

In the format provided by the authors and unedited.

Multimodal imaging of brain connectivity reveals predictors of individual decision strategy in statistical learning

Vasilis M. Karlaftis ^{1,6}, Joseph Giorgio^{1,6}, Petra E. Vértes², Rui Wang³, Yuan Shen⁴, Peter Tino⁵, Andrew E. Welchman^{1*} and Zoe Kourtzi ^{1*}

¹Department of Psychology, University of Cambridge, Cambridge, UK. ²Department of Psychiatry, Behavioural and Clinical Neuroscience Institute, University of Cambridge, Cambridge, UK. ³Key Laboratory of Mental Health, Institute of Psychology, Chinese Academy of Sciences, Beijing, China. ⁴School of Science and Technology, Nottingham Trent University, Nottingham, UK. ⁵School of Computer Science, University of Birmingham, Birmingham, UK. ⁶These authors contributed equally: Vasilis M. Karlaftis, Joseph Giorgio. *e-mail: aew69@cam.ac.uk; zk240@cam.ac.uk

Supplementary Methods

Stimuli: Stimuli comprised four symbols chosen from Ndjuká syllabary (**Figure 1a**) that were highly discriminable from each other and were unfamiliar to the participants. Each symbol subtended 8.5° of visual angle and was presented in black on a mid-grey background. Experiments were controlled using Matlab and the Psychophysics toolbox ^{31,2}. For the behavioural training sessions, stimuli were presented on a 21-inch CRT monitor (ViewSonic P225f 1280 x 1024 pixel, 85 Hz frame rate) at a distance of 45 cm. For the test sessions, stimuli were presented using a projector and a mirror set-up (1280 x 1024 pixel, 60 Hz frame rate) at a viewing distance of 67.5 cm. The physical size of the stimuli was adjusted so that the angular size was constant during training and test sessions.

Sequence design: We generated probabilistic sequences by using a temporal Markov model and varying the memory length (i.e. context length) of the sequence, following our previous work³. The model consists of a series of symbols, where the symbol at time i is determined probabilistically by the previous ' k ' symbols. We refer to the symbol presented at time i , $s(i)$, as the target and to the preceding k -tuple of symbols ($s(i-1)$, $s(i-2)$, ..., $s(i-k)$) as the context. The value of ' k ' is the order or level of the sequence:

$$P(s(i) | s(i-1), s(i-2), \dots, s(1)) = P(s(i) | s(i-1), s(i-2), \dots, s(i-k)), k < i$$

In our study, we used three levels of memory length; for $k=0,1,2$. The simplest $k=0^{th}$ order model is a memory-less source. This generates, at each time step i , a symbol according to symbol probability $P(s)$, without taking into account the context (i.e. previously generated symbols). The order $k=1$ Markov model generates symbol $s(i)$ at each time i conditional on the previously generated symbol $s(i-1)$. This introduces a memory in the sequence; i.e. the probability of a particular symbol at time i strongly depends on the preceding symbol $s(i-1)$. Unconditional symbol probabilities $P(s(i))$ for the case $k=0$ are now replaced with conditional

ones, $P(s(i)|s(i-1))$). Similarly, an order $k=2$ Markov model generates a symbol $s(i)$ at each time i conditional on the two previously generated symbols $s(i-1)$, $s(i-2)$: $P(s(i)|s(i-1),s(i-2))$.

At each time the symbol that follows a given context is determined probabilistically, thus generating stochastic Markov sequences. The underlying Markov model can be represented through the associated context-conditional target probabilities (**Figure 1b**). We used 4 symbols that we refer to as items A, B, C and D. The correspondence between items and symbols was counterbalanced across participants. Note, that we designed the stochastic sources from which the sequences were generated so that the memory-conditional uncertainty remains the same across levels. In particular, for the zero-order source, only two symbols are likely to occur most of the time; the remaining two symbols have very low probability (0.05); this is introduced to ensure that there is no difference in the number of symbols across levels. Of the two dominant symbols, one is more probable (probability 0.72) than the other (probability 0.18). This structure is preserved in Markov chain of order 1 and 2, where conditional on the previous symbols, only two symbols are allowed to follow, one with higher probability (0.80) than the other (0.20). This ensures that the structure of the generated sequences across levels differs mainly in the memory length (i.e. context length) rather than the context-conditional probabilities.

