ , 95%CI[336.4,340.4],  $\sigma_{\epsilon_2} = 58.6ms$ , 95%CI[51.2,59.7]; replication experiment:  $\mu_2 = 329.7ms$ , 95%CI[323.1,343.1],  $\sigma_{\epsilon_2} = 59.7$ , 95%CI[58.6,60.7]). Overall, participants demonstrated an increase in proportion of anticipatory responses as training progressed, as manifested in the positive group-level parameter (50% condition:  $\theta_{slope} = 1.59$ , 95%CI [0.92, 2.23]; 100% condition:  $\theta_{slope} = 3.87$ , 95%CI [2.38, 4.1]), with variation between participants (50% condition:  $\theta_{slope} = 1.82$ , 95%CI [1.38, 2.37]; 100% condition:  $\theta_{slope} = 2.77$ , 95%CI [2.25, 3.41]).

Posterior predictive checks (simulated distributions of RTs based on estimated model parameters) revealed a good fit of the model to the actual data. Simulated posterior distributions recreated the patterns observed empirically, showing more rapid transition to anticipatory responses in participants with higher  $\theta_{slope_i}$  parameter estimate (Supplementary Fig. S16).

Examining the association of  $\theta_{slope_i}$  parameter estimates with Go choices during probe, we found positive association with subsequent probe choices in the preliminary experiment (50% condition: log-OR = 0.03,  $Z = 2.45$ ,  $p = .007$ , OR = 1.03, 95% CI [1.02, 1.04]; 100% condition: log-OR = 0.04,  $Z = 3.75$ ,  $p < .001$ , OR = 1.04, 95% CI [1.03, 1.51]; two-sided mixed model logistic regression). In the replication study only in the 50% condition this parameter was found to be significant (log-OR = 0.017,  $Z = 2.12$ ,  $p = .038$ , OR = 1.01, 95% CI [1.008, 1.02]; 100% condition: log-OR = 0.003,  $Z = 0.71$ ,  $p = .24$ , OR = 1.003, 95% CI [0.998, 1.008]; one-sided mixed model logistic regression).

#### Using an alternative Bayesian model with unfixed $\theta_0$ parameter:

In a post-hoc analysis, conducted following all other analyses, we examined the impact of changing the model parameters, by adding a new parameter  $\theta_0$ , instead of the fixed one. The model had the same restrictions and priors. The adaptation required was to manage one  $\theta_0$  parameter for each contingency.

The conclusion did not change after the addition of the new parameter. We examined the correlation between choices, and the value  $\theta$  parameter would reach at the final training run (equal to  $\theta_0 + \theta_{slope}$ ). In both studies,  $\theta_0$  converged to a similar solution  $\theta_0 = -2 (\pm 2)$ , and  $\theta$  at the final training run predicted choices well.

#### Reaction time restrictions:

In a post-hoc analysis, conducted following all other analyses, we examined the models' performance under non-excluded missing values, as well as separating RT distributions (anticipatory and cue-dependent) with the minimum value of 1ms. We replaced the missing values with the averaged reaction time for the participant in a given run. This choice was made in order to diminish data corruption with irrelevant values. Also, we modeled restrictions on mean RT, which upper boundary of anticipatory responses set as 1ms below the lower boundary of cue-dependent RT. We found no changes in the model performances, both in convergence and prediction. It should be emphasized that the original model converged well with its restrictions, with a high separation boundary between the two distributions, thus the new restrictions should not have changed the original results. Moreover, missing values were around 6%, thus it also should not have affected results.

#### Controlling for initial value in Study 2

In an additional post-hoc analysis, we aimed to examine the impact of initial value on the preference modification effect during probe. In study 2 experiments, 80 stimuli were rank-ordered and categorized into ten value-groups of eight stimuli each (stimuli ranked 1-8 was categorized into the highest value group, stimuli 9-16 comprised the second-highest value group, etc.). The middle eight value groups, each contained four Go stimuli and four NoGo stimuli of similar initial value, which were pitted against each other in the probe trials (see details in methods section and Figure 9). Early CAT studies that used snack food stimuli and two value categories (high-value and low-value stimuli), reported a stronger preference modification effect for Go stimuli in the high value category, compared with the low-value stimuli (Bakkour et al., 2017; Schonberg et al., 2014). However, this differential value effect has not been consistently replicated in later experiments (Bakkour et al., 2016; Salomon et al., 2018, 2019), which also used non-consumable stimuli, similar to the face stimuli used in Study 2 of the current work.

Under the assumption that the effect of CAT on choices is consistent across value categories, we analyzed and reported in the main text the aggregated results from all value categories combined. To validate this assumption, we reanalyzed the probe data using categorical variables representing the eight value groups and examined the significance of an interaction term of value-group with contingency. To evaluate the statistical significance of the (value X contingency) interaction effect, we compared a full model consisting of both main effects and interaction, with a restricted model containing only the main effects. Similarly, to examine the main effect of value-group, we compared a model containing the two main effects to restricted model in which the coefficients associated with value-group were omitted.

In the preliminary experiment we did not find significant value-group and contingency interaction effects ( $\Delta AIC = -4.94$ ,  $\chi^2_{(7)} = 9.06$ ,  $p = .249$ ; likelihood ratio test). We did find a main effect for value group ( $\Delta AIC = 16.34$ ,  $\chi^2_{(7)} = 30.34$ ,  $p < .001$ ) as well as a main effect for contingency condition ( $\Delta AIC = 5.74$ ,  $\chi^2_{(1)} = 7.74$ ,  $p = .005$ ). In the replication experiment we identified a significant interaction effect ( $\Delta AIC = 4.35$ ,  $\chi^2_{(7)} = 18.35$ ,  $p = .010$ ), with no main effect for value group ( $\Delta AIC = -10.11$ ,  $\chi^2_{(7)} = 3.89$ ,  $p = .792$ ) and a significant main effect for contingency condition ( $\Delta AIC = 30.08$ ,  $\chi^2_{(1)} = 32.08$ ,  $p < .001$ ), see Supplementary Fig. S17.

Descriptively, the value group effects did not show a consistent or monotonic trend, neither within-experiment nor between experiments. In the preliminary experiment, a main effect of value group was manifested as smaller overall preference for Go stimuli in some value categories (e.g. 54.6% choices of Go stimuli in the 7<sup>th</sup> value group compared with 70.1% in the higher 5<sup>th</sup> value group; See Supplementary Fig. S17). This trend was not monotonic, i.e. we did not observe a more robust preference modification effect in the higher value groups, nor replicated in the following experiment. Likewise, in the replication experiment an interaction effect was detected, which indicated variability in the contingency effect across different value groups. Like in the preliminary experiment, this effect was not monotonic with the value group (higher value groups did not show consistently more or less robust contingency effects), nor was it found in the previous experiment. In contrast to the inconsistent effects of value group, the effect of cue-contingency was more robust and was replicated in both experiments, above and beyond the contribution of value categories.

In an additional post-hoc analysis, we further examined the impact of adding the (standardized) value difference between Go and NoGo stimuli as covariates to the probe analysis.

We found that in both experiments, value differences had a significant explanatory power when accounting for participant choices (preliminary experiment:  $OR = 2.00$ , 95%  $CI = [1.07, 1.33]$ ,  $Z = 3.24$ ,  $p = .002$ ; replication experiment:  $OR = 1.16$ , 95%  $CI = [1.09, 1.23]$ ,  $Z = 5.13$ ,  $p < .001$ ; two-sided mixed model logistic regression). Importantly however, the effect of enhanced preferences for Go stimuli remained consistent above and beyond value-differences effect (preliminary experiment, 50% contingency:  $OR = 1.33$ , 95%  $CI = [1.09, 1.62]$ ,  $Z = 2.86$ ,  $p = .004$ ; 100% contingency:  $OR = 1.81$ , 95%  $CI = [1.28, 2.57]$ ,  $Z = 3.38$ ,  $p < .001$ ; replication experiment, 50% contingency:  $OR = 1.41$ , 95%  $CI = [1.24, 1.59]$ ,  $Z = 5.50$ ,  $p < .001$ ; 100% contingency:  $OR = 1.97$ , 95%  $CI = [1.60, 2.43]$ ,  $Z = 6.40$ ,  $p < .001$ ; two-sided mixed model logistic regression).

#### Intent to treat approach for missing RT

In the computational model, trials in which participants did not respond to the Go cue were excluded from the analysis. These excluded trials accounted for 6.33% and 5.97% of trials in the preliminary and replication experiments, respectively. To make sure this decision had no critical impact on our results we performed a post-hoc analysis in which we implemented an intent to treat approach. Instead of excluding these trials, we allocated them an artificial value of  $RT =$

A repeated analysis with these previously excluded trials resulted in similar conclusions. The converged model estimated effective RTs as a time-dependent mixture of two RT distributions: one of early anticipatory responses (preliminary experiment:  $\mu_1 = -72.2ms$ , 95%  $CI[95.1, 39.2]$ ,  $\sigma_{\epsilon_1} = 263.6ms$ , 95%  $CI[246.0, 284.0]$ ; replication experiment:  $\mu_1 = -90.6ms$ , 95%  $CI[-104.2, -29.5]$ ,  $\sigma_{\epsilon_1} = 264.9ms$ , 95%  $CI[258.6, 277.6]$ ), and late cue-dependent responses (preliminary experiment:  $\mu_2 = 339.1ms$ , 95%  $CI[336.4, 340.4]$ ,  $\sigma_{\epsilon_2} = 58.6ms$ , 95%  $CI[51.2, 59.7]$ ; replication experiment:  $\mu_2 = 329.7ms$ , 95%  $CI[323.1, 343.1]$ ,  $\sigma_{\epsilon_2} = 59.7$ , 95%  $CI[58.6, 60.7]$ ). Overall, participants demonstrated an increase in proportion of anticipatory responses as training progressed, as manifested in the positive group-level parameter (50% condition:  $\theta_{slope} = 1.59$ , 95%  $CI [0.92, 2.23]$ ; 100% condition:  $\theta_{slope} = 3.87$ , 95%  $CI [2.38, 4.1]$ ), with variation between participants (50% condition:  $\theta_{slope} = 1.82$ , 95%  $CI [1.38, 2.37]$ ; 100% condition:  $\theta_{slope} = 2.77$ , 95%  $CI [2.25, 3.41]$ ).

Positive association with subsequent probe choice both in preliminary experiment (50% condition:  $\log-OR = 0.03$ ,  $Z = 2.45$ ,  $p = .024$ ,  $OR = 1.03$ , 95%  $CI [1.02, 1.04]$ ; 100% condition:

log-OR = 0.04,  $Z = 3.75$ ,  $p = .001$ , OR = 1.04, 95% CI [1.03, 1.51]; two-sided mixed model logistic regression) and in the replication experiment (50% condition: log-OR = 0.03,  $Z = 2.45$ ,  $p = .024$ , OR = 1.03, 95% CI [1.02, 1.04]; 100% condition: log-OR = 0.04,  $Z = 3.75$ ,  $p = .001$ , OR = 1.04, 95% CI [1.03, 1.51]; two-sided mixed model logistic regression).

#### Cross validation of Study 2 predictions

To validate the generalizability of our prediction model, we examined the effect of using out-of-sample data predictions. In a five-fold cross validation (CV) analysis, the data of all users were split into five groups. In each CV iteration, predictions of the proportion of Go choices was generated to the validation group, based on a mixed logistic regression model that was trained using the data of remaining four groups. In case our model suffered from high overfit, we would expect that the CV predictions would be dissimilar to the model reported in the main manuscript, which was based on the entire dataset.

