

1                                   **Introducing the CORTEX database:**  
2                   **COntext-dependent Reinforcement learning and Transfer EXperiments**

3                                   **Isabelle Hoxha (isabellehoxha@gmail.com)**

4                                   Leiden University, Wassenaarseweg 52, 2333 AK Leiden, The Netherlands  
5  
6

7                                   **Stefano Palminteri (stefano.palminteri@gmail.com)**

8                                   INSERM, Ecole Normale Supérieure, 29 rue d'Ulm, 75005 Paris, France  
9  
10

## Abstract

Learning - Transfer paradigms are particularly relevant to study value learning, as they can uncover contextual value learning. However, there is no consensus within and across laboratories on the implementation of the experimental variable nor the exact computational processing underlying context-dependence, thus making it unclear whether models such as reference-point and range adaptation genuinely generalize across contexts or are merely tailored to specific task structures. We therefore created an extremely large ( $n > 2500$  human participants) yet extremely curated dataset to establish a standardized framework and systematically evaluate the applicability of competing value learning models. We showed that the range adaptation model is the best fitting model for half of the participants, with task specifications such as feedback type modulating the relevance of in-context learning models.

**Keywords:** open science; reinforcement learning; range adaptation; reference-point centering

## Introduction

Increasing evidence suggest that learning values in multi-armed bandit tasks is done in context, i.e. the values are not learned with their absolute magnitude but rather in relation to the alternatives (Palminteri et al., 2015; Bavard et al., 2018; Lebreton et al., 2019; Hayes & Wedell, 2023b). However, task specifications such as feedback type, number of alternatives, or magnitudes displayed, seem to alter the predictive power of contextual learning models. Here, we addressed this issue by compiling to a unique format 31 experiments from 10 published studies, gathering more than 2500 human participants and 700,000 trials, with transfer testing after the initial learning phase. These datasets were subsequently fit using three reinforcement-learning model variants which take into account the context in different ways.

## Methods

**Data format.** The dataset only gathers bandit tasks in published experiments. The datasets were selected on

the basis of the presence of a learning phase, where fixed ensembles of contingencies are presented together, and a transfer phase, where all pairs of stimuli are presented together. While all experiments have their specificities, a common subset of variables was kept: the experiment value (i.e. the expected value across the full experiment), the agent number, the trial number, the trial order within the presented ensemble, the session number, a binary marker of whether the trial was a transfer or not, the stimuli presented, the expected values of the stimuli, the availability of the stimuli (i.e. whether there was a forced choice or observational trial), the choice, the outcomes (obtained and foregone when available), the response times, and the regret associated with the trial. Any missing information was replaced by a NaN entry.

**Models.** At the present time, we compared the fitness of three reinforcements learning models on the collected data.

**Absolute model.** This model corresponds to a classical Q-learning rule. At each trial, when participants receive feedback  $r$ , Q-values are adapted according to:

$$Q_i \leftarrow Q_i + \alpha_i(r_i - Q_i)$$

We set  $\alpha_i = \alpha_c$  when option  $i$  was chosen, and  $\alpha_i = \alpha_u$  when option  $i$  was not chosen.

**Reference-point centering.** Before the updating Q-values, reference-point centering assumes that the rewards are compared to the mean of received rewards  $V$  within a context:

$$\tilde{r} = r - V$$

The referenced rewards  $\tilde{r}$  are subsequently used to update Q-values. The mean value is then updated based on the mean of observed rewards  $\bar{r}$ :

$$V \leftarrow \alpha_{ref}(\bar{r} - V)$$

**Range adaptation.** The rewards are compared to the range of the rewards received within a context, with the range adapting at each trial. In this work, the following variant was implemented:

$$\tilde{r} = \left( \frac{r - R_{\min}}{R_{\max} - R_{\min} + 1} \right)^{w_r}$$

The range boundaries are updated at rate  $\alpha_R$  if the rewards are outside of the range, following:

$$\begin{cases} R_{\min} \leftarrow R_{\min} + \alpha_R(\min(r) - R_{\min}) \\ R_{\max} \leftarrow R_{\max} + \alpha_R(\max(r) - R_{\max}) \end{cases}$$

## Results

**Datasets.** The datasets included are summarized on Table 1. In total, data from 10 published studies, spanning 31 experiments and 2534 agents have been included, for a total of 700,950 trials.