In particular, for level-0 (zero-order), the Markov model was based on the probability of symbol occurrence: one symbol had a high probability of occurrence, one low probability, while the remaining two symbols appeared rarely (**Figure 1b**). For example, the probabilities of occurrence for the four symbols A, B, C and D were 0.18, 0.72, 0.05 and 0.05, respectively. Presentation of a given symbol was independent of the items that preceded it. For level-1 (first-order) and level-2 (second-order), the target depended on one or two immediately preceding items, respectively (**Figure 1b**). Given a context, only one of two targets could follow; one had a high probability of being presented and the other a low probability (e.g., 80% vs. 20%). For

example, when Symbol A was presented, only symbols B or C were allowed to follow, and B had a higher probability of occurrence than C.

Note, that we designed the stochastic sources from which the sequences were generated so that the memory-conditional uncertainty remains the same across levels. In particular, for the zero-order source (level-0), only two symbols are likely to occur most of the time; the remaining two symbols have very low probability (0.05); this is introduced to ensure that there is no difference in the number of symbols across levels. Of the two dominant symbols, one is more probable (probability 0.72) than the other (probability 0.18). This structure is preserved in Markov chain of order 1 (level-1) and 2 (level-2), where conditional on the previous symbols, only two symbols are allowed to follow, one with higher probability (0.80) than the other (0.20). This ensures that the structure of the generated sequences across levels differs mainly in the memory length (i.e. context length) rather than the context-conditional probabilities.

Procedure: Participants were initially familiarized with the task through a brief practice session (8 minutes) with random sequences (i.e. all four symbols were presented with equal probability 25% in a random order). Following this, participants took part in multiple behavioural training and test sessions that were conducted on different days. In addition, they participated in two brain imaging sessions, one before the first training session and one after the last training session. Participants were trained with structured sequences and tested with both structured and random sequences to ensure that training was specific to the trained sequences.

In the first test session (pre-training), participants were presented with level-0, level-1 and level-2 sequences and random sequences. Participants were then trained with level-0 sequences, and subsequently with level-1 and level-2 sequences. Training on level-0 sequences involves learning frequency statistics (i.e. participants are required to learn the occurrence

probability of each symbol), whereas training on level-1 and level-2 sequences involves learning context-based statistics (i.e. participants are required to learn the probability of a given symbol appearing depends on the preceding symbol(s)). For each level, participants completed a minimum of 3 and a maximum of 5 training sessions (840-1400 trials). Each training session comprised five blocks of structured sequences (56 trials per block) and lasted one hour. Training at each level ended when participants reached plateau performance (i.e. performance did not change significantly for two sessions). Participants were given feedback (i.e. score in the form of Performance Index) at the end of each block, rather than per-trial error feedback, which motivated them to continue with training. A post-training test session followed training per level (i.e. on the following day after completion of training) during which participants were presented with structured sequences determined by the statistics of the trained level and random sequences (90 trials each). In contrast to the training sessions, no feedback was given during test. The mean time interval ($\pm$ standard deviation) between the pre-training and the post-training test sessions was 23.3 ($\pm$ 2.5) days.

For each trial, a sequence of 8-14 symbols appeared in the center of the screen, one at a time in a continuous stream (**Figure 1a**). This variable trial length ensured that participants maintained attention during the whole trial. The end of each trial was indicated by a red dot cue. Following this, all four symbols were shown in a 2x2 grid. The positions of test stimuli were randomized from trial to trial. Participants were asked to indicate which symbol they expected to appear following the preceding sequence by pressing a key corresponding to the location of the predicted symbol.