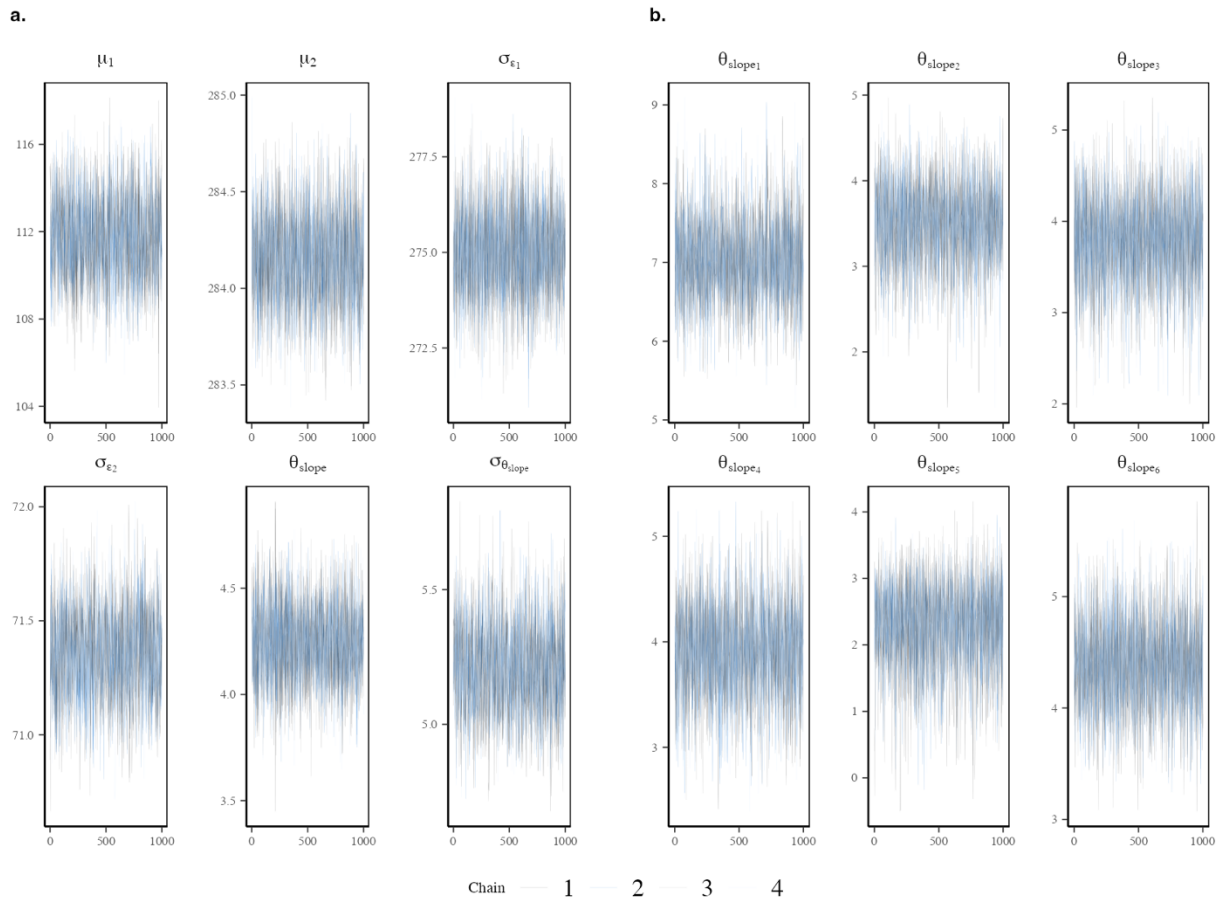
Examining the correlation between the predictions made by the full data model (reported in the main manuscript), to the predictions made for the held-back data in the CV analysis, we found that in both experiments, the two prediction models resulted in nearly identical predictions (preliminary experiment:  $r = 0.975$ , mean absolute prediction difference = 1.51%, range = [0.04%, 4.53%]; replication experiment:  $r = 0.994$ , mean absolute prediction difference = 0.53% range = [0.002%, 2.43%]; see Fig. S18.). These results imply that the prediction model was not sensitive to the specific data it was exposed to, thus these results are reassuring that the prediction model is likely to generalize well to out-of-sample data as well.

#### predictive power of $\theta_{slope_t}$

We found that choice RT often had a significant explanatory power, wherein faster choices were associated with higher likelihood of choosing to Go stimuli (preliminary experiment – 50% contingency: log-OR = 0.35,  $Z = 1.31$ ,  $p = .19$ , OR = 1.42, 95% CI [0.84, 2.40], 100% contingency: log-OR = -0.72,  $Z = -2.33$ ,  $p = .019$ , OR = 0.49, 95% CI [0.26, 0.89]; replication experiment – 50% contingency: log-OR = -0.58,  $Z = -3.50$ ,  $p < .001$ , OR = 0.56, 95% CI [0.41, 0.78], 100% contingency: log-OR = -1.01,  $Z = -5.59$ ,  $p < .001$ , OR = 0.36, 95% CI [0.26, 0.52]; two-sided

mixed logistic regression). Nonetheless  $\theta_{slope_i}$  was consistently found to be predictive of choosing Go stimuli, above and beyond choice RT (preliminary experiment – 50% contingency: log-OR = 0.12,  $Z = 2.44$ ,  $p = .007$ , OR = 1.13, 95% CI [1.02, 1.25], 100% contingency: log-OR = 0.26,  $Z = 3.36$ ,  $p < .001$ , OR = 1.30, 95% CI [1.12, 1.52]; replication experiment – 50% contingency: log-OR = 0.1,  $Z = 2.42$ ,  $p = .008$ , OR = 1.11, 95% CI [1.02, 1.21], 100% contingency: log-OR = 0.12,  $Z = 3.53$ ,  $p < .001$ , OR = 1.12, 95% CI [1.05, 1.20]).

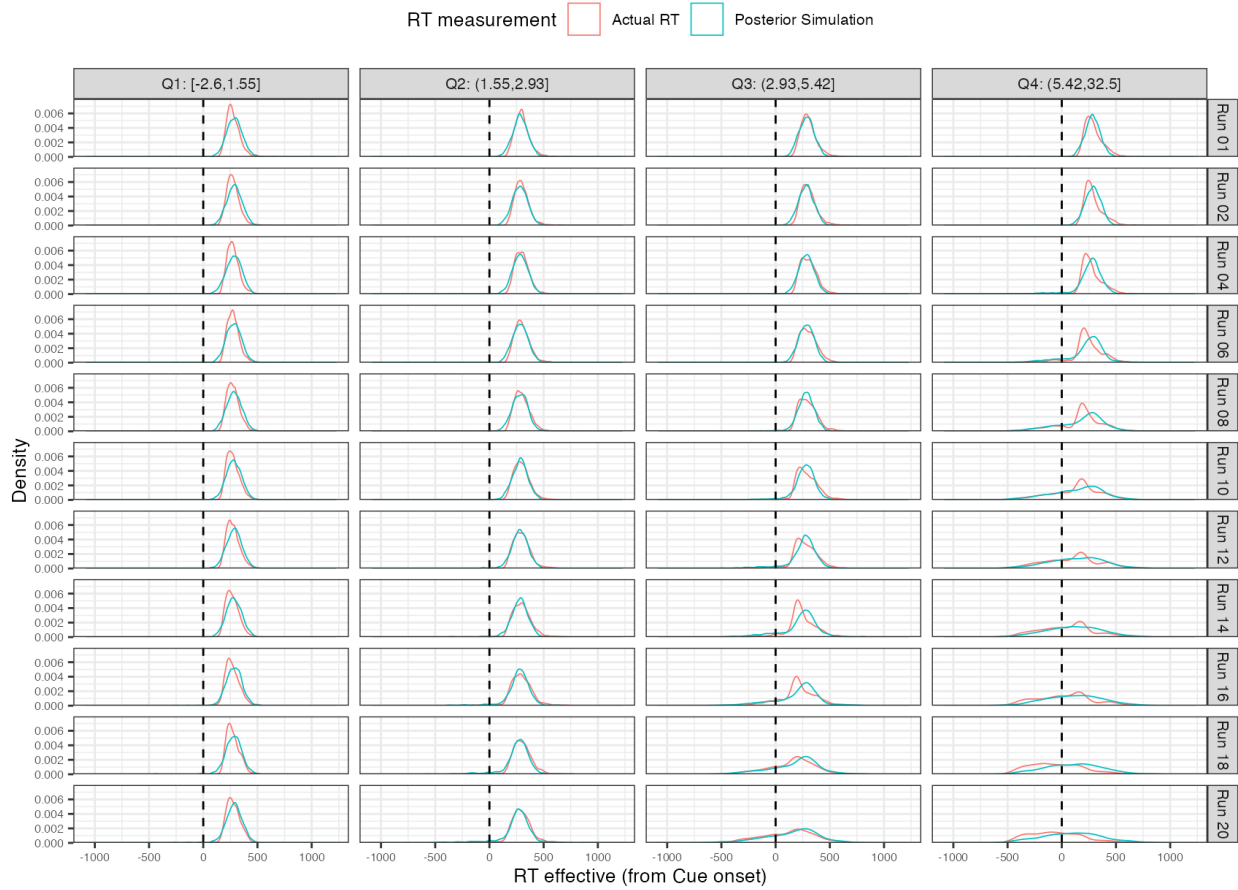
**Fig. S1.**



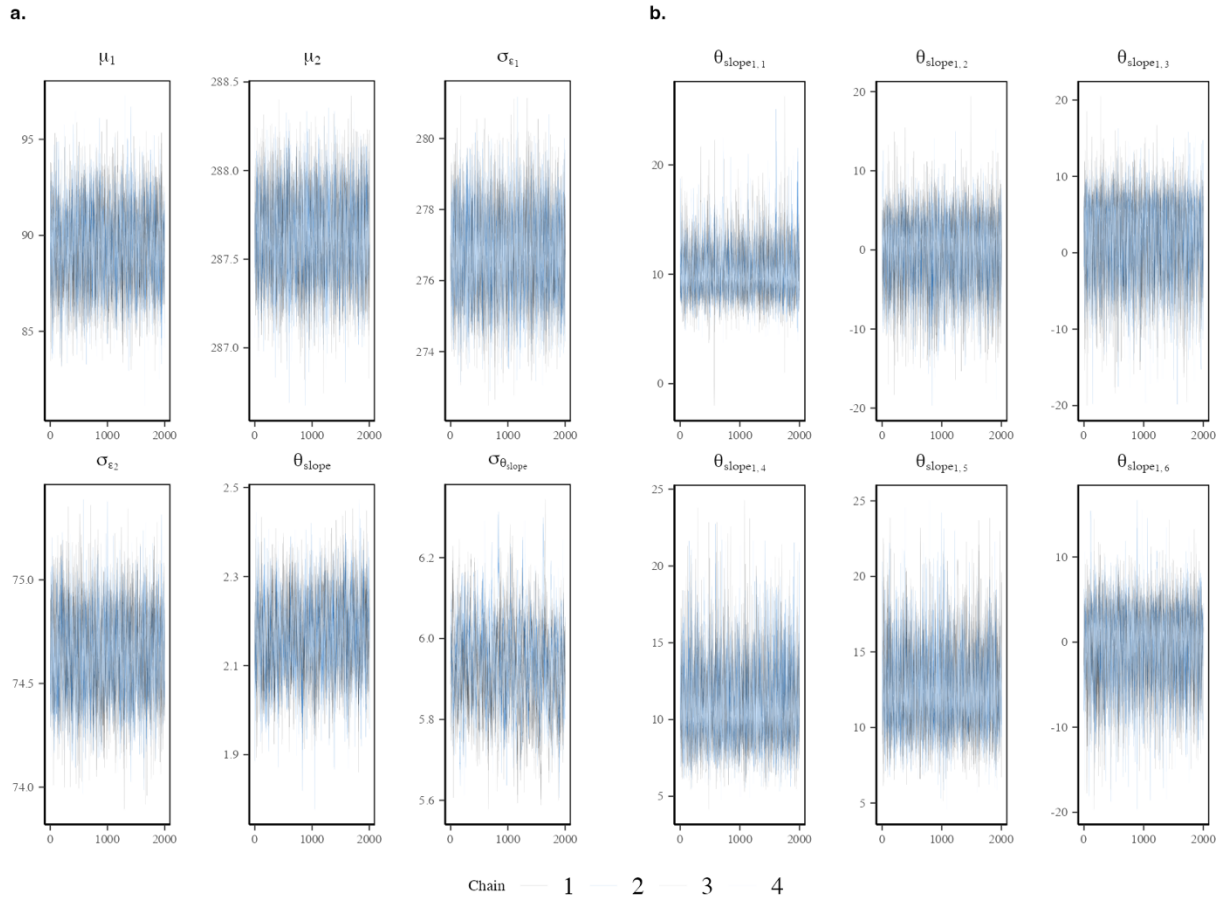
Trace plots of the Stan model of the meta-analysis. (a) Six hyper-parameters defined the shape of the RT data: two means ( $\mu_1, \mu_2$ ) and standard deviations ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ ) for the normal distribution of anticipatory and cue-dependent responses, and a mean and  $SD$  ( $\theta_{slope}, \sigma_{\theta_{slope}}$ ) for the distribution from which individualized  $\theta_{slope_i}$  parameters would be drawn for each participant. (b) Trace plot

of the first six participants'  $\theta_{slope_i}$  parameters. A well-converged parameter is characterized with all four chains (in different colors) reaching stable solution around the same estimate for each parameter ('hairy caterpillar' pattern).