Table 1: Included datasets

| Paper                        | N.exp      | N.agents |
|------------------------------|------------|----------|
| Bavard & Palminteri, 2023    | 4 (1-4)    | 500      |
| Hayes & Wedell, 2023         | 1 (5)      | 64       |
| Vandendriessche et al., 2023 | 1 (6)      | 26       |
| Bavard et al., 2018          | 2 (7, 8)   | 60       |
| Gueguen et al., 2024         | 1 (9)      | 44       |
| Hayes & Wedell, 2023b        | 2 (10,11)  | 222      |
| Hayes & Wedell, 2023c        | 1 (12)     | 50       |
| Salem-Garcia et al., 2023    | 10 (13-22) | 207      |
| Bavard et al., 2021          | 8 (23-30)  | 800      |
| Anlló et al., 2024           | 11 (31)    | 561      |

**Fitting results.** We present here the best winning model on the basis of the Bayesian Information Criterion (BIC) and the Akaike Information Criterion (AIC). We represent the share of each winning model per participant across all experiments and within experiment.

On the basis of the AIC, the range adaptation model was the best for 1254 participants, the reference-point centering model for 186 participants and the absolute model for 1094 participants. On the basis of the BIC, the absolute model was this time the best, being the best fitting model for 1481 participants, followed by the range adaptation model (916 participants) and the reference-point centering model (137 participants).

Importantly, the best fitting model depended on participants and experiment. We observed that the absence of feedback in the transfer phase enhanced context-dependent learning (over global learning).

Anlló, H., Bavard, S., Benmarrakchi, F., Bonagura, D., Cerrotti, F., Cicue, M., Gueguen, M., Guzmán, E. J., Kadieva, D., Kobayashi, M., Lukumon, G., Sartorio, M., Yang, J., Zinchenko, O., Bahrami,

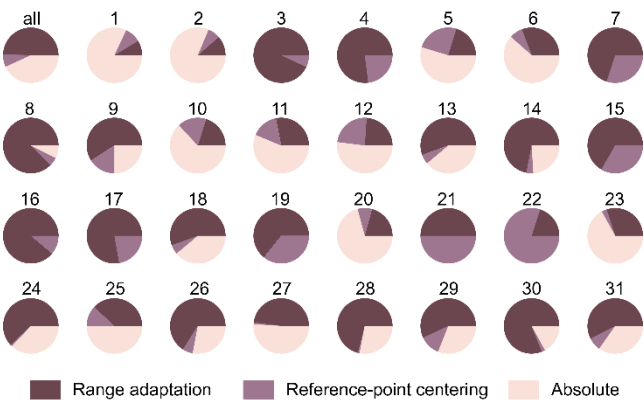


Figure 1: Proportion of best fitting model on the basis of AIC for all experiments and split per experiment.

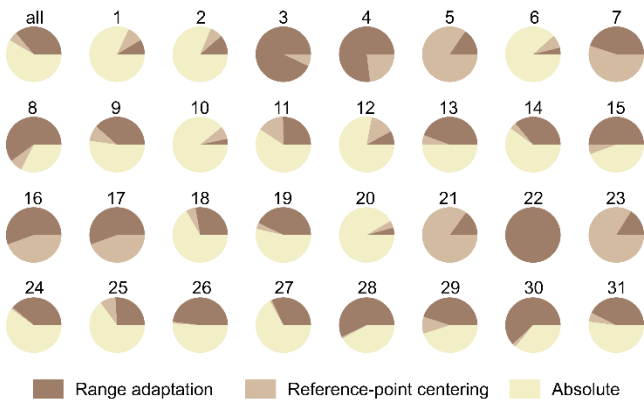


Figure 2: Best fitting model on the basis of BIC for all experiments and split per experiment.

## Conclusion

Over the 31 selected experiments, we showed that the range adaptation model was the best selected model. We noted that the selection of the best model depended on the task, and therefore the absolute encoding model was sometimes better suited, in particular when a unique range was presented to participants. The next steps of this study will include an analysis of parameter recovery as a function of task specifications, as well as other variants of in-context value learning.