Psychophysical training: To ensure that sequences in each block were representative of the Markov model order per level, we generated 10,000 Markov sequences per level comprising

672 items per sequence. To quantify how close the generated sequence was to the ideal Markov model, we estimated the Kullback-Leibler divergence (KL divergence) as follows:

$$KL = \sum_{target} Q(target) \log \left(\frac{Q(target)}{P(target)} \right)$$

for the level-0 model, and

$$KL = \sum_{context} Q(context) \sum_{target} Q(target|context) \log \left(\frac{Q(target|context)}{P(target|context)} \right)$$

for the level-1 and level-2 models, where $P()$ refers to probabilities or conditional probabilities derived from the presented sequence and $Q()$ refers to those specified by the ideal Markov model. KL divergence is a standard measure of distance between distributions and values close to 0 indicate small differences between the distributions. We selected fifty sequences with the lowest KL divergence (i.e. these sequences matched closely the Markov model per level). The sequences presented to the participants during the experiments were selected randomly from this sequence set.

For each trial, a sequence of 8-14 symbols appeared in the center of the screen, one at a time in a continuous stream, each for 300ms followed by a central white fixation dot (ISI) for 500ms (**Figure 1a**). This variable trial length ensured that participants maintained attention during the whole trial. Each block comprised equal number of trials with the same number of items. The end of each trial was indicated by a red dot cue that was presented for 500ms. Following this, all four symbols were shown in a 2x2 grid. The positions of test stimuli were randomized from trial to trial. Participants were asked to indicate which symbol they expected to appear following the preceding sequence by pressing a key corresponding to the location of the predicted symbol. Participants learned a stimulus-key mapping during the familiarization phase: key ‘8’, ‘9’, ‘5’ and ‘6’ in the number pad corresponded to the four positions of the test stimuli —upper left, upper right, lower left and lower right, respectively. After the participant’s response, a white circle appeared on the selected item for 300ms to indicate the participant’s

choice, followed by a fixation dot for 150ms (ITI) before the start of the next trial. If no response was made within 2s, a null response was recorded and the next trial started.

Test sessions: The pre-training test session (Pre) included nine runs (i.e. three runs per level), the order of which was randomized across participants. Test sessions after training per level included nine runs of structured sequences determined by the same statistics as the corresponding trained level and random sequences. Each run comprised five blocks of structured and five blocks of random sequences presented in a random counterbalanced order (2 trials per block; a total of 10 structured and 10 random trials per run), with an additional two 16s fixation blocks, one at the beginning and one at the end of each run. Each trial comprised a sequence of 10 stimuli which were presented for 250ms each, separated by a blank interval during which a white fixation dot was presented for 250ms. Following the sequence, a response cue (central red dot) appeared on the screen for 4s before the test display (comprising four test stimuli) appeared for 1.5s. Participants were asked to indicate which symbol they expected to appear following the preceding sequence by pressing a key corresponding to the location of the predicted symbol. A white fixation was then presented for 5.5s before the start of the next trial.

Performance index: We assessed participant responses in a probabilistic manner. We computed a performance index per context that quantifies the minimum overlap (min: minimum) between the distribution of participant responses and the distribution of presented targets estimated across 56 trials per block by:

$$PI(\text{context}) = \sum \min (P_{\text{resp}}(s_t|\text{context}_t), P_{\text{pres}}(s_t|\text{context}_t))$$

where t is the trial index and the target s is from the symbol set A, B, C and D.

The overall performance index is then computed as the average of the performance indices across contexts, $PI(\text{context})$, weighted by the corresponding context probabilities:

$$PI = \sum PI(\text{context}) \cdot P(\text{context}).$$

To compare across different levels, we defined a normalized PI measure that quantifies relative participant performance above random guessing. We computed a random guess baseline; i.e. performance index PI_{rand} that reflects participant responses to targets with a) equal probability of 25% for each target per trial for level-0 ($PI_{\text{rand}} = 0.53$); b) equal probability for each target for a given context for level-1 ($PI_{\text{rand}} = 0.45$) and level-2 ($PI_{\text{rand}} = 0.44$). To correct for differences in random-guess baselines across levels, we subtracted the random guess baseline from the performance index ($PI_{\text{normalized}} = PI - PI_{\text{rand}}$).