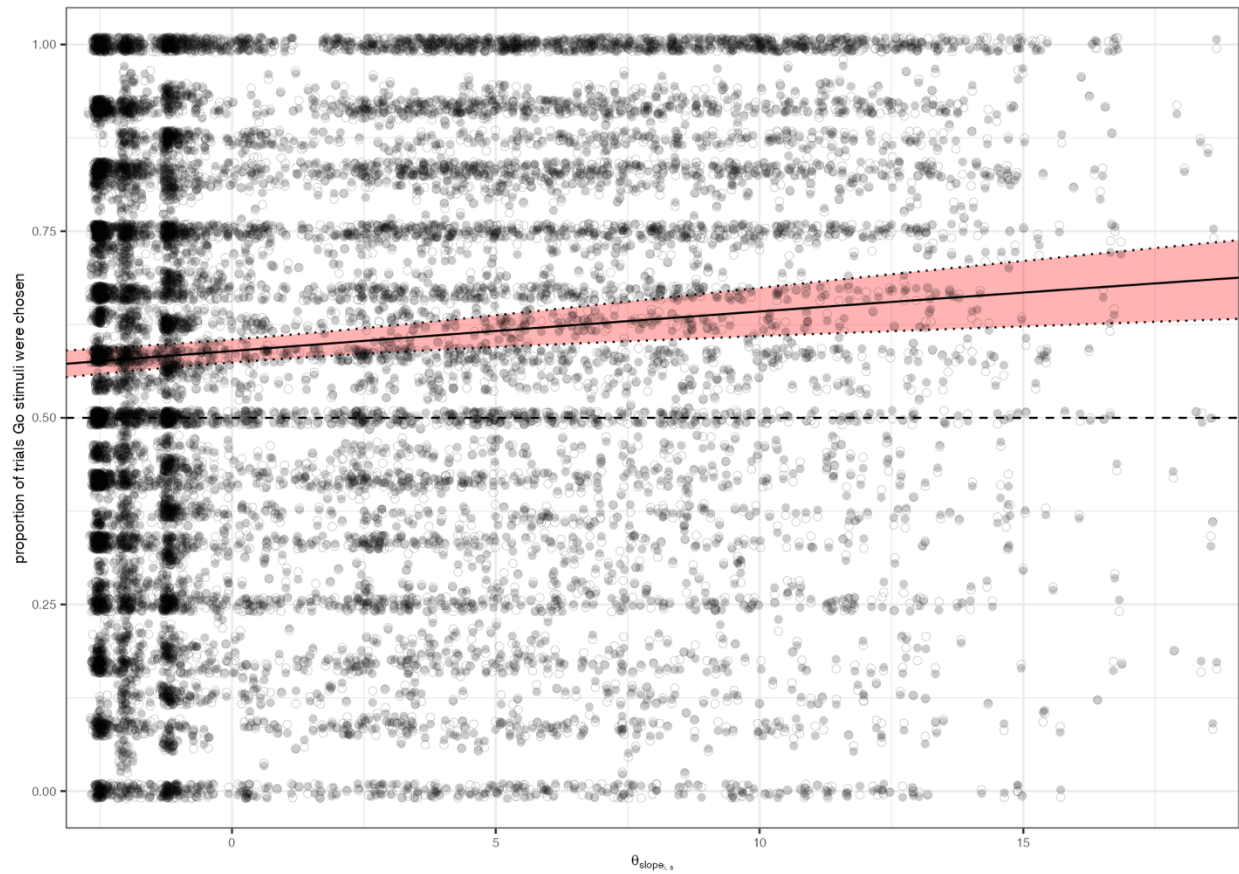
**Fig. S2.**



Actual RT distributions (red) versus simulated posterior distributions (teal), by  $\theta_{slope_i}$  quantile group and run. Participants were categorized into four equal quantile groups, according to their  $\theta_{slope_i}$  parameter estimates (denoted here as Q1-Q4; columns). Throughout the different training runs (rows; included here every other run) the posterior simulated RT distributions fitted well to the transition pattern.

343 **Fig. S3.**

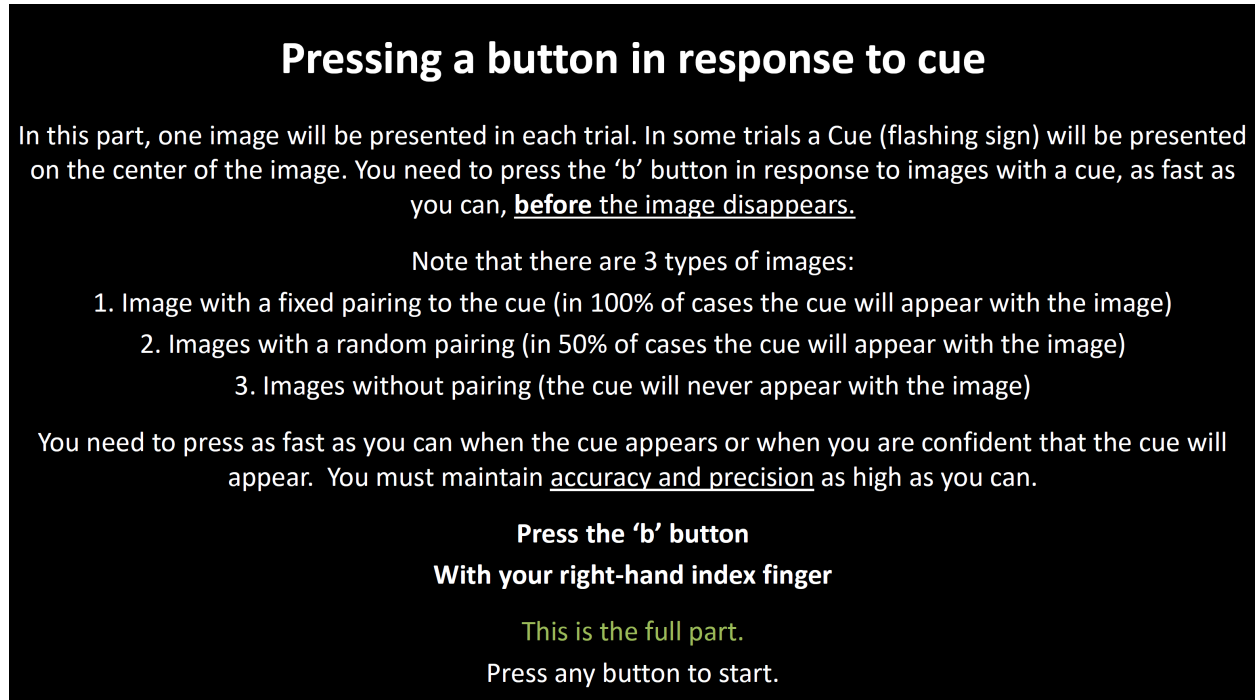
344 Trace plots of the meta-analysis Stan model, with stimulus-level  $\theta_{slope_{i,s}}$  parameter. (a) Six hyper-  
 345 parameters defined the shape of the RT data: two means ( $\mu_1, \mu_2$ ) and standard deviations ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ )  
 346 for the normal distribution of anticipatory and cue-dependent responses, and a mean and *SD*  
 347 ( $\theta_{slope}, \sigma_{\theta_{slope}}$ ) for the distribution from which individualized  $\theta_{slope_{i,s}}$  parameters were drawn for  
 348 each participant. (b) Trace plot of the first six  $\theta_{slope_{i,s}}$  parameters, associated with six stimuli of  
 349 the first participant in the meta-analysis. A well-converged parameter is characterized with all four  
 350 chains (in different colors) reaching stable solution around the same estimate for each parameter  
 351 ('hairy caterpillar' pattern).  
 352

353 **Fig. S4.**

354 Association of probe choices with stimulus-level computational marker in study 1. Larger  $\theta_{slope_{i,s}}$   
 355 estimates parameter estimates were positively associated with subsequent preference modification  
 356 in probe phase, across the different experiments. Dots represent individual stimuli (small vertical  
 357 jitter was added to better visualize high density areas). Trend lines and surrounding color margins  
 358 represent estimated preference modification effect and 95% CI, respectively (mixed model logistic  
 359 regression). Horizontal dashed line represents 50% chance level.  
 360  
 361

362

363 **Fig. S5.**



**Pressing a button in response to cue**

In this part, one image will be presented in each trial. In some trials a Cue (flashing sign) will be presented on the center of the image. You need to press the 'b' button in response to images with a cue, as fast as you can, before the image disappears.

Note that there are 3 types of images:

1. Image with a fixed pairing to the cue (in 100% of cases the cue will appear with the image)
2. Images with a random pairing (in 50% of cases the cue will appear with the image)
3. Images without pairing (the cue will never appear with the image)

You need to press as fast as you can when the cue appears or when you are confident that the cue will appear. You must maintain accuracy and precision as high as you can.

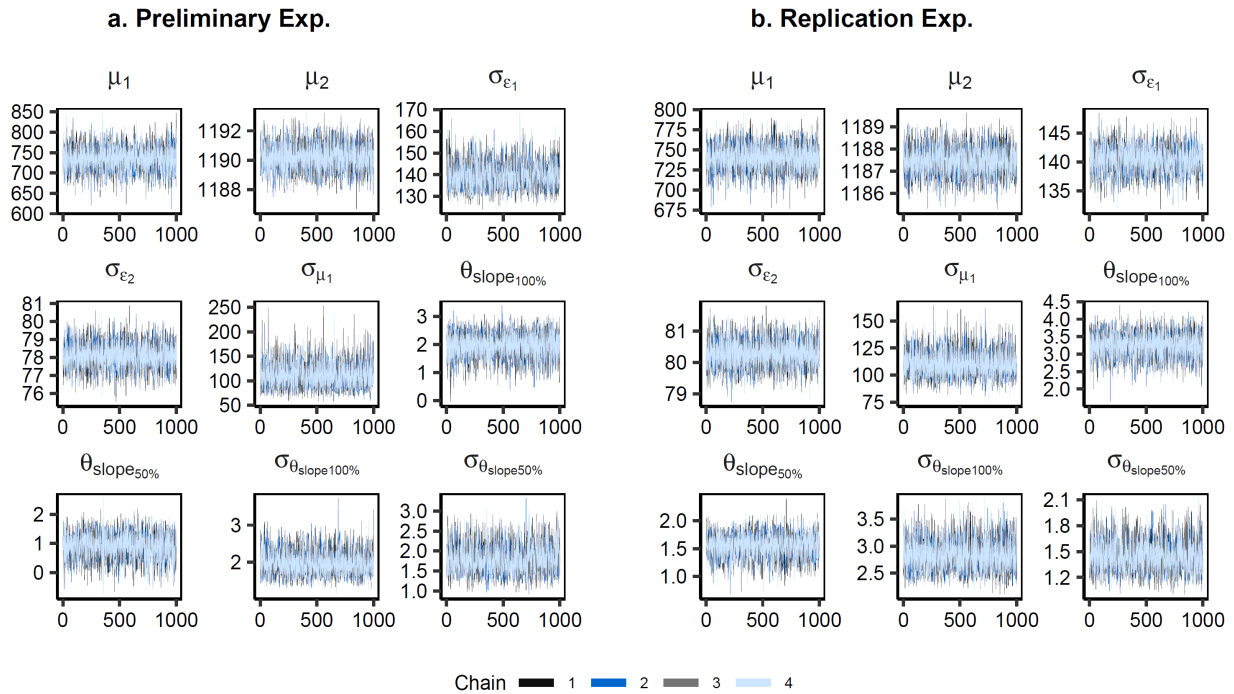
**Press the 'b' button**  
**With your right-hand index finger**

