



The contribution of striatal pseudo-reward prediction errors to value-based decision-making

Ernest Mas-Herrero^{a,*}, Guillaume Sescousse^{b,c,1}, Roshan Cools^b, Josep Marco-Pallarés^{d,e,**}

^a Montreal Neurological Institute, McGill University, Montreal, QC, H3A 2B4, Canada

^b Donders Institute for Brain, Cognition and Behaviour, Radboud University, Nijmegen, the Netherlands

^c Lyon Neuroscience Research Center - INSERM U1028 - CNRS UMR5292, PSYR2 Team, University of Lyon, Lyon, France

^d Cognition and Brain Plasticity Group [Bellvitge Biomedical Research Institute – IDIBELL], L'Hospitalet de Llobregat, Barcelona, Spain

^e Department of Cognition, Development and Educational Psychology, Institute of Neurosciences, University of Barcelona, Barcelona, Spain

ABSTRACT

Most studies that have investigated the brain mechanisms underlying learning have focused on the ability to learn simple stimulus-response associations. However, in everyday life, outcomes are often obtained through complex behavioral patterns involving a series of actions. Parallel learning systems might be important to reduce the complexity of the learning problem in such scenarios, as proposed in the framework of hierarchical reinforcement learning (HRL). The key feature of HRL is the decomposition of complex sets of action into subgoals. These subgoals are associated with the computation of pseudo-reward prediction errors (PRPEs), which allow the reinforcement of actions that led to a subgoal before the final goal itself is achieved. Here we wanted to test the hypothesis that, despite not carrying any rewarding value *per se*, pseudo-rewards might generate a bias in choice behavior in the absence of any advantage. Second, we also hypothesized that this bias might be related to the strength of PRPE striatal representations. In order to test these ideas, we developed a novel decision-making paradigm to assess reward prediction errors (RPEs) and PRPEs in two studies (fMRI study: $n = 20$; behavioral study: $n = 19$). Our results show that the participants developed a preference for the most pseudo-rewarding option throughout the task, even though it did not lead to more monetary rewards. fMRI analyses revealed that this preference was predicted by individual differences in the relative striatal sensitivity to PRPEs vs RPEs. Together, our results indicate that pseudo-rewards generate learning signals in the striatum and subsequently bias choice behavior despite their lack of association with actual reward.

1. Introduction

Reinforcement Learning (RL) theories have provided invaluable insights into reward-guided learning and decision-making. A core feature of most RL theories is that actions are reinforced according to reward prediction errors (RPEs), that is, according to whether obtained outcomes are better or worse than expected (Sutton and Barto, 1998). It is now widely accepted that dopaminergic neurons projecting to the striatum play a pivotal role in the computation of RPEs (Schultz et al., 1997).

Although the computational principles of standard RL theories have accounted for a wide range of experimental findings in simple learning paradigms (Niv and Schoenbaum, 2008; Lee et al., 2012; Mas-Herrero and Marco-Pallarés, 2016), they provide limited efficacy in complex environments and situations involving a series of actions (Botvinick et al., 2009; Botvinick, 2012; Vilà-Balló et al., 2017). For instance, when attending a conference in a new town, a series of actions is required to get from the hotel to the conference venue. If you get lost on the way, you

will need to evaluate which of these actions need(s) to be adjusted next time. As the number of required actions increases, this operation becomes increasingly difficult for standard RL algorithms. One of the computational approaches that has emerged to solve this *credit assignment* problem is hierarchical reinforcement learning (HRL) (Botvinick et al., 2009; Botvinick, 2012; Hengst, 2012).

In HRL, sequences of actions (e.g. exit hotel, turn right, go straight and turn left) are packaged together into subroutines evaluated according to their own subgoals (e.g. reach bus stop). Attaining a subgoal yields a so-called *pseudo-reward*, and any deviation from subgoal expectations generates a pseudo-reward prediction error (PRPE). Thus, in HRL there are two value functions – tracking rewards and pseudo-rewards – that are learnt in parallel. HRL reduces the complexity of the learning problem by (1) focusing on a small set of decisions (the number of subroutines) rather than a large sequence of primitive actions and (2) confining behavioral adjustment to the actions within one subroutine when subgoal expectations are not met (Fig. 1). Notably, PRPEs reinforce actions leading to the

* Corresponding author.

** Corresponding author. Department of Basic Psychology – IDIBELL, L'Hospitalet de Llobregat, University of Barcelona, Campus Bellvitge, Barcelona, 08097, Spain. E-mail addresses: ernesto.masherrero@mcgill.ca (E. Mas-Herrero), josepmarco@ub.edu (J. Marco-Pallarés).

¹ Both authors equally contributed to this work.

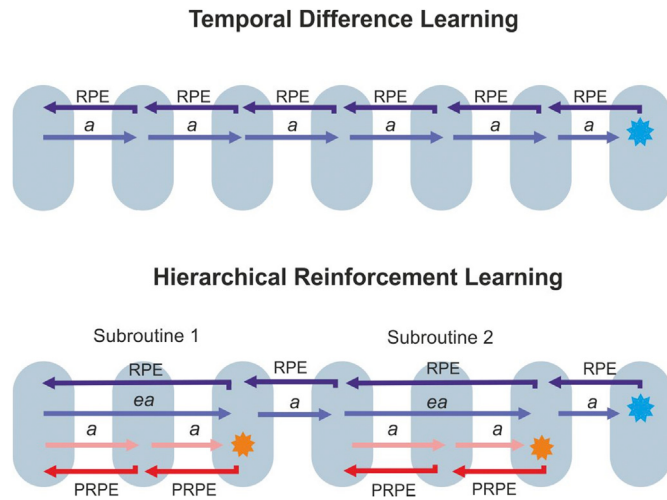


Fig. 1. Main differences between classical Reinforcement Learning, as implemented in flat Reinforcement Learning (FRL) models, and Hierarchical Reinforcement Learning (HRL). Flat RL models aim to predict all events occurring in between an initial primitive action (*a*) and the final reward (blue star). All these events generate RPEs which are back-propagated in time. In order to accelerate learning, HRL algorithms divide primitive actions into subroutines of extended actions (*ea*) leading to subgoals/pseudo-rewards (orange star). RPEs are computed following the attainment of each subgoal/pseudo-reward. In doing this, RPE signals are back-propagated more rapidly (in the above example, 6 jumps are needed in FRL to reach the initial action from the final outcome, while only 4 jumps are required in HRL). Importantly, HRL includes a second value function tracking the attainment of subgoals/pseudo-rewards. This value function uses PRPEs to reinforce those actions leading to a certain subgoal. The goal of the present study is to examine the parallel neural computation of RPEs and PRPEs and how they influence choice behavior (adapted from Botvinick et al., 2009).

subgoal even before the final reward is obtained, therefore speeding up learning. Note however that the benefits of HRL over flat RL mainly occur when subgoals are pre-learned and are indeed appropriate for the accomplishment of the goal (Botvinick et al., 2009; Van Dijk et al., 2011). Importantly, while the concept of HRL resonates with the ubiquity of hierarchy in human behavior, its cognitive and neural nature is yet to be determined. The first fMRI studies that tackled this question showed that both RPEs and PRPEs are computed in the ventral striatum (Ribas-Fernandes et al., 2011; Diuk et al., 2013; although note that a recent study by Ribas-Fernandes et al., 2019, could not replicate their earlier finding).