Strategy choice and strategy index: To quantify each participant's strategy, we compared individual participant response distributions (response-based model) to two baseline models: (i) a probability matching model, where probabilistic distributions of possible outcomes are derived from the Markov models that generated the presented sequences (Model-matching), and (ii) a probability maximization model, where only the most likely outcome is allowed for each context (Model-maximization). We used KL divergence to quantify how close the response distribution is to matching and maximization distributions. KL divergence close to 0 indicates small difference between the distributions. KL is defined as follows:

$$KL = \sum_{\text{target}} M(\text{target}) \log\left(\frac{M(\text{target})}{R(\text{target})}\right)$$

for the level-0 model, and

$$KL = \sum_{\text{context}} M(\text{context}) \sum_{\text{target}} M(\text{target}|\text{context}) \log\left(\frac{M(\text{target}|\text{context})}{R(\text{target}|\text{context})}\right)$$

for the level-1 and level-2 models, where $R(\cdot)$ and $M(\cdot)$ denote the probability distribution or conditional probability distribution derived from the human responses and the models (i.e. probability matching or maximization) respectively, across all the conditions.

We quantified the difference between the KL divergence from the response-based model to Model-matching and the KL divergence from the response-based model to Model-maximization. We refer to this quantity as strategy choice indicated by $\Delta\text{KL}(\text{Model-maximization, Model-matching})$ and it reflects the participant's preference towards matching or maximization. We then derived an individual strategy index by calculating the integral of each participant's strategy curve across trials and subtracting it from the integral of the exact matching curve across trials, as defined by Model-matching. We defined the integral curve difference (ICD) between individual strategy and exact matching as the individual strategy index. That is, strategy index close to zero indicates a strategy closer to matching, while higher positive values indicate deviation from matching towards maximization.

Supplementary Figure 1 illustrates how the response probability distributions may yield negative or positive strategy index values. For example, for level-1, Table A shows the context-target probability distribution that defines the matching model; a participant response distribution matching this model would indicate exact matching strategy. Table B represents the exact maximization model; that is, a participant whose response distribution follows this model chooses consistently the most probable outcome. Table C represents a random response model; that is, the participant chooses all context-target contingencies with equal probability. Participants may demonstrate this random distribution of responses at the beginning of learning before they have extracted the structure of the sequence or the exact context-target contingencies. Following training, participants may show response distributions closer to matching or deviating from matching towards maximization. Underestimating the probability of the most probable context-target contingency (e.g. Table D) will result in response distributions between the matching and the random model and yield a negative strategy index. In contrast, overestimating the probability of the most probable context-target contingency (e.g.

Table E) will result in response distributions between the matching and maximization models and yield a positive strategy index.

Further, response distributions during training (i.e. strategy choice per block: $\Delta\text{KL}(\text{Model-maximization}, \text{Model-matching})$) from three representative participants are shown in comparison to these models (matching, maximization, random) (**Supplementary Figure 1c**). Note that the strategy index is computed as the integral between the values of participant strategy choice and the matching model across blocks. As a result, calculating the strategy index for a participant that starts with a strategy closer to random and then deviates closer to the matching model may result in a negative (e.g. participant A) or a positive value (e.g. participant B). For example, data from a participant A that underestimates the probability of the most probable context-target contingency during most of the training blocks yield a negative strategy index. However, data from a participant B that overestimates the probability of the most probable context-target contingency in some of the training blocks yield a positive strategy index, as the integral becomes positive when the participant strategy crosses the matching model curve. In contrast, strategy choice data for a participant C that deviates from matching towards maximization yields a higher positive strategy index.