**This is the full part.**

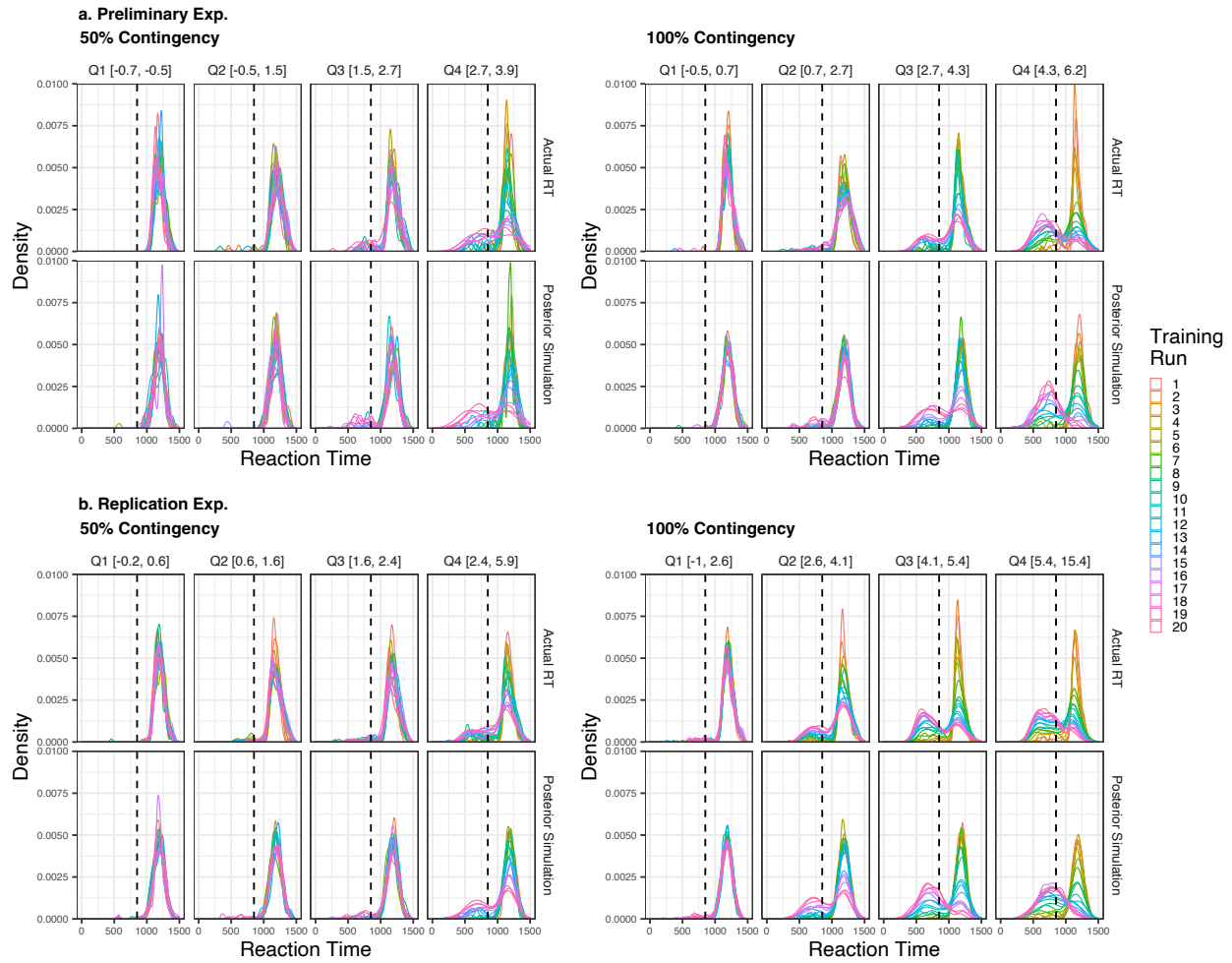
Press any button to start.

364

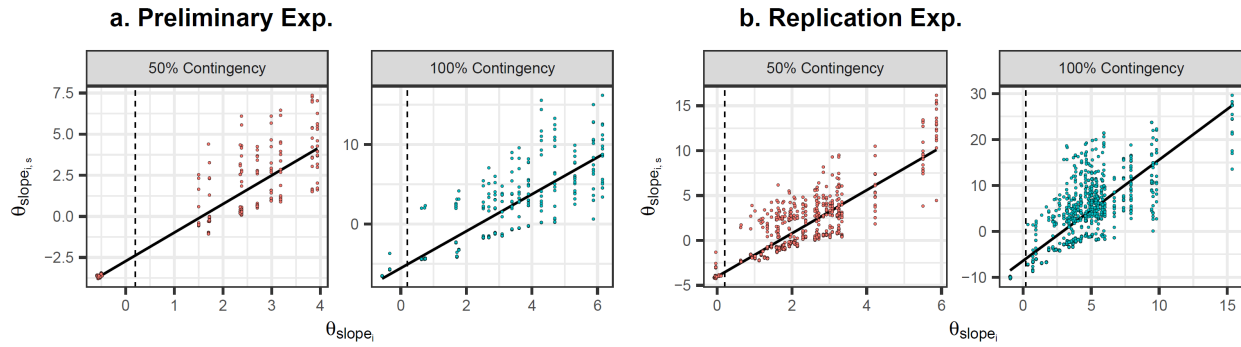
365 A translated version of the cue-approach training phase instructions presented on the screen before  
366 the CAT task in two experiments included in Study 2. Participants were clearly instructed that they  
367 may respond before the cue onset, when they anticipate that a cue will appear in association with  
368 the image.

369 **Fig. S6.**

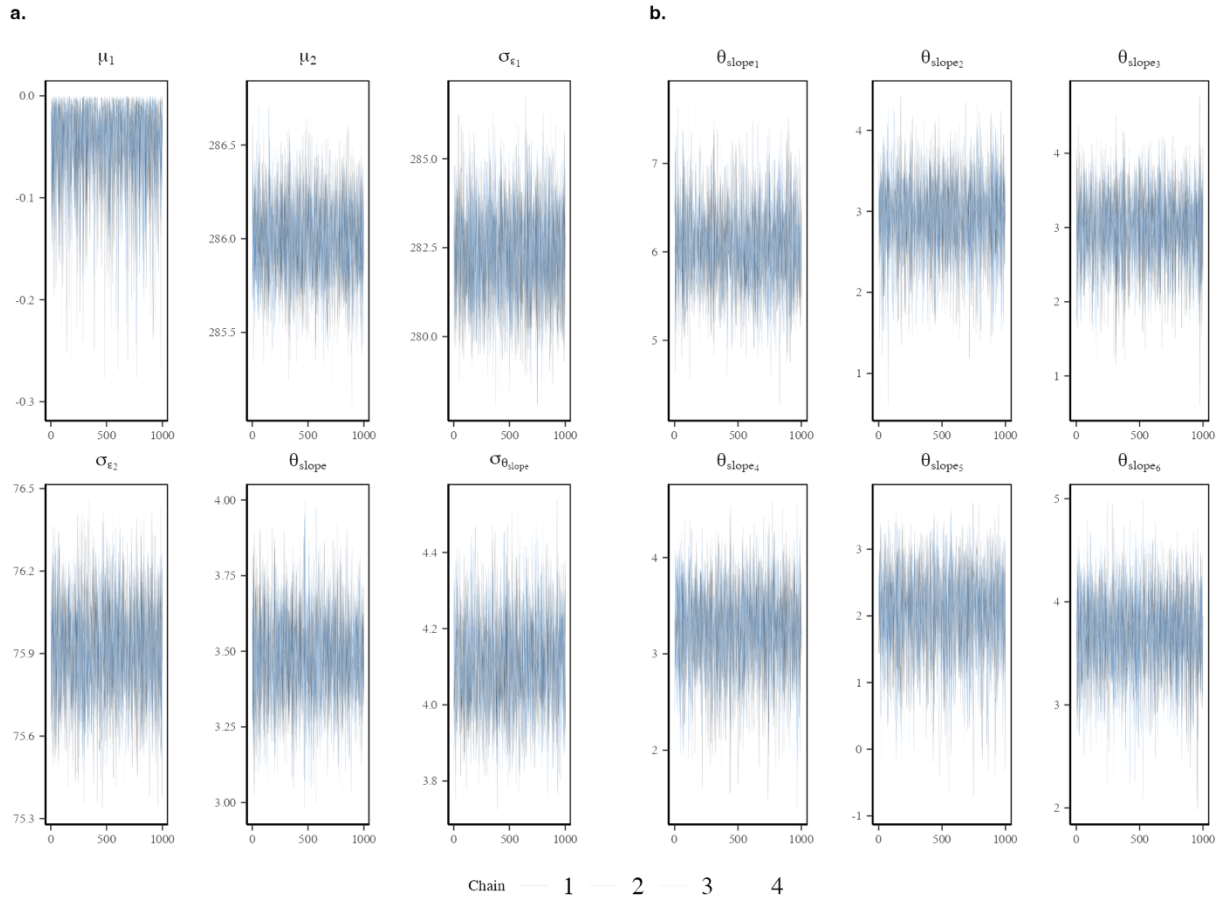
Trace plots of the Stan models in Study 2: (a) Preliminary experiment and (b) Replication experiment. Nine hyper-parameters defined the RT model: two means ( $\mu_1, \mu_2$ ) and  $SD$  ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ ) for the normal distribution of anticipatory and cue-dependent RTs, respectively. A mean ( $\theta_{slope_{100\%}}, \theta_{slope_{50\%}}$ ) and  $SD$  ( $\sigma_{\theta_{slope_{100\%}}}, \sigma_{\theta_{slope_{50\%}}}$ ) for the distribution from which  $\theta_{slope_i}$  parameters would be drawn for each  $i^{\text{th}}$  participant in the 100% and 50% contingency. A well-converged parameter is characterized with all four chains (in different colors) reaching stable solution around the same estimate for each parameter ('hairy caterpillar' pattern).

378 **Fig. S7.**

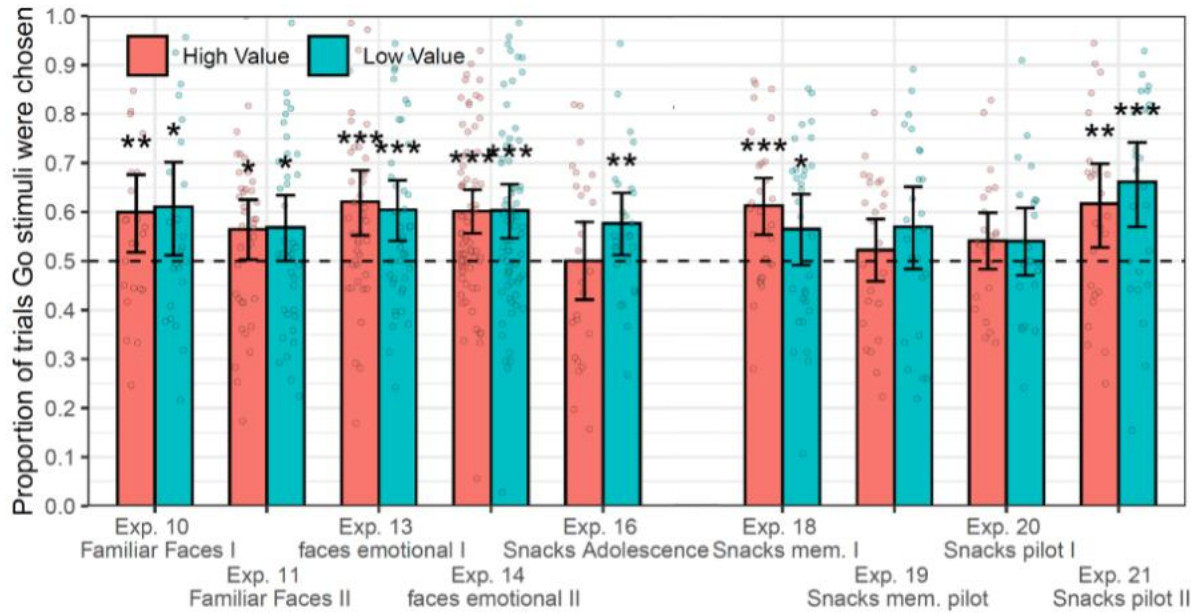
Actual RT distributions versus simulated posterior distributions, by  $\theta_{slope_i}$  quantile group, in Study 2 (a) preliminary experiment and (b) replication experiment. Participants were categorized into four equal quantile groups, according to their  $\theta_{slope_i}$  parameter estimates (denoted here as Q1-Q4; columns). Participants with higher  $\theta_{slope_i}$  were characterized with faster transition to anticipatory responses (top row). Posterior simulated RT distributions using mixture of Gaussians (bottom row) captured relatively well this transition pattern. Vertical dashed line represents cue-onset.