However, the reinforcing properties of pseudo-rewards may also potentially lead to a preference for non-beneficial pseudo-rewarding options. Going back to the previous example, you might prefer to take a bus in front of your hotel followed by a 10 min walk to the conference venue, rather than a bus that would first require a 10 min walk but would then drop you off even more rapidly at the conference venue. This may be because in the first option, the first subgoal (taking the bus) is attained more rapidly and thus associated with more pseudo-reward.

Here, we aimed to study whether individuals may develop such biases towards pseudo-rewards and whether these depend on striatal sensitivity to RPEs and PRPEs. We developed a novel fMRI task in which participants had to collect as much money as possible from locked boxes. In order to collect the money, participants first had to unlock these boxes (subgoal). Once unlocked (pseudo-reward), each box could contain money (reward) or not. Crucially, participants had to choose between two keys to unlock the boxes. These two keys eventually led to similar amounts of money but, by design, one of them unlocked more boxes than the other. We hypothesized that participants would be biased towards the former, associated with more pseudo-reward, even though this decision was not related to higher monetary gains. In addition, if the

striatum computes both RPEs and PRPEs, we hypothesized that the relative striatal sensitivity to RPEs vs PRPEs would predict participants' bias.

2. Materials and methods

2.1. Participants

Twenty-three students from the University of Barcelona ($M = 21.9$ years, $SD = 2.7$, 10 males) participated in the fMRI experiment. Three participants had to be excluded due to excessive head motion (see *fMRI data analysis* section). All participants were paid 20€ per hour and a monetary bonus depending on their performance. All participants gave written informed consent, and all procedures were approved by the University of Barcelona ethics committee.

2.2. Experimental procedure

We developed a novel learning task inducing both reward and pseudo-reward prediction errors (Fig. 2). Participants were encouraged to accumulate as much money as possible. The money was placed into two boxes which were locked. Each box could only be unlocked with its own key. Thus, in each trial, the participants had to choose between two keys, in order to unlock the padlock and find out whether there was money in the box. Unlocking the box can be defined as a pseudo-reward as it is a non-rewarding state that participants must go through to reach the reward. Crucially, one box could be unlocked 70% of the time but only contained money 30% of the time, while the other box could only be unlocked 30% of the time but contained money 70% of the time. Thus, the two boxes were associated with the same final reward probability ($70\% \times 30\% = 21\%$), but one of them was associated with a higher pseudo-reward rate (70% vs 30%). With this design we aimed to study whether pseudo-rewards may bias participants' decision-making towards a preference for the most pseudo-rewarding option, despite the fact that this option did not confer any advantage in terms of monetary reward (Fig. 2).

Participants had to select, in less than 2 s, one of the two keys by pressing either the left (index finger) or right button (middle finger) of a response pad. If participants did not respond in time, a question mark appeared in the center of the screen for 1000 ms (these trials were discarded from the analyses). After a delay (1200 ms), a 'pseudo-feedback' indicated whether the padlock was unlocked or not (1000 ms). Then, 500 ms after the pseudo-feedback offset, the padlock was presented either on the left or the right side of the screen. Participants were told that by correctly indicating its position (left or right button) they would be 'removing the key' from the padlock. Participants were asked to perform this action as fast as possible (within a time limit of 1000 ms), regardless of whether it had been unlocked or not. This was included to ensure that participants were paying attention throughout the trial. If they did not respond within the time limit, the final outcome was always negative. If they were successful, a 'reward feedback' was presented after a delay of 1200 ms, indicating whether the box contained money or not (green tick = 0.25€, red cross = 0€, 1000 ms). In trials in which the key was not properly removed or the padlock remained unlocked, a red cross indicating that the participant did not accumulate money was also presented in order to maintain the same structure in all trials (this 'non-informative' feedback was not further analyzed). After reward feedback a fixation cross remained on the screen until the end of the trial (mean trial duration was set to 11 s) followed by variable inter-trial interval (randomly jittered between 100 and 2500 ms, with 300 ms increments). Thus, participants were rewarded only if the following requirements were met: the box was unlocked, the key was properly removed in time and the box contained money. In other cases, the reward feedback was always a non-rewarding outcome. As mentioned previously, one of the keys (the most pseudo-rewarding key, referred to as Key+) had a greater probability of unlocking its associated boxes ($p = .7$) but those boxes had