Further, we provide a mathematical description of strategy index variability. In particular, we generated synthetic response data from a virtual participant and present a two-parameter model characterizing the participant response distribution. Response distribution (denoted as P) is described as the mixture of two components, P_1 and P_2 . To control the contribution of these two components, we define a parameter β as the weight of the two components ($0 \leq \beta \leq 1$): $P = \beta P_1 + (1 - \beta) P_2$. The first component is the random model (i.e. equal probabilities for all context-target contingencies). Participants may follow this random model of responses at the beginning of training before they have learned the sequence structure and relative probabilities. The second component reflects the probability distribution of the items

in the sequence presented to the participant, e.g. $P_2 = [0.2, 0.8, 0, 0]$. This specification assumes that (1) only two items have non-zero probability; (2) the high probable target is four times more frequent than the less probable target. To capture how the participants learn these contingencies, we parameterized this distribution as follows: $P_2 = [1-\alpha, \alpha, 0, 0]$, where $0 \leq \alpha \leq 1$. In particular, for (i) $\alpha = 1$, the participant predicts always the most probable target (i.e. maximization); (ii) $\alpha = 0.8$, the participant responses match the target distribution (i.e. matching); (iii) $\alpha = 0.5$, the participant predicts equally the two possible (non-zero probability) targets; (iv) $\alpha < 0.5$, the participant predicts the less probable target more frequently than the more probable target. In sum, we formulate our synthetic response model as follows: $P = \beta [0.25, 0.25, 0.25, 0.25] + (1-\beta) [1-\alpha, \alpha, 0, 0]$.

To illustrate how the strategy index varies with parameters α and β , we computed the strategy index for all possible combinations of α and β values, where α and β vary between 0 and 1. This generated a strategy index surface as a function of α and β (**Supplementary Figure 2**). In particular, for $\beta = 1$ the strategy index is invariant to the parameter α and reflects equal responses for all targets (i.e. random model); yielding a strategy index value of -0.26. For $\beta = 0$, the model is reduced to $P = [1-\alpha, \alpha, 0, 0]$ and is fully described by the P_2 component (see above). Therefore, (i) for $\alpha = 1$ the model describes a maximization response (i.e. strategy index = 0.63), (ii) for $\alpha = 0.8$ it describes a matching response (i.e. strategy index = 0), (iii) for $\alpha = 0.5$ it describes a random response between the two possible targets (i.e. strategy index = -0.26) and (iv) for $\alpha < 0.5$ it describes predictions of the less probable target more frequently than the more probable target (i.e. strategy index < -0.26). Further, for $0.5 < \alpha < 0.8$ the participant would underestimate the probability of the most probable target and yield a strategy index between -0.26 and 0; whereas for $0.8 < \alpha < 1$ the participant would overestimate the probability of the most probable target and yield a strategy index between 0 and 0.63. Note that the strategy index increases monotonically with α for a fixed β .

Supplementary Figure 2 presents data from three representative participants based on this two-parameter model. In particular, we present the evolution of their strategy index across training blocks as a walk on the model surface. That is, we fitted the two-parameter model on the participants' response data per block and estimated the parameters α and β per participant and block. We then computed the participant strategy index as the difference between the participant strategy choice and the matching model. In particular, we observed that all participants started close to the random model ($\beta \approx 1$) and then deviated towards higher α and lower β values. However, the trajectory and end point of the individual participants varied and therefore yielded different strategy index values. That is, participant A showed $0.5 < \alpha < 0.8$ throughout most of the training blocks (i.e. underestimated the highly probable targets) while $\alpha \approx 0.8$ (i.e. close to matching) at the end of the training, yielding a negative strategy index. In contrast, participant B showed $\alpha \approx 0.8$ consistently across blocks and therefore yielded a strategy index close to 0 (i.e. matching). Finally, participant C overestimated the highly probable targets (i.e. $0.8 < \alpha < 1$) and yielded a higher positive strategy index (i.e. closer to maximization).