387 **Fig. S8.**

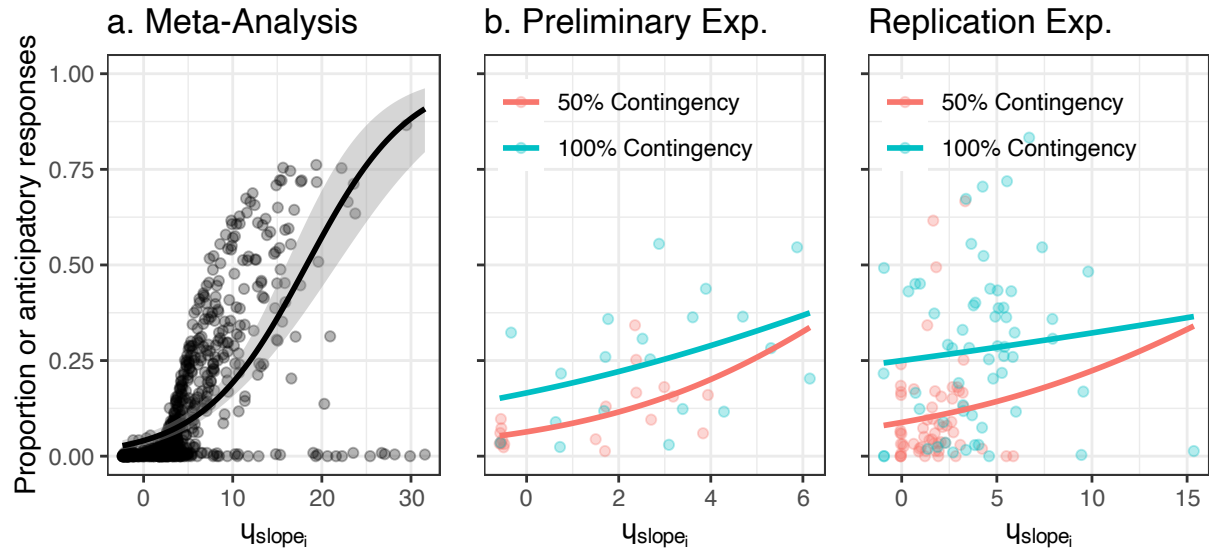
388  
 389 Correlation of  $\theta_{slope_i}$  parameter fitted individually for each participant and  $\theta_{slope_{i,s}}$  fitted  
 390 individually for each Go stimulus within each participant, in Study 2 preliminary experiment (a)  
 391 and replication experiment (b). Each dot represents a stimulus within participant. Each participant  
 392 was fitted with one  $\theta_{slope_i}$ , and 16  $\theta_{slope_{i,s}}$  parameter estimates, per contingency condition. Very  
 393 low  $\theta_{slope_i}$  parameter estimates were also characterized with low variability in  $\theta_{slope_{i,s}}$  estimates.  
 394 Participants with small  $\theta_{slope_i}$  estimate in either condition, were excluded from further logistic  
 395 regression analysis (vertical dashed line represents a threshold of  $\theta_{slope_i} = 0.2$ ).

396 **Fig. S9.**

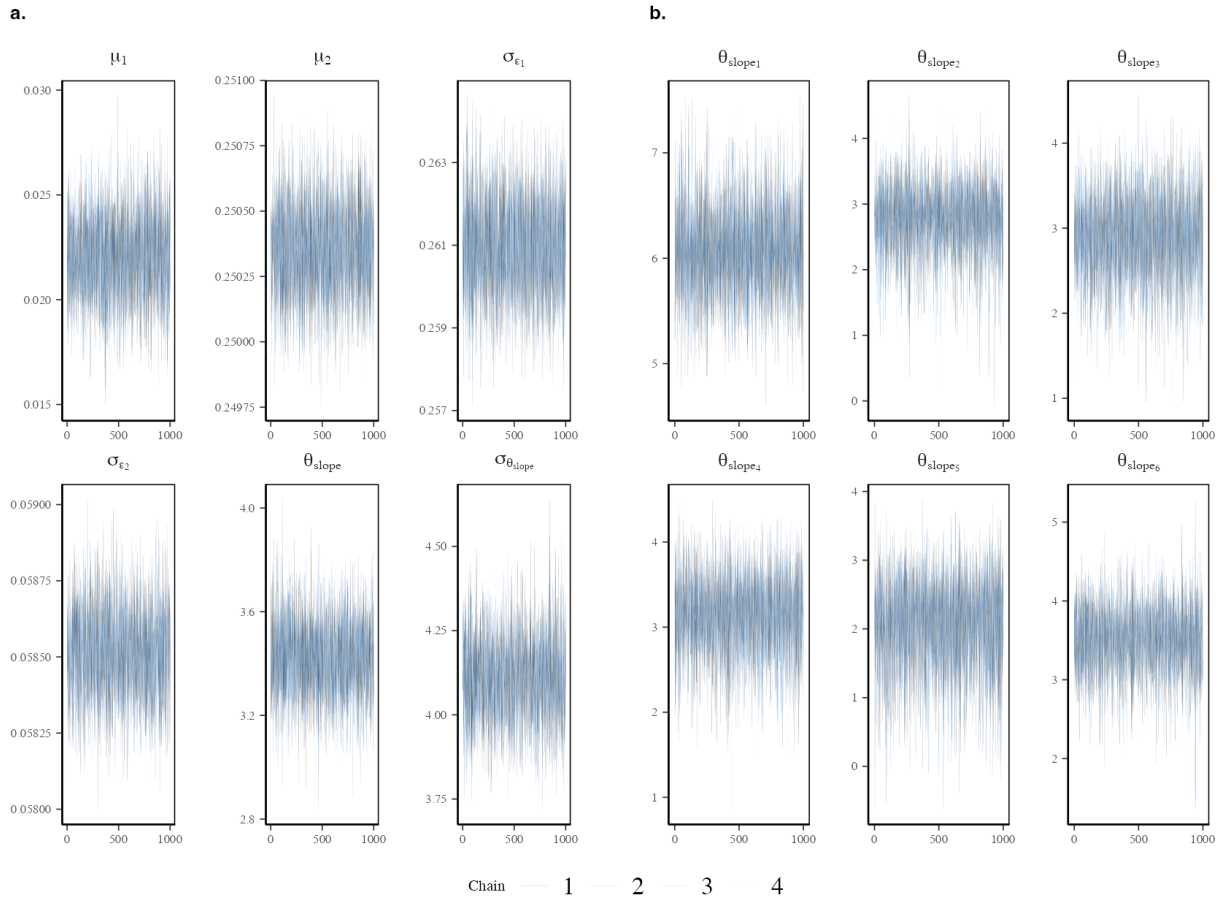
397 Trace plots of the Stan model of the meta-analysis, with upper limit of  $\mu_1 < 0$ . (a) Six hyper-  
 398 parameters defined the shape of the RT data: two means ( $\mu_1, \mu_2$ ) and standard deviations ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ )  
 399 for the normal distribution of anticipatory and cue-dependent responses, and a mean and *SD*  
 400 ( $\theta_{slope}, \sigma_{\theta_{slope}}$ ) for the distribution from which individualized  $\theta_{slope_i}$  parameters would be drawn  
 401 for each participant. The MCMC values of  $\mu_1$  converged around the upper limit restriction. (b)  
 402 Trace plot of the first six participants'  $\theta_{slope_i}$  parameters. A well-converged parameter is  
 403 characterized with all four chains (in different colors) reaching stable solution around the same  
 404 estimate for each parameter ('hairy caterpillar' pattern).  
 405

**Fig. S10.**

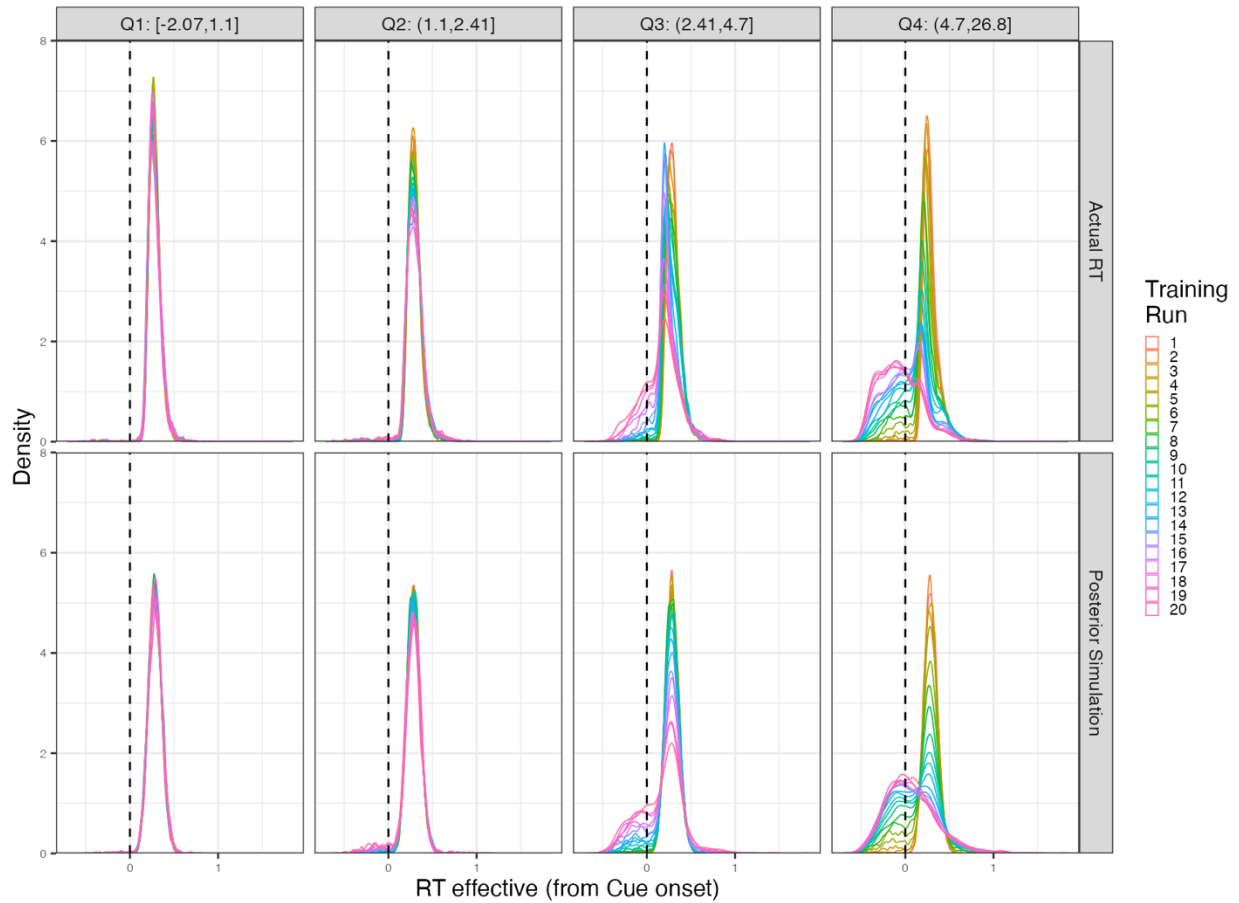
Probe results of unpublished CAT experiments. Proportion of trial participants chose Go stimuli over NoGo stimuli, by experiment and value category. Go stimuli were pitted against NoGo stimuli of similar initial value (both either high- or low-value). In Experiments 13 and 14 (emotional faces), the value category indicates the effect of facial expression (high value = positive affect, low value = neutral affect). Dots represent individual participants. Bars, error bars and asterisks represent the results of a logistic regression analysis, summarizing the mean estimated effect, 95% confidence interval and statistical significance, respectively. \*\*\* -  $p < 0.001$ , \*\* =  $p < 0.01$ , \* -  $p < 0.05$ ; one-sided mixed-model logistic regression. Models' results were transformed from log-odds to probabilities units (odds =  $\exp(\log\text{-odds})$ ; prob. =  $\text{odds}/(1 + \text{odds})$ ).