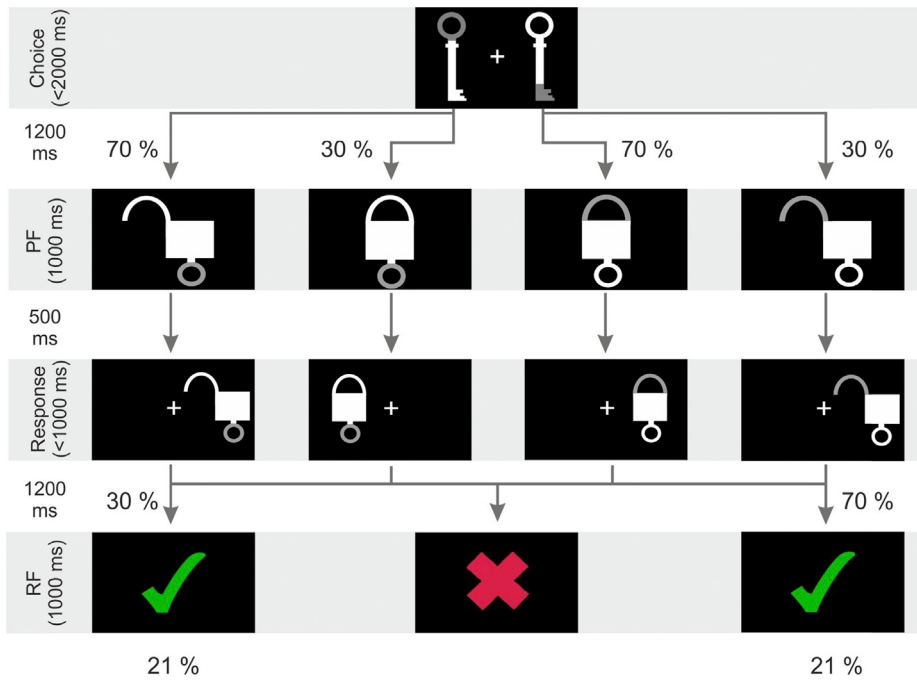


Fig. 2. Task design. In the first-stage participants have to choose between two different keys which, unbeknownst to them, lead to equal amounts of reward. Each key unlocks a different type of box. In the second-stage a locked or unlocked padlock is presented depending on whether participants successfully unlocked the box or not. After using the key, and before the outcome is presented, participants have to remove the key from the padlock by indicating the position of the lock (third stage). Finally, either a green tick (reward) or a red cross (no-reward) is presented. The key on the left unlocks its associated boxes 70% of the time but these boxes contain money only 30% of the time. In contrast, the key on the right unlocks its associated boxes only 30% of the time but these boxes contain money 70% of the time. Overall, both keys are thus equally rewarded (21% of trials). RF = Reward Feedback; PF = Pseudo-Reward feedback.

a low probability of containing money ($p = .3$). The other key (referred to as Key-) presented the opposite pattern: it unlocked a smaller fraction of boxes ($p = .3$) but these boxes were more likely to contain money ($p = .7$). As a consequence both keys were rewarded with the same overall probability ($p = .21$). Reward and pseudo-reward outcomes were probabilistically determined on each trial. Participants were told that a monetary bonus proportional to the amount of money accumulated during the task would be paid out at the end of experiment.

The task consisted of three blocks of 69 trials. Each block comprised of 46 forced-choice trials and 23 free-choice trials that were intermixed. Forced-choice trials were similar to the free-choice trials explained above with the difference that only one key appeared in the screen, either in the left or right side (random order). Participants had to press the corresponding button to select it. In half of the forced-choice trials, participants were forced to select the most pseudo-rewarding key. In the other half, they had to select the least pseudo-rewarding key. These forced-choice trials were introduced to ensure that participants would sample both options equally. In the remaining free-choice trials, participants could freely choose between the two different keys.

2.3. Behavioral analysis

In order to study the impact of pseudo-rewards on participants' choices we used two behavioral measures. First, we quantified the preference for the Key+, by computing the proportion of free-choices in which participants selected the Key+ over the Key-. To illustrate the effect of learning, we divided free-choice trials into 11 bins of 6 trials each, and plotted the preference for the Key+ across bins. A repeated-measures ANOVA with time (i.e. bins) as a within-participant factor was performed to assess whether individuals' preferences remained stable through the task or changed with learning. In addition, we also performed a binomial test for each individual to test whether the proportion of Key+ choices was significantly different from chance level ($p = .5$) at the individual level.

Second, we performed a modified logistic regression to predict participants' choices based on both the amount of reward and pseudo-reward obtained in previous trials (following a similar procedure as in Diuk et al., 2013). Specifically, we regressed participants' free choices (KEY+/KEY-) onto the linear combination of rewards and pseudo-rewards obtained in

the last four trials in which KEY+/KEY- choices were made (including both free and forced-choice trials). For each trial, we first computed the value of each key as follows:

$$V_t = \beta_{rewards} \sum_{j=1}^4 R_j + \beta_{pseudo} \sum_{j=1}^4 Pr_j. \quad (1)$$

Where R_j was +1 or 0 depending on whether the participant received money or not on the j th-to-last time (s)he selected that key, and Pr_j was +1 or 0 depending on whether the participant was pseudo-rewarded (unlocked the box) or not the j th-to-last time (s)he selected that key. Participants' free-choices were then logistically regressed on these values using a softmax equation:

$$p(Key+) = \frac{e^{v(Key+)}}{e^{v(Key+)} + e^{v(Key-)}}. \quad (2)$$

We optimized the two β parameters by minimizing the negative log likelihood using the *fminunc* function in MATLAB. Then, we tested whether the estimated parameters were significantly different from 0 at the group level using one-sample t-tests. The two β parameters represent the respective weights of past rewards and pseudo-rewards in participant's choices. If participants are guiding their behavior based on the availability of both rewards and pseudo-rewards we would expect both β s to be positive and significantly different from 0.

2.4. Reinforcement learning model

In order to estimate RPEs and PRPEs for the fMRI analysis, we implemented a temporal difference learning (TD) model in which two actions values are computed in parallel based on the history of rewards (V_r) and pseudo-rewards (V_{pr}), respectively (Diuk et al., 2013). V_r is calculated according to reward prediction error. In particular, two RPEs are computed in the current task: when the pseudo-feedback is presented (RPE1) and when the final outcome is reached (RPE2). In contrast, V_{pr} is calculated according to PRPEs. PRPEs only arise when the pseudo-feedback is presented and indicates whether the subgoal has been accomplished or not. Thus, at the time of the pseudo-feedback two distinct prediction errors are computed, the RPE1 and the PRPE, related to the attainment of rewards and pseudo-rewards, respectively.

RPE. The action value Vr of each key (in the example below, the Key+) was updated when the pseudo-feedback was presented:

$$Vr(\text{KEY+})_{t+1} = Vr(\text{KEY+})_t + \alpha \text{RPE1}. \quad (3)$$

Where α is the learning rate, the subscript t represent number of trial, and RPE1 is the RPE term generated following the first action:

$$\text{RPE1} = Vr(\text{state}^i)_t - Vr(\text{KEY+})_t. \quad (4)$$

That is, RPE1 is generated by comparing the value of the state presented just after the selection of a key ($Vr(\text{state}^i)_t$) and the action value of that key ($Vr(\text{KEY+})_t$) (note that for simplicity we omitted the reward term r which is equal to 0 for RPE1). In the current task four different states may be presented after the key selection: KEY + box unlocked, KEY + box locked, KEY- box unlocked, KEY- box locked (represented by j superscript). These states acquire value after the final feedback is presented through a classic TD rule:

$$Vr(\text{state}^i)_{t+1} = Vr(\text{state}^i)_t + \alpha \text{RPE2} \quad (5)$$

Where RPE2 is the prediction error term computed after the second action:

$$\text{RPE2} = r - Vr(\text{state}^i)_t. \quad (6)$$

Where r was +1 or 0 depending whether the box contained money or not.