## References

B., Silva Concha, J., Hertz, U., Konova, A. B., Li, J., ... Palminteri, S. (2024). Comparing experience- and description-based economic preferences across 11 countries. *Nature*

146 *Human Behaviour*, 8(8), 1554–1567.  
 147 <https://doi.org/10.1038/s41562-024-01894-9>  
 148 Bavard, S., Lebreton, M., Khamassi, M., Coricelli, G., &  
 149 Palminteri, S. (2018). Reference-point centering  
 150 and range-adaptation enhance human  
 151 reinforcement learning at the cost of irrational  
 152 preferences. *Nature Communications*, 9(1),  
 153 4503. [https://doi.org/10.1038/s41467-018-](https://doi.org/10.1038/s41467-018-06781-2)  
 154 06781-2  
 155 Bavard, S., & Palminteri, S. (2023). The functional form  
 156 of value normalization in human reinforcement  
 157 learning. *eLife*, 12, e83891.  
 158 <https://doi.org/10.7554/eLife.83891>  
 159 Bavard, S., Rustichini, A., & Palminteri, S. (2021). Two  
 160 sides of the same coin: Beneficial and  
 161 detrimental consequences of range adaptation in  
 162 human reinforcement learning. *Science*  
 163 *Advances*, 7(14), eabe0340.  
 164 <https://doi.org/10.1126/sciadv.abe0340>  
 165 Gueguen, M. C. M., Anlló, H., Bonagura, D., Kong, J.,  
 166 Hafezi, S., Palminteri, S., & Konova, A. B.  
 167 (2024). Recent Opioid Use Impedes Range  
 168 Adaptation in Reinforcement Learning in Human  
 169 Addiction. *Biological Psychiatry*, 95(10), 974–  
 170 984.  
 171 <https://doi.org/10.1016/j.biopsych.2023.12.005>  
 172 Hayes, W. M., & Wedell, D. H. (2023a). Effects of  
 173 blocked versus interleaved training on relative  
 174 value learning. *Psychonomic Bulletin & Review*,  
 175 30(5), 1895–1907.  
 176 <https://doi.org/10.3758/s13423-023-02290-6>  
 177 Hayes, W. M., & Wedell, D. H. (2023b). Reinforcement  
 178 learning in and out of context: The effects of  
 179 attentional focus. *Journal of Experimental*  
 180 *Psychology: Learning, Memory, and Cognition*,  
 181 49(8), 1193–1217.  
 182 <https://doi.org/10.1037/xlm0001145>  
 183 Hayes, W. M., & Wedell, D. H. (2023c). Testing models  
 184 of context-dependent outcome encoding in  
 185 reinforcement learning. *Cognition*, 230, 105280.  
 186 <https://doi.org/10.1016/j.cognition.2022.105280>  
 187 Lebreton, M., Bacić, K., Palminteri, S., & Engelmann, J.  
 188 B. (2019). Contextual influence on confidence  
 189 judgments in human reinforcement learning.  
 190 *PLOS Computational Biology*, 15(4), e1006973.  
 191 <https://doi.org/10.1371/journal.pcbi.1006973>  
 192 Palminteri, S., Khamassi, M., Joffily, M., & Coricelli, G.  
 193 (2015). Contextual modulation of value signals in  
 194 reward and punishment learning. *Nature*  
 195 *Communications*, 6(1), 8096.  
 196 <https://doi.org/10.1038/ncomms9096>  
 197 Salem-Garcia, N., Palminteri, S., & Lebreton, M.  
 198 (2023). Linking confidence biases to  
 199 reinforcement-learning processes.  
 200 *Psychological Review*, 130(4), 1017–1043.  
 201 <https://doi.org/10.1037/rev0000424>  
 202 Vandendriessche, H., Demmou, A., Bavard, S.,  
 203 Yadaq, J., Lemogne, C., Mauras, T., &  
 204 Palminteri, S. (2023). Contextual influence of  
 205 reinforcement learning performance of  
 206 depression: Evidence for a negativity bias?  
 207 *Psychological Medicine*, 53(10), 4696–4706.  
 208 <https://doi.org/10.1017/S0033291722001593>  
 209