MRI data acquisition: Scanning was conducted using a 3T Philips Achieva MRI scanner with a 32-channel head coil. T1-weighted anatomical data (175 slices; $1 \times 1 \times 1 \text{ mm}^3$ resolution) were collected during the first scanning session. Resting-state echo-planar imaging (EPI) data (gradient echo-pulse sequences) were acquired in both scanning sessions (whole brain coverage; 180 volumes; TR=2s; TE=35ms; 32 slices; $2.5 \times 2.5 \times 4 \text{ mm}^3$ resolution; SENSE). The benefit of non-isotropic resolution is acquisition speed; that is, faster acquisition of fewer slices at higher in-plane resolution (keeping voxel volume constant and signal-to-noise ratio similar). This is advantageous for resting-state fMRI (rs-fMRI) that requires relatively high temporal resolution. We employed standard pipelines (i.e. SPM) that have been extensively used to model fMRI data at non-isotropic resolution. We employed a well-established volumetric

analysis (i.e. Group Independent Component Analysis-GICA) to investigate functional connectivity at rest that has been developed and validated on non-isotropic data⁴⁻⁸. Finally, a recent study⁹ has shown highly similar ICA results between isotropic and anisotropic datasets.

We collected rs-fMRI from three runs that each lasted for 6 minutes. Participants were instructed to keep their eyes open and maintain fixation to a white dot presented at the center of the screen. Diffusion Tensor Imaging (DTI) data were also collected in both scanning sessions and the acquisition consisted of 60 isotropically-distributed diffusion weighted directions ($b=1500 \text{ smm}^{-2}$; $TR=9.5\text{s}$; $TE=78\text{ms}$; 75 slices; $2\times 2\times 2 \text{ mm}^3$ resolution; SENSE) plus a single volume without diffusion weighting ($b=0 \text{ smm}^{-2}$, denoted as b_0). The DTI sequence was repeated twice during each session, once following the Anterior-to-Posterior phase-encoding direction and once the Posterior-to-Anterior direction. This acquisition scheme was implemented to allow correction of susceptibility-induced geometric distortions¹⁰.

DTI connectivity-based segmentation of striatum: Previous work across species^{11,12} has shown that dissociable cortical projections from anatomically-defined striatal subdivisions mediate distinct brain functions. To investigate learning-dependent changes in these corticostriatal connections, we defined the striatum (i.e. caudate and putamen) anatomically from the Automated Anatomical Labeling (AAL) atlas¹³. We then conducted a DTI connectivity-based segmentation to segment the striatum into finer subdivisions (i.e. segments) based on their whole-brain connectivity profile¹⁴.

We pre-processed and analyzed the DTI data in FSL 5.0.8 (FMRIB Software Library, <http://fsl.fmrib.ox.ac.uk/fsl/fslwiki/>). We first corrected the data for susceptibility distortions, eddy currents and motion artifacts (FSL topup and FSL eddy)¹⁵ and rotated the gradient directions (bvecs) to correct for the estimated motion rotation^{16,17}. We generated a distribution model in each voxel using *FSL BedpostX*¹⁸ with default parameters.

To simulate tracts from a seed defined in MNI space, we computed the transformation matrix from MNI to native space per participant (FSL flirt). We followed a 4-step registration procedure: (a) aligned the non-weighted diffusion volume (b0) of each session to their midspace and create a midspace-template (rigid-body)^{19,20}, (b) aligned the midspace-template to the anatomical (T1) scan (affine), (c) aligned the T1 image to the MNI template (affine) and (d) inverted and combined all the transformation matrices of the previous steps to obtain the MNI-to-native registration. The results of each step were visually inspected to ensure that the alignment was successful.

We then simulated tracts (i.e. probabilistic streamlines) starting from the seed area (i.e. striatum) to the rest of the brain (i.e. target area) using the *ProbtrackX* algorithm²¹. Following a hypothesis-free classification method²², we down-sampled the target area (AAL atlas excluding the seed: bilateral caudate and putamen) to 4x4x4 mm³ resolution. As the seed areas were in MNI space, we provided the MNI-to-native transformation matrix and used the *omatrix2* option to create a seed-by-target connectivity matrix (the *ProbtrackX* algorithm transforms the seed from MNI to native space and performs the probabilistic tractography simulation in native space; the results are then transformed back into MNI space). We used a mid-sagittal exclusion mask to prevent tracts from crossing hemispheres²¹ and length correction to account for the distance-from-the-seed bias towards shorter connections²². The parameters we used in *ProbtrackX* are: 5000 samples per voxel, 2000 steps per sample until conversion, 0.5mm step length, 0.2 curvature threshold, 0.01 volume fraction threshold and loopcheck enabled to prevent tracts from forming loops. We repeated this procedure for each hemisphere (**Supplementary Figure 3**).