PRPE. PRPE were computed after each pseudo-feedback as:

$$\text{PRPE} = Pr - Vpr(\text{KEY+})_t. \quad (7)$$

Where Pr was +1 or 0 depending on whether the box was unlocked or not on that trial, and $Vpr(\text{KEY+})_{t-1}$ is the expectancy of obtaining a pseudo-reward from the Key+. Vpr represents action values according to the attainment of pseudo-rewards only. These action values are learned following the same rule than Vr :

$$Vpr(\text{KEY+})_{t+1} = Vpr(\text{KEY+})_t + \alpha \text{PRPE} \quad (8)$$

The learning rate (α) was set to 0.5 (Seymour et al., 2004; Wilson and Niv, 2015).

2.5. fMRI data acquisition

fMRI data was collected using a 3T whole-body MRI scanner (General Electric MR750 GEM E). Conventional high-resolution structural images (MP RAGE sequence, repetition time (TR) = 4.7 ms, echo time (TE) = 4.8 ms, inversion time 450 ms, flip angle 12°, 1 mm isotropic voxels) were followed by functional images sensitive to blood oxygenation level-dependent (BOLD) contrast (echo planar T2*-weighted gradient echo sequence, TR = 2500 ms, TE = 29 ms, flip angle 90°). Three functional runs were acquired, each consisting of 316 whole-brain volumes (40 slices, 3.1 mm in-plane resolution, 3.1 mm thickness, no gap, positioned to cover all but the most superior region of the brain and the cerebellum). In order to reduce susceptibility artifacts in the orbitofrontal cortex and the anterior parts of the ventral striatum, slices were orientated with an angle of 30° with the plane intersecting the anterior and the posterior commissures (Weiskopf et al., 2006).

2.6. fMRI data analysis

Pre-processing was carried out using Statistical Parametric Mapping software (SPM8, Wellcome Department of Imaging Neuroscience, University College, London, UK, www.fil.ion.ucl.ac.uk/spm/). Functional images were first sinc interpolated in time to correct for slice timing differences and were spatially realigned. Realigned images were then spatially smoothed with a 4 mm FWHM kernel before they were motion-adjusted using the ArtRepair toolbox (Mazaika et al., 2007). Specifically,

functional images with significant motion artifacts were identified based on scan-to-scan motion (head position change exceeding 0.5 mm or global mean BOLD signal change exceeding 1.3%) and replaced by linear interpolation between the closest non-outlier volumes. Three participants were removed from the analysis as they presented more than 20% of volumes with artifacts in one or more runs due to excessive motion (see Mazaika et al., 2007) and, thus, the final sample consisted of 20 participants. In the remaining participants, 0.72% (SD = 1.11) of volumes were corrected for excessive motion on average. Then, the bias-corrected structural image was coregistered to the mean functional image and segmented by means of the Unified Segmentation implemented in SPM8. The resulting normalization parameters were applied to all functional images. Finally, functional images were spatially smoothed with a 7 mm FWHM kernel (the two-step smoothing of 4 mm and 7 mm is roughly equivalent to an overall smoothing of 8 mm, see Mazaika et al., 2007).

For the statistical analysis an event-related design matrix was specified. Seven regressors were included for each trial: key presentation, response, pseudo-feedback, presentation of the padlock to the left or the right, response and informative or non-informative feedback. Each response regressor was associated with a parametric regressor indicating whether the response was given with the middle or index finger. The pseudo-feedback regressor was associated with two parametric regressors modelling RPE1 and the PRPE computed at the time of the pseudo-feedback. The feedback regressor was associated with a parametric regressor modelling the reward prediction error computed at the time of the final outcome (RPE2). In order to achieve a uniform scaling of the output regression parameters, the parametric regressors representing RPE1, PRPE and RPE2 were standardized to a mean of 0 and a standard deviation of 1 (Erdeniz et al., 2013). The two parametric regressors included in the pseudo-feedback regressor were moderately correlated ($M = 0.61$). In order to avoid multicollinearity issues, the PRPE regressor was orthogonalized with respect to the RPE1. With this procedure we aimed to examine whether variations in PRPEs may still account for variations in the BOLD signal after removing the variance shared with RPE1. All regressors were subsequently convolved with the canonical hemodynamic response function and entered in a first level analysis. A high-pass filter with a cut-off of 128 s was applied to the time series. Three main contrasts of interest, testing the slopes of RPE1, PRPE and RPE2 regressors, were built at the first level. These contrast images were introduced into separate second-level group analysis based on one-sample t-tests. In order to study the conjunction of these three contrasts, first-level images were also introduced into a flexible factorial design. The conjunction analysis was formulated as a conjunction null hypothesis (Friston et al., 2005; Nichols et al., 2005) and should therefore only yield activations that were significantly present in all the contrasts introduced in the conjunction. Finally, in order to assess the relationship between participants' behavior and the BOLD signal in RPE- and PRPE-sensitive regions, we also performed a simple regression analysis using the proportion of free-choices in which participants selected the Key+ as a covariate and PRPE-RPE1 contrast. As mentioned above, we use this measure as an index of the participants' preference for the pseudo-rewarding option over the entire task. Based on our *a priori* hypothesis about the role of the ventral striatum in computing prediction errors, we restricted the correction for multiple comparisons to this region. Specifically, all results are reported at a voxel-level threshold of $p < .05$ family-wise error (FWE) corrected within a small ventral striatal volume defined *a priori* based on an anatomical mask of the bilateral Nucleus Accumbens derived from the probabilistic atlas of Hammers et al. (2003). Stereotaxic coordinates are reported in MNI space.

3. Results

3.1. Behavior

Participants presented on average a preference for the Key+

($M = 0.58$, $STD = 0.14$), which was significantly above the chance level of 0.50 [$t(19) = 2.46$, $p = .02$] (Fig. 3A). This result is particularly interesting given that both keys were strictly equivalent in terms of final reward probability. Indeed, while some individuals showed a clear preference towards the Key+, none of participants seemed to develop such a bias towards the Key- (Fig. 3A). We used a conservative binomial tests to objectivize this observation at the individual level: while 8 out of 20 participants showed a significant preference towards the most pseudo-rewarding key (i.e. a proportion of Key + choices significantly above 0.50), none of the participants showed a preference towards the least pseudo-rewarding key. This preference for the Key+ was acquired throughout the task as suggested by a clear effect of time [$F(1,19) = 8.66$, $p < .01$, Fig. 3B]. Across all participants, the preference for the Key + significantly increased from the first to the second half of the experiment [$t(19) = 2.63$, $p = .016$]. Indeed, while participants' preference was not significantly different from chance level ($p = .05$) during the first half [$M = 0.54$, $SD = 0.17$; $t(19) = 1.03$, $p = .315$], it was significantly above chance level during the second half [$M = 0.62$, $SD = 0.15$; $t(19) = 3.56$, $p = .002$].