**Fig. S11.**

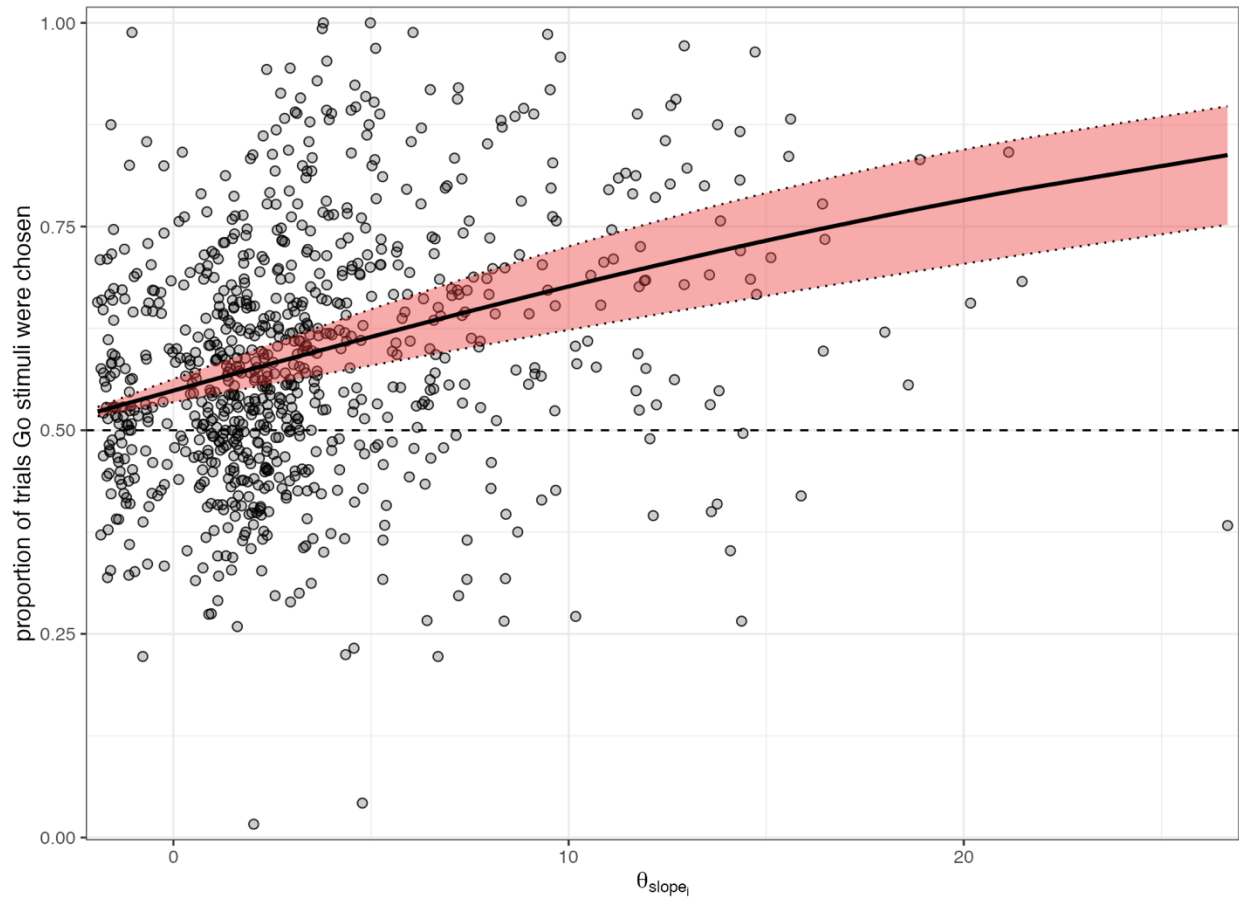
Association of  $\theta_{slope_i}$  computational parameter and a simple alternative marker of learning, using the proportion of anticipatory responses. (a) association in study 1 meta-analysis (anticipatory response threshold = 144.52ms). (b) association in study 2 preliminary and replication experiments (anticipatory response threshold = 210.2ms). Points represents individual participants; line represents a logistic regression model trendline predicting the proportion of anticipatory responses using  $\theta_{slope_i}$ .

427 **Fig. S12.**

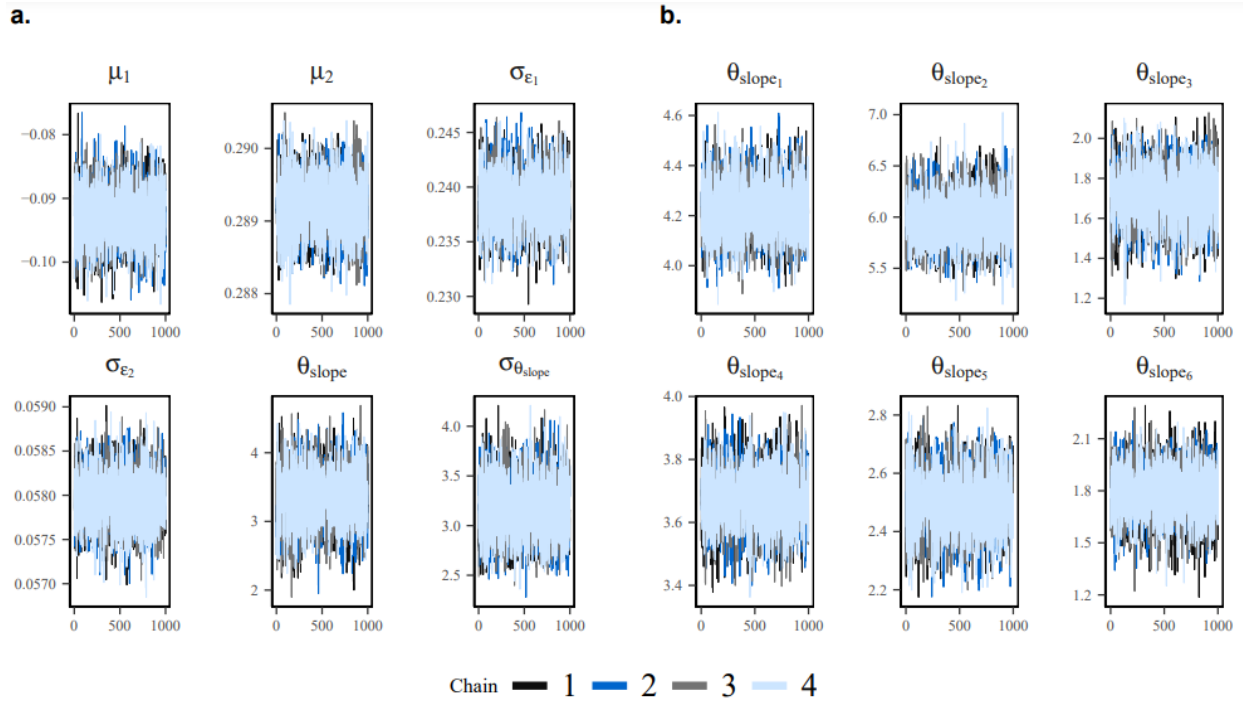
Trace plots of the Stan log-normal model of the meta-analysis. (a) Six hyper-parameters defined the shape of the RT data: two means ( $\mu_1, \mu_2$ ) and standard deviations ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ ) for the normal distribution of anticipatory and cue-dependent responses, and a mean and *SD* ( $\theta_{slope}, \sigma_{\theta_{slope}}$ ) for the distribution from which individualized  $\theta_{slope_i}$  parameters would be drawn for each participant. (b) Trace plot of the first six participants'  $\theta_{slope_i}$  parameters. A well-converged parameter is characterized with all four chains (in different colors) reaching stable solution around the same estimate for each parameter ('hairy caterpillar' pattern).

437 **Fig. S13.**

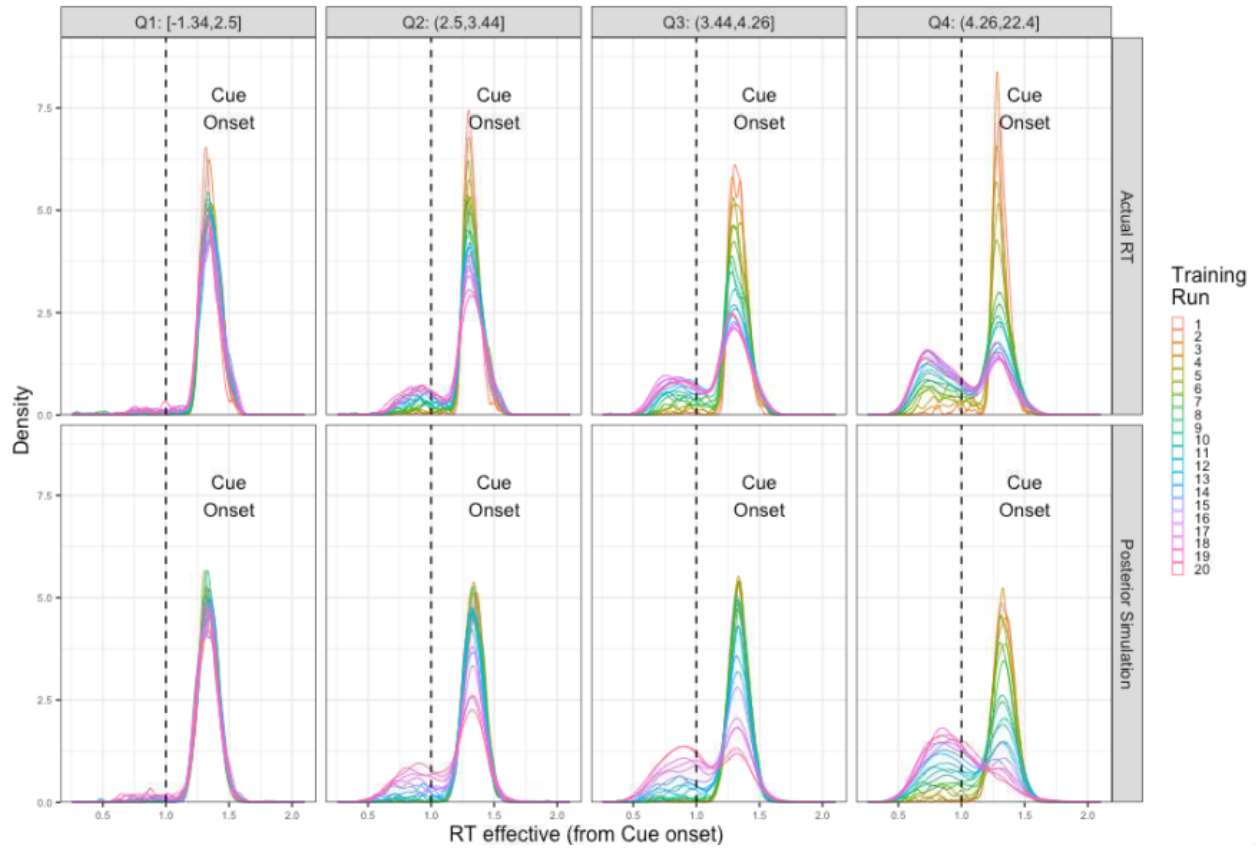
438  
 439 Actual RT distributions versus simulated posterior distributions, by  $\theta_{slope_i}$  quantile group for  
 440 Study 1 meta-analysis. Participants were categorized into four equal quantile groups, according to  
 441 their parameter estimates (denoted here as Q1-Q4; columns). Participants with higher parameter  
 442 estimates were characterized with faster transition to anticipatory responses (top row). Posterior  
 443 simulated RT distributions using mixture of log-normal (bottom row) recreated relatively well this  
 444 transition pattern. Vertical dashed line represents cue-onset.