The logistic regression analysis revealed a significant effect of both the history of rewards [$t(19) = 2.3$, $p = .03$] and pseudo-rewards, [$t(19) = 2.6$, $p = .02$] on participants' choices, suggesting that participants were influenced by both types of feedback (Fig. 3C). As expected, the beta parameter quantifying the weight of pseudo-rewards in the choice process showed a strong correlation with participants' preference: the stronger the preference for the Key+, the stronger the influence of pseudo-rewards in participants' choices [$r(18) = 0.94$, $p < .001$].

3.2. Striatal representation of prediction errors

Fig. 4 shows that brain activity in the bilateral ventral striatum correlated significantly with RPEs at the time of both the pseudo-feedback (RPE1, $p < .001$ FWE Small Volume Correction, SVC; left: $x = -10$, $y = 9$, $z = -10$; right: $x = 9$, $y = 9$, $z = -7$) and the final feedback (RPE2, $p < .001$ FWE SVC; left: $x = -10$, $y = 9$, $z = -10$; right: $x = 9$, $y = 9$, $z = -7$). Ventral striatal activity also correlated with PRPE at the time of the pseudo-feedback ($p < .01$, FWE SVC; left: $x = -9$, $y = 9$, $z = -7$; right: $x = 9$, $y = 9$, $z = -7$). In addition, the conjunction of the three contrasts (PRPE, RPE1 and RPE2) revealed significant activation in both the left and right ventral striatum ($p < .01$, FWE SVC; left: $x = -7$, $y = 9$, $z = -7$; right: $x = 9$, $y = 9$, $z = -7$). The brain T-maps illustrating the main fMRI results can be accessed at <https://neurovault.org/collections/4487/>. In a complementary analysis, we extracted the mean contrast estimates within the bilateral NAcc for each of the three main prediction errors in each participant, and then tested

whether the group average estimates were significantly different from zero using two-tailed t-tests. As expected, we found significant positive effects, confirming that brain activity in the NAcc is correlated with PRPE [$t(19) = 3.98$, $p < .001$], RPE1 [$t(19) = 4.13$, $p < .001$], and RPE2 [$t(19) = 4.67$, $p < .001$].

Since the ventral striatum is encoding both PRPEs and RPEs, we wondered whether the relative sensitivity to these two error terms would predict the behavioral preference for the Key + across participants. To assess this relative striatal sensitivity to PRPEs vs RPEs, we contrasted the corresponding parametric regressors reflecting how steeply these error terms scale with BOLD signal (PRPE-RPE1), and used the preference for the Key + as a covariate in a whole-brain simple regression analysis. We observed that the relative sensitivity to PRPEs vs RPEs was positively correlated with participants' preference specifically in the left ($p < .001$ FWE SVC; $x = -10$, $y = 9$, $z = -7$) and the right ventral striatum ($p < .05$ FWE SVC; $x = 9$, $y = 6$, $z = -7$). In other words, the higher the preference for the Key+, the greater the sensitivity to PRPEs compared with RPEs in the ventral striatum following pseudofeedback (Fig. 5). It is important to note that this relationship was not merely driven by brain-behavior correlations between participants' preferences and NAcc sensitivity to either RPEs or PRPEs alone, since no significant voxels were found when participants' preference was regressed against the parametric regressors. Complementary ROI analysis within the bilateral NAcc also showed a positive correlation between participants' performance and the relative NAcc sensitivity to PRPEs vs RPEs [$r(18) = 0.57$, $p = .009$]. Such a significant relationship was not present when using the contrast estimates from the PRPE [$r(18) = 0.39$, $p = .09$], the RPE1 [$r(18) = -0.365$, $p = .11$], or the RPE2 [$r(18) = 0.09$, $p = .71$] alone.

3.3. Control experiment

We performed an additional behavioral experiment to rule out the possibility that the observed behavioral bias was merely driven by participants' prior experience with keys and locks. Given that the sequence key-lock-open is quite common in everyday life, opening boxes may already have a pre-learned reward value that could explain the observed preference for the pseudo-rewarding key. To rule out this explanation, we tested a new group of 19 participants engaging in the same task as the fMRI experiment, but now using arbitrary cues instead of keys and locks (Fig. S1, more details in Supplementary Information). Notably, participants also developed a preference towards the most pseudo-rewarding option ($t(18) = 2.6$, $p = .018$, Fig. S2). These results replicate our previous findings using a more arbitrary design in a new group of participants, and thus confirm that participants' bias towards pseudo-rewards is not driven by pre-learned reward values.

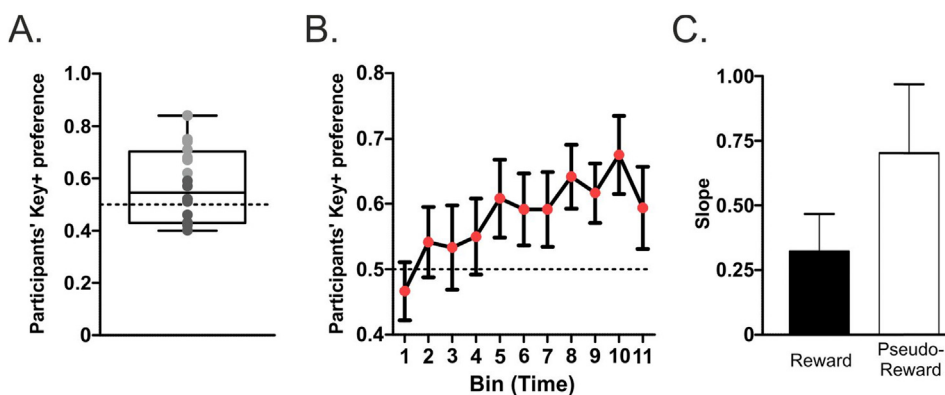


Fig. 3. Behavioral results across the 3×23 free choice trials of the task ($n = 20$). **A)** Box and whisker plot for the proportion of choices of the most pseudo-rewarding key (KEY+). The box extends from the 25th to 75th percentiles. The line in the middle of the box is plotted at the median. The whiskers represent the largest and smallest values no further than $1.5 \times IQR$ (interquartile range). Grey dots represent individuals who showed a significant preference towards the KEY+, while black dots represent individuals who did not show a preference for the KEY+ (based on a binomial test). **B)** Participants' preference across trials (mean $\pm$ SEM). Each bin comprises 6 trials. Note a clear increase from the initial trials to the last trials of the task, indicating that the preference for the KEY+ is acquired through learning of contingencies. **C)** Average regression weights for pseudo-reward and reward from the logistic regression analysis.

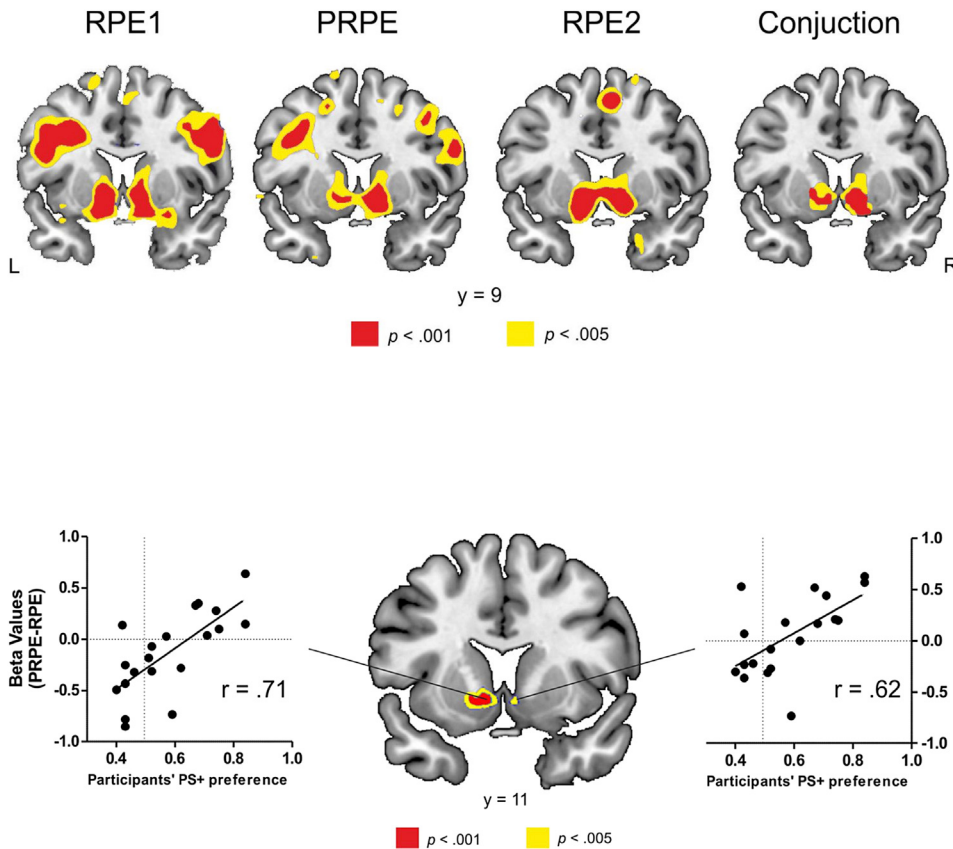


Fig. 4. Striatal representation of reward and pseudo-reward prediction errors. T-maps showing the brain regions in which a positive relationship was observed across trials between BOLD activity and the various errors terms, namely RPE1 (reward prediction error at the time of pseudo-feedback), PRPE (pseudo-reward prediction error at the time of pseudo-feedback) and RPE2 (reward prediction error at the time of final feedback). The last panel on the right illustrates the conjunction of the three previous relationships in the ventral striatum. Note that, in all T-maps, peak activations in the bilateral ventral striatum survive a voxel-wise threshold of $p < .01$ family-wise error (FWE) corrected within a small volume defined *a priori* based on an anatomical mask of the bilateral Nucleus Accumbens (Hammers et al., 2003). Coordinates are in MNI space.

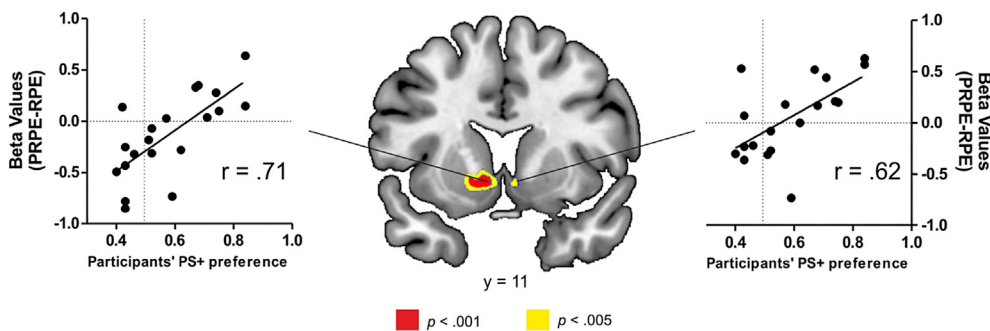


Fig. 5. Brain-behavior relationship. T-map displaying the striatal voxels in which a significant relationship was found between relative striatal sensitivity to PRPEs vs RPEs and the behavioral preference for the Key + across participants. Note that peak activations in the bilateral ventral striatum survive a threshold of $p < .05$ family-wise error (FWE) corrected within a small volume defined *a priori* based on an anatomical mask of the bilateral Nucleus Accumbens (Hammers et al., 2003). The scatter plot, shown for purely illustrative purposes, illustrates the same relationship in the striatal voxels found to be significant in the voxel-wise analysis. Coordinates are in MNI space.

4. Discussion

In the present study, we aimed to investigate whether pseudo-rewards could bias choice behavior –in a context where obtaining these pseudo-rewards does not confer any advantage– and whether such a bias might reflect striatal sensitivity to pseudo-rewards. In order to test these hypotheses, we developed a novel fMRI task in which participants had to first accomplish a subgoal (unlocking a box = pseudo-reward) to obtain a probabilistic monetary reward. Participants presented a significant preference towards the key that was unlocking more boxes (i.e. more pseudo-rewarding) even though it did not lead to more monetary reward. Importantly, the inclusion of forced-choice trials ensured that participants sampled both options equally, and thus that the behavioral preference did not result from imbalanced learning between the two arms of the task (Niv et al., 2002). This finding were replicated later in a second group of participants engaging in a modified version of the task using arbitrary cues instead of keys and locks; this suggests that the observed bias towards the most pseudo-rewarding option is not driven by pre-learned reward values. At the brain level, we observed a parallel representation of reward and pseudo-reward prediction errors in the ventral striatum. Notably, the relative striatal sensitivity to PRPEs vs RPEs predicted individual differences in the behavioral preference for the most pseudo-rewarding key. Our findings reveal a clear bias in participants' choices driven mainly by the availability of pseudo-rewards. However, this bias cannot be considered sub-optimal in the present study since both options were equally rewarded and the bias was thus not associated with a cost. Nevertheless, it provides evidence that individuals tend to rely on pseudo-reward information in order to guide their choices during learning.