445 **Fig. S14.**

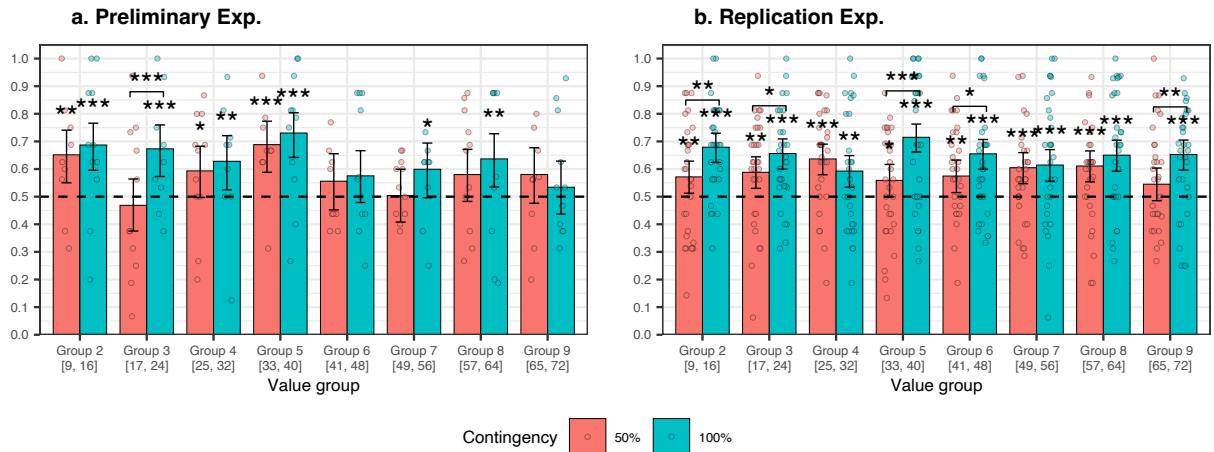
446  
 447 Meta-analysis results - computational marker and preference modification effect. Participants who  
 448 transitioned faster to anticipatory responses during cue-approach training (larger  $\theta_{slope_i}$  estimates)  
 449 also demonstrated stronger preference modification effects (proportion of trials Go stimuli were  
 450 chosen). Trend line and surrounding red margins represent estimated preference modification  
 451 effect and 95% CI, respectively (mixed model logistic regression). Dots represent individual  
 452 participants.  
 453

**Fig. S15.**

Trace plots of the Stan log-normal model of the designated experiments. (a) Six hyper-parameters defined the shape of the RT data: two means ( $\mu_1, \mu_2$ ) and standard deviations ( $\sigma_{\epsilon_1}, \sigma_{\epsilon_2}$ ) for the normal distribution of anticipatory and cue-dependent responses, and a mean and *SD* ( $\theta_{\text{slope}}, \sigma_{\theta_{\text{slope}}}$ ) for the distribution from which individualized  $\theta_{\text{slope}_i}$  parameters would be drawn for each participant. (b) Trace plot of the first six participants'  $\theta_{\text{slope}_i}$  parameters. A well-converged parameter is characterized with all four chains (in different colors) reaching stable solution around the same estimate for each parameter ('hairy caterpillar' pattern).

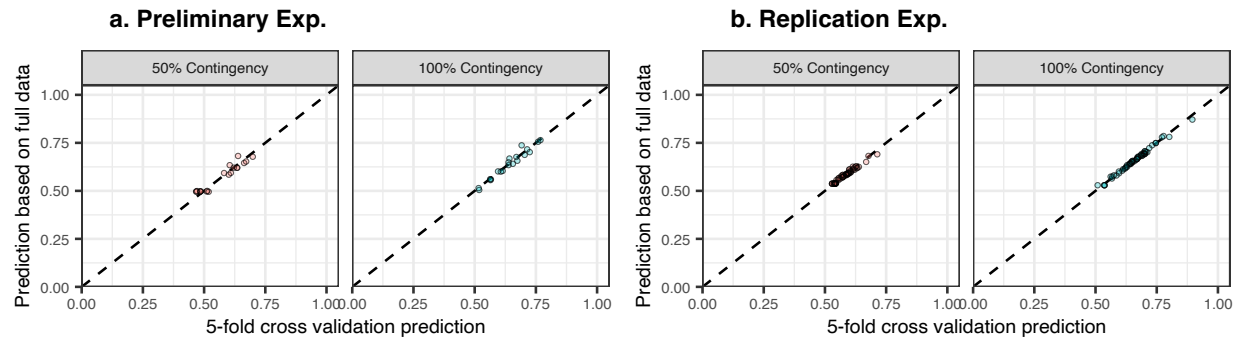
464 **Fig. S16.**

465 Actual RT distributions versus simulated posterior distributions, by  $\theta_{slope_i}$  quantile group for  
 466 Study 2. Participants were categorized into four equal quantile groups, according to their parameter  
 468 estimates (denoted here as Q1-Q4; columns). Participants with higher parameter estimates were  
 469 characterized with faster transition to anticipatory responses (top row). Posterior simulated RT  
 470 distributions using mixture of log-normal (bottom row) recreated relatively well this transition  
 471 pattern. Vertical dashed line represents cue-onset.  
 472

**Fig. S17.**

Study 2 probe results, by contingency and value-group. Proportion of trials participants chose Go stimuli over NoGo stimuli of similar initial value, in the preliminary experiment (a) and replication experiment (b). Dots represent individual participants, error-bars represents 95% CI based on a mixed model logistic regression. In square brackets are the value index of the stimuli in the group (lower index indicates higher value). No monotonic trend was observed for value group main effect or interaction with contingency conditions. Statistical significance is denoted with asterisks (\*  $p < 0.05$ , \*\*  $p < 0.01$ , \*\*\*  $p < 0.001$ ; one-sided mixed model logistic regression). Dashed line represents 50% chance level.

**Fig. S18.**



Study 2 five-fold cross validation results. The predicted proportion of trials participants chose Go stimuli over NoGo stimuli based on the full model using the entire dataset (y-axis), plotted against prediction made by CV model (x-axis) which used the data of 80% of the participants to predict the proportion of the held back 20% participant. Results presented for the preliminary experiment (a) and replication experiment (b). Dots represent individual participants; dashed line indicates the identity line. The CV model resulted in nearly identical prediction as the full model.

## Computational model of non-reinforced learning

### Code S1.

```
493 // Stan model - Meta analysis (28 experiments). Effective RT: 2 Gaussians mixture model
494 data {
495   // Input: data shape
496   int N_subjects; // N = 828 valid participants
497   int N_trials; // maximal number of trials
498   int N_trials_valid[N_subjects]; // actual number of trials of each participant
499   // Input: Dependent (RT-effective) and independent (Run) data
500   real RT[N_trials, N_subjects]; // dependent variable: effective RT (RT - Cue)
501   real Run[N_trials, N_subjects]; // training run [0,1] indicating (1st, 20th run)
502   real theta_b0; // b0 = -3.1, mixture proportion defined at 1st run (set to 0.1%)
503 }
504 parameters {
505   // Group-level parameters of the two Gaussians
506   real <upper=150> mu_fix_1; // early Gaussian Mean (upper limit at cue-onset)
507   real <lower=200> mu_fix_2; // late Gaussian Mean (lower limit at cue-onset)
508   real <lower=0> sigma_e[2]; // SD of the two Gaussians
509   // Group-level (fixed-effect) parameters: theta-slope individualized learning parameter
510   real theta_slope_fix; // Mean of individualized theta-slope parameter
511   real <lower=0> theta_slope_sd; // SD of individualized theta-slope parameter
512   // Participant-level (random-effect) parameter: individualized learning parameter
513   Real theta_slope[N_subjects]; // parameter to determine mixture proportion
514 }
515 model {
516   // Priors
517   mu_fix_1 ~ normal(-500,500);
518   mu_fix_2 ~ normal(500,500);
519   sigma_e ~ normal(0,1000);
520   theta_slope_fix ~ normal(0,1);
521   theta_slope_sd ~ normal(0,0.7);
522   // Likelihood
523   for (i in 1:N_subjects){ // go over participants
524     real theta; // mixture proportion
525     // participant-level [random-effect] from group-level [fix-effect] parameters
526     theta_slope[i] ~ normal(theta_slope_fix,theta_slope_sd);
527     for (t in 1:N_trials_valid[i]){ // go over participant's trials
528       /* define a linear function of training run with individualized theta-slope use the
529       inverse of the Normal distribution CDF to go from [-inf,inf] to [0,1] range */
530       theta = Phi_approx(theta_b0 + theta_slope[i] * Run[t,i]); // Phi = pnorm function
531       // maximize likelihood of parameters' fit to RT data
532       target += log_mix(theta, // mixture proportion
533         normal_lpdf(RT[t,i] | mu_fix_1, sigma_e[1]), // anticipatory RT
534         normal_lpdf(RT[t,i] | mu_fix_2, sigma_e[2])); // cue-dependent RT
535     }
536   }
537 }
538 }
539 Stan model used in Study 1 meta-analysis.
```

## Computational model of non-reinforced learning

### Code S2.

```
540 // Stan model - Study 2. RT: mixture model, with stimulus-level Learning parameters
541 data {
542   // Input: data shape
543   int N_subjects; // Number of participants
544   int N_trials; // maximal number of Go trials
545   int N_stim; // number of unique Go stimuli-participants combinations
546   int N_trials_valid[N_subjects]; // actual number of trials of each participant
547   // Input: Dependent (RT-effective) and independent (Run and contingency) data
548   real RT[N_trials, N_subjects]; // dependent variable: effective RT (RT - Cue)
549   real Run[N_trials, N_subjects]; // training run [0,1] indicating [1st, 20th]
550   int Stim[N_trials, N_subjects]; // stimulus indicator
551   real theta_b0; // b0 = -3.1, mixture proportion defined at 1st run (set to 0.1%)
552 }
553 parameters {
554   // Group-level parameters of the two Gaussians
555   real <upper=150> mu1_fix; // early Gaussian Mean (upper limit: 150ms after cue)
556   real <lower=200> mu2_fix; // late Gaussian Mean (Lower limit: 200ms after cue)
557   real <lower=0> sigma_e[2]; // SD of the two Gaussians
558   // Group-level (fixed-effect) parameters: theta-slope individualized learning parameter
559   real theta_slope_fix; // individualized theta-slope parameter Mean
560   real <lower=0> theta_stim_sd; // SD of stimulus-level theta-slope parameter
561   // Stimulus-level (random-effect) parameters: individualized learning parameters
562   real theta_slope_stim [N_stim]; // stimuli individualized parameter
563 }
564 model {
565   //priors
566   mu1_fix ~ normal(500,500);
567   mu2_fix ~ normal(1000,500);
568   sigma_e ~ normal(0,1000);
569   theta_slope_fix ~ normal(0,0.7);
570   theta_stim_sd ~ cauchy(0,1);
571   //Likelihood
572   // Stimulus-level [random-effect] from group-level [fix-effect] parameters
573   theta_slope_stim ~ normal(theta_slope_fix, theta_stim_sd);
574   for (i in 1:N_subjects){ // go over participants
575     for (t in 1:N_trials_valid[i]){ // go over participant's trials
576       // identify stimulus index
577       int stim_i;
578       stim_i = Stim[t,i];
579       /* define a linear function of training run with individualized theta-slope use the
580       inverse of the Normal distribution CDF to go from [-inf,inf] to [0,1] range */
581       // maximize likelihood of parameters' fit to RT data
582       target += log_mix(Phi_approx(
583         theta_b0 + theta_slope_stim[stim_i] * Run[t,i]), // mixture proportion
584         normal_lpdf(RT[t,i] | mu1_fix, sigma_e[1]), // anticipatory RT
585         normal_lpdf(RT[t,i] | mu2_fix, sigma_e[2])); // cue-dependent RT
586     }
587   }
588 }
589 }
```