Our fMRI results showed that both RPEs and PRPEs are

simultaneously encoded within the ventral striatum. This finding supports the idea that subgoal-related PEs recruit the same brain circuitry as goal-related PEs, i.e. they engage similar brain regions and elicit similar ERP components as classic signed and unsigned RPEs (Ribas-Fernandes et al., 2011, 2019; Diuk et al., 2013; Shahnazian et al., 2018). Moreover, it highlights the role of the striatum for computing prediction errors at several levels of complexity, and that dissociable prediction errors are integrated in this same structure (Daw et al., 2011; Diuk et al., 2013). In particular, our findings are in line with those previously reported by Diuk et al. (2013). This is remarkable given the differences in design between studies and the stringent orthogonalization procedure we used, which ensured that the PRPE regressor explained a significant amount of unique BOLD variance over and beyond the classic RPE. Additionally, Diuk et al. (2013) used a RL model in which RPEs only arose at the end of a trial when the outcome was presented. According to their model, reward expectancies did not vary from the first decision until the final outcome. In the context of the current task, their model would not anticipate potential outcomes based on whether the box was unlocked or not. In contrast, we have assumed a TD perspective, that is, a continuous-time model of learning in which any piece of information provided between the first decision and the final outcome can lead to a RPE and can be used to generate new reward expectations.

In addition, we have showed that the relative striatal sensitivity to PRPEs vs RPEs influences participants' behavior. Specifically, greater relative striatal sensitivity to PRPEs vs RPEs was associated with stronger preference for the most pseudo-rewarding key. In an analogous fashion, Daw et al. (2011) have shown that striatal sensitivity to model-free vs model-based RPEs signals predicted participants' preference for model-free vs model-based learning strategy. Together, these results not only suggest that the ventral striatum encodes simultaneous learning

signals, but also that these learning signals play a strong role in shaping behavior. However, the idea that the striatum may encode multiple and concomitant prediction error signals poses the question of how are these learning signals distinguished. This important question, which has been raised previously by Diuk et al. (2013) and Daw et al. (2011), will need to be addressed in future studies. One possibility is that the striatum is involved in integrating prediction error signals in order to guide behavior, rather than playing a role in learning *per se*. This line of reasoning agrees with previous studies suggesting that the ventral striatum does not have an exclusive role in encoding reward prediction errors, but is involved in encoding the behavioral relevance of events (Klein-Flügge et al., 2011; Li and Daw, 2011; FitzGerald et al., 2014). For instance, Klein-Flügge et al. (2011) have shown that, in contrast to midbrain activity reflecting RPE signals in the absence of a behavioral policy, the ventral striatum responds mainly to those events that are relevant to guide behavior. Some authors have further theorized that striatal activity might mediate the dynamic attribution of incentive salience to events, causing them and their associated actions to become more relevant for future decisions (Berridge, 2007).

It might be noted that the present results could also reflect sign-tracking behavior. Sign-tracking refers to the tendency of individuals to direct their behavior towards a stimulus that has become attractive and desirable as a result of a learned association between that stimulus (conditioned stimulus, CS) and a reward (such as food). In this context, individuals may become attracted towards the CS even when the CS is located far from the paired reward or in an opposite direction (Flagel et al., 2009). Thus, the CS does not only act as a predictor of reward but also acquires motivational properties, becoming a “motivational magnet”. Previous studies have shown that there are large individual differences in this type of behavior, with some individuals engaging almost exclusively with the CS (sign-tracking) while others use this information to predict the location of the reward (goal-tracking) (Flagel et al., 2009; Robinson et al., 2014; Morrison et al., 2015). Similarly, in the current task, some participants may have assigned motivational properties to the pseudo-reward while others may have not. Thus, pseudo-rewards may also induce sign-tracking behavior and may, in part, explain the large individual differences in solving complex structured tasks in our everyday life. However, this account remains speculative, and further studies comparing performance in well-established sign-tracking paradigms with participants' preference in the current task are required to better understand the relationship between pseudo-reward and sign-tracking behaviors.

It is important to emphasize that the behavioral bias towards the Key + cannot be explained by the desire to reduce uncertainty as reported in theories such as “temporal resolution of uncertainty” (Kreps and Porteus, 1978) or “information seeking” (Bromberg-Martin and Hikosaka, 2009). Previous studies have shown that animals avoid being uncertain about future outcomes (Bromberg-Martin and Hikosaka, 2009) and often prefer to have this uncertainty resolved earlier rather than later (Eliasz and Schotter, 2007). However, in the present study, the option leading to the biggest reduction in final outcome uncertainty was the least pseudo-rewarding key, which led to a fully predictable outcome (no-reward following locked box) in 70% of the cases. Thus, choosing the Key-provided more advance information about the upcoming outcome than the Key+, which was associated with full uncertainty resolution at the time of pseudo-feedback in only 30% of the trials.

In sum, we have shown that pseudo-rewards can influence participants' choices independently of the reward value associated with them. Thus, although they are critical for speeding up learning in complex environments, they might potentially lead to behavioral biases towards the accomplishment of subgoals. Additionally, using a novel paradigm we have shown that RPEs and PRPEs are simultaneously encoded in the ventral striatum, and that the relative sensitivity of the striatum to PRPEs vs RPEs predicted participants' choices, emphasizing the relevance of the striatum in complex learning problems and decision-making.

Acknowledgments

The present project has been funded by the Spanish Government (MINECO Grants PSI2012-37472 and PSI2015-69664-P to J.M.-P.). J.M.-P. was supported by the ICREA Academia program. E.M.-H. was supported by the FPI program (BES-2010-032702). G.S. was supported by a Veni grant from the Netherlands Organisation for Scientific Research (NWO). We want to thank Dr. Martijn van Otterlo and Dr. Payam Piray for their insightful comments and suggestions during the first stages of the project.

Appendix A. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.neuroimage.2019.02.052>.