590 Stan model used in Study 1 with stimulus-level learning parameter.

## Computational model of non-reinforced learning

### Code S3.

```
591 // Stan model - Study 2. RT: 2 Gaussians mixture model, with 2 contingency conditions
592 data {
593   // Input: data shape
594   int N_subjects; // n=20 or n = 59
595   int N_trials; // maximal number of Go trials
596   int N_trials_valid[N_subjects]; // actual number of trials of each participant
597   // Input: Dependent (RT-effective) and independent (Run and contingency) data
598   real RT[N_trials, N_subjects]; // dependent variable: RT
599   real Run[N_trials, N_subjects]; // training run [0,1] indicating [1st, 20th]
600   real Contingency[N_trials, N_subjects]; // 50% (0) or 100% (1) contingency condition
601   real theta_b0; // b0 = -3.1, mixture proportion defined at 1st run (set to 0.1%)
602   real Cue; // fixed at 850ms
603 }
604 parameters {
605   // Group-Level parameters of the two Gaussians
606   real <lower=0, upper=Cue-100> mu1_fix; // early Gaussian Mean (limit: 100ms before cue)
607   real <lower=Cue> mu2_fix; // late Gaussian Mean (Lower limit at cue-onset)
608   real <lower=0> sigma_e[2]; // SD of the two Gaussians
609   real <lower=0> mu1_sd[1]; // SD of individualized early Gaussian mean parameter
610   // Group-Level (fixed-effect) parameters: theta-slope individualized learning parameter
611   real theta_fix[2]; // individualized theta-slope parameter Means (2 conditions)
612   real <lower=0> theta_sd[2]; // SDs of individualized theta-slope parameter
613   // Participant-Level (random-effect) parameters: individualized learning and early mean
614   real <upper=Cue+100> mu1_rand[N_subjects]; // individualized mean for early RT
615   real theta_slope_100[N_subjects]; // individualized learning-parameter - 100% contingency
616   real theta_slope_50[N_subjects]; // individualized learning-parameter - 50% contingency
617 }
618 model {
619   //priors
620   mu1_fix ~ normal(500,500);
621   mu2_fix ~ normal(1000,500);
622   sigma_e ~ normal(0,1000);
623   mu1_sd ~ normal(0,150);
624   theta_fix ~ normal(0,0.7);
625   theta_sd ~ normal(0.7,0.7);
626   //Likelihood
627   for (i in 1:N_subjects){ // go over participants
628     // participant-Level [random-effect] from group-Level [fix-effect] parameters
629     mu1_rand[i] ~ normal(mu1_fix,mu1_sd); // individualized early RT mean
630     theta_slope_100[i] ~ normal(theta_fix[1],theta_sd[1]); // 100% contingency
631     theta_slope_50[i] ~ normal(theta_fix[2],theta_sd[2]); // 50% contingency
632     for (t in 1:N_trials_valid[i]){ // go over participant's trials
633       /* define a linear function of training run with individualized theta-slope use the
634       inverse of the Normal distribution CDF to go from [-inf,inf] to [0,1] range */
635       // maximize likelihood of parameters' fit to RT data
636       target += Log_mix(Phi_approx(
637         theta_b0 + (Contingency[t,i]*theta_slope_100[i] +
638         (1-Contingency[t,i])*theta_slope_50[i]) * Run[t,i]), // mixture proportion
639         normal_Lpdf(RT[t,i] | mu1_rand[i], sigma_e[1]), // anticipatory RT
640         normal_Lpdf(RT[t,i] | mu2_fix, sigma_e[2])); // cue-dependent RT
641     }
642   }
643 }
644 }
645 Stan model used in Study 2 main-analysis.
```

## Computational model of non-reinforced learning

### Code S4.

```

646 // Stan model - Study 2. RT: mixture model, with stimulus-level Learning parameters (each
647 contingency separately)
648 data {
649   // Input: data shape
650   int N_subjects; // n=20 or n = 59
651   int N_trials; // maximal number of Go trials
652   int N_stim; // number of Go stimuli (=16)
653   int N_trials_valid[N_subjects]; // actual number of trials of each participant
654   // Input: Dependent (RT-effective) and independent (Run and contingency) data
655   real RT[N_trials, N_subjects]; // dependent variable: RT
656   real Run[N_trials, N_subjects]; // training run [0,1] indicating [1st, 20th]
657   int Stim[N_trials, N_subjects]; // stimulus indicator
658   real theta_b0; // b0 = -3.1, mixture proportion defined at 1st run (set to 0.1%)
659   real Cue; // fixed at 850ms
660 }
661 parameters {
662   // Group-Level parameters of the two Gaussians
663   real <lower=0, upper=Cue+100> mu1_fix; // early Gaussian Mean (Limit: 100ms after cue)
664   real <lower=Cue> mu2_fix; // Late Gaussian Mean (lower limit at cue-onset)
665   real <lower=0> sigma_e[2]; // SD of the two Gaussians
666   // Group-Level (fixed-effect) parameters: theta-slope individualized learning parameter
667   real theta_fix; // individualized theta-slope parameter Means
668   real <lower=0> theta_sub_sd; // SD of participant-level theta-slope parameter
669   real <lower=0> theta_stim_sd; // SD of stimulus-level theta-slope parameter
670   // Participant-Level (random-effect) parameters: individualized learning parameters
671   real theta_slope_sub [N_subjects]; // participants individualized parameter
672   real theta_slope_stim [N_subjects, N_stim]; // stimuli individualized parameter
673 }
674 model {
675   //priors
676   mu1_fix ~ normal(500,500);
677   mu2_fix ~ normal(1000,500);
678   sigma_e ~ normal(0,1000);
679   theta_fix ~ normal(0,0.7);
680   theta_sub_sd ~ cauchy(0,1);
681   theta_stim_sd ~ cauchy(0,1);
682   //likelihood
683   // participant-level [random-effect] from group-level [fix-effect] parameters
684   theta_slope_sub ~ normal(theta_fix, theta_sub_sd);
685   for (i in 1:N_subjects){ // go over participants
686     // stimulus-level parameter model variability around participant-level parameter
687     theta_slope_stim[i,] ~ normal(theta_slope_sub[i], theta_stim_sd);
688     for (t in 1:N_trials_valid[i]){ // go over participant's trials
689       // identify stimulus index
690       int stim_i;
691       stim_i = Stim[t,i];
692       /* define a linear function of training run with individualized theta-slope use the
693       inverse of the Normal distribution CDF to go from [-inf,inf] to [0,1] range */
694       // maximize likelihood of parameters' fit to RT data
695       target += Log_mix(Phi_approx(
696         theta_b0 + theta_slope_stim[i, stim_i] * Run[t,i]), // mixture proportion
697         normal_Lpdf(RT[t,i] | mu1_fix, sigma_e[1]), // anticipatory RT
698         normal_Lpdf(RT[t,i] | mu2_fix, sigma_e[2])); // cue-dependent RT
699     }
700   }
701 }
702 }

```

Stan model used in Study 2 with stimulus-level learning parameter. The model was evaluated separately for each condition (100% contingency and 50% contingency).

## Computational model of non-reinforced learning

### Code S5.

```
// Stan model - Meta analysis (28 experiments). Effective RT: 2 Log-normal mixture model
data {
  // Input: data shape
  int N_subjects; // N = 828 valid participants
  int N_trials; // maximal number of trials
  int N_trials_valid[N_subjects]; // actual number of trials of each participant
  // Input: Dependent (RT-effective) and independent (Run) data
  real RT[N_trials, N_subjects]; // dependent variable: effective RT (RT - Cue)
  real Run[N_trials, N_subjects]; // training run [0,1] indicating (1st, 20th run)
  real theta_b0; // b0 = -3.1, mixture proportion defined at 1st run (set to 0.1%)
}
parameters {
  // Group-Level parameters of the two Gaussians
  real <upper=0.14> mu_fix_1; // early Gaussian Mean (upper limit at cue-onset)
  real <lower=0.18> mu_fix_2; // Late Gaussian Mean (Lower limit at cue-onset)
  real <lower=0> sigma_e[2]; // SD of the two Gaussians
  // Group-Level (fixed-effect) parameters: theta-slope individualized learning parameter
  real theta_fix; // Mean of individualized theta-slope parameter
  real <lower=0> theta_sd; // SD of individualized theta-slope parameter
  // Participant-Level (random-effect) parameter: individualized learning parameter
  Real theta_b1[N_subjects]; // parameter to determine mixture proportion
}
model {
  // Priors
  mu_fix_1 ~ normal(-0.693,1);
  mu_fix_2 ~ normal(0.4,0.4);
  sigma_e ~ normal(0,1);
  theta_fix ~ normal(0,1);
  theta_sd ~ normal(0,0.7);
  // Likelihood
  for (s in 1:N_subjects){ // go over participants
    real theta; // mixture proportion
    // participant-level [random-effect] from group-level [fix-effect] parameters
    theta_b1[s] ~ normal(theta_fix,theta_sd);
    for (t in 1:N_trials_valid[s]){ // go over participant's trials
      /* define a linear function of training run with individualized theta-slope use the
      inverse of the Normal distribution CDF to go from [-inf,inf] to [0,1] range */
      theta = Phi_approx(theta_b0 + theta_b1[s] * Run[t,s]); // Phi = pnorm function
      // maximize likelihood of parameters' fit to RT data
      target += log_mix(theta, // mixture proportion
        Lognormal_Lpdf(RT[t,s] | mu_fix_1, sigma_e[1]), // anticipatory RT
        Lognormal_Lpdf(RT[t,s] | mu_fix_2, sigma_e[2])); // cue-dependent RT
    }
  }
}
```

Alternative Stan model used in Study 1 using log-normal distributions.

754 **Table S1.** Results - Unpublished experiments

| Experiment         | <i>n</i> | Training runs | Prop. Go stimuli were chosen | <i>p</i> <sup>1</sup> | OR (95% CI)       |
|--------------------|----------|---------------|------------------------------|-----------------------|-------------------|
| Exp. 10            | 25       | 20            | High value 58.4%             | 0.009                 | 1.50 [1.07, 2.09] |
| Familiar Faces I   |          |               | Low value 58.4%              | 0.014                 | 1.57 [1.05, 2.35] |
| Exp. 11            | 39       | 20            | High value 55.6%             | 0.021                 | 1.30 [1.01, 1.67] |
| Familiar Faces II  |          |               | Low value 55.5%              | 0.025                 | 1.32 [1.00, 1.73] |
| Exp. 13            | 42       | 20            | Pos. affect 60.0%            | 3.1e-04               | 1.64 [1.23, 2.17] |
| Emotional Faces I  |          |               | Neut. affect 58.7%           | 7.0e-04               | 1.53 [1.18, 1.98] |
| Exp. 14            | 70       | 20            | Pos. affect 59.1%            | 6.3e-06               | 1.51 [1.26, 1.82] |
| Emotional Faces II |          |               | Neut. affect 58.4%           | 1.9e-04               | 1.52 [1.21, 1.91] |
| Exp. 16            | 26       | 12            | High value 50.1%             | 0.494                 | 1.00 [0.73, 1.38] |
| Snacks Adolesce.   |          |               | Low value 56.8%              | 0.010                 | 1.36 [1.05, 1.77] |
| Exp. 19            | 23       | 12            | High value 53.8%             | 0.080                 | 1.18 [0.94, 1.49] |
| Snacks pilot I     |          |               | Low value 53.5%              | 0.126                 | 1.18 [0.89, 1.55] |
| Exp. 20            | 25       | 20            | High value 59.7%             | 0.005                 | 1.61 [1.12, 2.32] |
| Snacks pilot II    |          |               | Low value 63.8%              | 3.7e-04               | 1.95 [1.32, 2.88] |

755 *Notes:*756 <sup>1</sup> One-sided logistic regression.

## SI References

- Bakkour, A., Leuker, C., Hover, A. M., Giles, N., Poldrack, R. A., & Schonberg, T. (2016). Mechanisms of choice behavior shift using cue-approach training. *Frontiers in Psychology*, 7, Article 421. <https://doi.org/10.3389/fpsyg.2016.00421>
- Bakkour, A., Lewis-Peacock, J. A., Poldrack, R. A., & Schonberg, T. (2017). Neural mechanisms of cue-approach training. *NeuroImage*, 151, 92–104. <https://doi.org/10.1016/j.neuroimage.2016.09.059>
- Botvinik-Nezer, R., Bakkour, A., Salomon, T., Shohamy, D., & Schonberg, T. (2021). Memory for individual items is related to nonreinforced preference change. *Learning & Memory*, 28(10), 348–360. <https://doi.org/10.1101/LM.053411.121>
- Botvinik-Nezer, R., Salomon, T., & Schonberg, T. (2020). Enhanced bottom-up and reduced top-down neural mechanisms drive long-lasting non-reinforced behavioral change. *Cerebral Cortex*, 30(3), 858–874. <https://doi.org/10.1093/cercor/bhz132>
- Lundqvist, D., Flykt, A., & Öhman, A. (1998). Karolinska directed emotional faces. *Cognition and Emotion*.
- Salomon, T., Botvinik-Nezer, R., Gutentag, T., Gera, R., Iwanir, R., Tamir, M., & Schonberg, T. (2018). The cue-approach task as a general mechanism for long term non-reinforced behavioral change. *Scientific Reports*, 8, article number: 3614, 1-13. <https://doi.org/10.1038/s41598-018-21774-3>
- Salomon, T., Botvinik-Nezer, R., Oren, S., & Schonberg, T. (2019). Enhanced striatal and prefrontal activity is associated with individual differences in nonreinforced preference change for faces. *Human Brain Mapping*, 41(4), 1043–1060. <https://doi.org/10.1002/hbm.24859>
- Schonberg, T., Bakkour, A., Hover, A. M., Mumford, J. A., Nagar, L., Perez, J., & Poldrack, R. A. (2014). Changing value through cued approach: an automatic mechanism of behavior change. *Nature Neuroscience*, 17(4), 625–630. <https://doi.org/10.1038/nn.3673>