References

- Berridge, K.C., 2007. The debate over dopamine's role in reward: the case for incentive salience. *Psychopharmacology (Berlin)* 191, 391–431.
- Botvinick, M.M., 2012. Hierarchical reinforcement learning and decision making. *Curr. Opin. Neurobiol.* 22, 956–962.
- Botvinick, M.M., Niv, Y., Barto, A.C., 2009. Hierarchically organized behavior and its neural foundations: a reinforcement learning perspective. *Cognition* 113, 262–280.
- Bromberg-Martin, E.S., Hikosaka, O., 2009. Midbrain dopamine neurons signal preference for advance information about upcoming rewards. *Neuron* 63 (1), 119–126.
- Daw, N.D., Gershman, S.J., Seymour, B., Dayan, P., Dolan, R.J., 2011. Model-based influences on humans' choices and striatal prediction errors. *Neuron* 69, 1204–1215.
- Diuk, C., Tsai, K., Wallis, J., Botvinick, M., Niv, Y., 2013. Hierarchical learning induces two simultaneous, but separable, prediction errors in human basal ganglia. *J. Neurosci.* 33, 5797–5805.
- Eliasz, K., Schotter, A., 2007. Experimental testing of intrinsic preferences for noninstrumental information. *Am. Econ. Rev.* 97 (2), 166–169.
- Erdeniz, B., Rohe, T., Done, J., Seidler, R., 2013. A simple solution for model comparison in bold imaging: the special case of reward prediction error and reward outcomes. *Front. Neurosci.* 7, 116.
- FitzGerald, T.H.B., Schwartenbeck, P., Dolan, R.J., 2014. Reward-related activity in ventral striatum is action contingent and modulated by behavioral relevance. *J. Neurosci.* 34 (4), 1271–1279.
- Flagel, S.B., Akil, H., Robinson, T.E., 2009. Individual differences in the attribution of incentive salience to reward-related cues: implications for addiction. *Neuropharmacology* 56, 139–148.
- Friston, K.J., Penny, W.D., Glaser, D.E., 2005. Conjunction revisited. *Neuroimage* 25, 661–667.
- Hammers, A., Allom, R., Koepf, M.J., Free, S.L., Myers, R., Lemieux, L., Mitchell, T.N., Brooks, D.J., Duncan, J.S., 2003. Three-dimensional maximum probability atlas of the human brain, with particular reference to the temporal lobe. *Hum. Brain Mapp.* 19, 224–247.
- Hengst, B., 2012. Hierarchical approaches. In: *Reinforcement Learning*. Springer, Berlin, Heidelberg, pp. 293–323.
- Klein-Flügge, M.C., Hunt, L.T., Bach, D.R., Dolan, R.J., Behrens, T.E.J., 2011. Dissociable reward and timing signals in human midbrain and ventral striatum. *Neuron* 72, 654–664.
- Kreps, D.M., Porteus, E.L., 1978. Temporal resolution of uncertainty and dynamic choice theory. *Econometrica: J. Econometr. Soc.* 185–200.
- Lee, D., Seo, H., Jung, M.W., 2012. Neural basis of reinforcement learning and decision making. *Annu. Rev. Neurosci.* 35, 287–308. <https://doi.org/10.1146/annurev-neuro-062111-150512>.
- Li, J., Daw, N.D., 2011. Signals in human striatum are appropriate for policy update rather than value prediction. *J. Neurosci.: Off. J. Soc. Neurosci.* 31 (14), 5504–5511.
- Mas-Herrero, E., Marco-Pallarés, J., 2016. Theta oscillations integrate functionally segregated sub-regions of the medial prefrontal cortex. *Neuroimage* 143, 166–174.
- Mazaika, P., Whitfield-Gabriel, S., Reiss, A., 2007. Artifact repair for fMRI data from high motion clinical subjects. In: *Presentation at Human Brain Mapping Conference*.
- Morrison, S.E., Bamkole, M.A., Nicola, S.M., 2015. Sign tracking, but not goal tracking, is resistant to outcome devaluation. *Front. Neurosci.* 9.
- Nichols, T., Brett, M., Andersson, J., Wager, T., Poline, J.-B., 2005. Valid conjunction inference with the minimum statistic. *Neuroimage* 25, 653–660.
- Niv, Y., Joel, D., Meilijson, I., Ruppin, E., 2002. Evolution of reinforcement learning in uncertain environments: a simple explanation for complex foraging behaviors. *Adapt. Behav.* 10, 5–24.
- Niv, Y., Schoenbaum, G., 2008. Dialogues on prediction errors. *Trends Cognit. Sci.* 12 (7), 265–272. <https://doi.org/10.1016/j.tics.2008.03.006>.
- Ribas-Fernandes, J.J., Solway, A., Diuk, C., McGuire, J.T., Barto, A.G., Niv, Y., Botvinick, M.M., 2011. A neural signature of hierarchical reinforcement learning. *Neuron* 71 (2), 370–379.
- Ribas-Fernandes, J.J., Shahnazian, D., Holroyd, C.B., Botvinick, M.M., 2019. Subgoal-and goal-related reward prediction errors in medial prefrontal cortex. *J. Cognit. Neurosci.* 31 (1), 8–23.

- Robinson, M.J., Anselme, P., Fischer, A.M., Berridge, K.C., 2014. Initial uncertainty in Pavlovian reward prediction persistently elevates incentive salience and extends sign-tracking to normally unattractive cues. *Behav. Brain Res.* 266, 119–130.
- Schultz, W., Dayan, P., Montague, P.R., 1997. A neural substrate of prediction and reward. *Science* 275, 1593–1599.
- Seymour, B., O'Doherty, J.P., Dayan, P., Koltzenburg, M., Jones, A.K., Dolan, R.J., Frackowiak, R.S., 2004. Temporal difference models describe higher-order learning in humans. *Nature* 429, 664–667.
- Shahnazian, D., Shulver, K., Holroyd, C.B., 2018. Electrophysiological responses of medial prefrontal cortex to feedback at different levels of hierarchy. *Neuroimage* 183, 121–131.
- Sutton, R., Barto, A., 1998. *Reinforcement Learning: an Introduction* (Adaptive Computation and Machine Learning). The MIT Press.
- van Dijk, S.G., Polani, D., Nehaniv, C.L., 2011. Hierarchical behaviours: getting the most bang for your bit. In: Kampis, G., Karsai, I., Szathmáry, E. (Eds.), *Advances in Artificial Life. Darwin Meets von Neumann*. Springer Berlin Heidelberg, pp. 342–349.
- Vilà-Balló, A., Mas-Herrero, E., Ripollés, P., Simó, M., Miró, J., Cucurell, D., et al., 2017. Unraveling the role of the hippocampus in reversal learning. *J. Neurosci.* 37 (28), 6686–6697.
- Weiskopf, N., Hutton, C., Josephs, O., Deichmann, R., 2006. Optimal EPI parameters for reduction of susceptibility-induced BOLD sensitivity losses: a whole-brain analysis at 3 T and 1.5 T. *Neuroimage* 33 (2), 493–504.
- Wilson, R.C., Niv, Y., 2015. Is model fitting necessary for model-based fMRI? *PLoS Comput. Biol.* 11 (6), e1004237.