This analysis generated a connectivity matrix from each voxel in the seed area to every voxel in the target area. Defining the seed in the MNI space guaranteed the same number of voxels in the seed across participants (after the data were transformed back from native to MNI

space), alleviating differences in individual brain size. Subsequently, we concatenated the connectivity matrices across participants and groups and correlated the connectivity values from and to each voxel in the seed; generating a seed-by-seed correlation matrix. We then performed k-means clustering on the correlation matrix for 2 to 8 classes (squared Euclidean distance). Lastly, we converted each class to a binary mask in MNI space to create the striatal segments and down-sampled them to the resting-state resolution (3x3x4 mm³) for further analysis.

To find the optimal number of clusters, we computed the mean silhouette value per clustering by averaging the values across voxels²³. The silhouette value shows how similar each voxel is to voxels of its class compared to voxels of other classes. Therefore, we selected the highest number of clusters that shows the maximum mean silhouette value averaged for the two hemispheres. This method resulted in 4 striatal segments per hemisphere (average silhouette value of 0.4) that corresponded to known anatomical subdivisions of the striatum (**Figure 3a, Supplementary Table 1**).

Resting-state data pre-processing: We pre-processed the resting-state data in SPM12.2 software package (<http://www.fil.ion.ucl.ac.uk/spm/software/spm12/>) following the optimized pipeline described in recent work⁵. We first processed the T1-weighted anatomical images by applying brain extraction and segmentation (SPM segment). From the segmented T1 we created a white matter (WM) mask and a cerebrospinal fluid (CSF) mask. For each resting-state run, we corrected the EPI data for slice scan timing (i.e. to remove time shifts in slice acquisition, SPM slice timing) and motion (least squares correction, SPM realign). We co-registered all EPI runs to the first run per participant (rigid body) and subsequently to the T1 image (rigid body, resliced to 1 x 1 x 1 mm³) and calculated the mean CSF and WM signal per volume (SPM coregister & reslice). We then aligned the T1 image to the MNI space (affine)

and applied the same transformation to the EPI data (SPM normalise). We resliced the aligned EPI data to $3 \times 3 \times 4 \text{ mm}^3$ resolution and applied spatial smoothing with a 5mm isotropic FWHM Gaussian kernel (SPM smooth). Finally, we despiked any secondary motion artifacts using the Brain Wavelet Toolbox²⁴, regressed out the signal from CSF and the motion parameters (translation, rotation and their squares and derivatives²⁵) and applied linear detrending²⁶. Note that the pipeline we followed⁵ does not include the global signal as a nuisance regressor, consistent with a recent review²⁷ suggesting that global signal regression may not be appropriate for comparisons between sessions and groups.

Independent Component Analysis (ICA): We used spatial GICA^{6,28} to extract participant- and session-specific hemodynamic source locations using the Group ICA fMRI Toolbox (GIFT) (<http://mialab.mrn.org/software/gift/>). Pre-processed EPI data from both groups (i.e. training, no-training control) from both sessions (i.e. Pre, Post) were included in the GICA. Following pre-processing of each run, the mean value per voxel was removed and dimensionality reduction was performed. We used the Minimum Description Length criteria (MDL)²⁹ to estimate the dimensionality and determine the number of independent components. We used a two-level dimensionality reduction procedure using Principal Component Analysis (PCA); first at the participant level and then at the group level. The ICA estimation (Infomax algorithm) was run 20 times and the component stability was estimated using ICASSO³⁰.

This procedure resulted in 22 spatially independent components. We then generated participant-specific spatial maps for each component using GICA3 back reconstruction⁴. Lastly, participant and group spatial maps were scaled to z maps for further analysis³¹. We then used a quantitative method, as described in previous work³², to remove components of non-neuronal origin. We first thresholded the group spatial maps at $z=1.0$ and calculated the spatial correlation of each component with CSF and grey matter (GM) probabilistic maps (as extracted

from the MNI template). We rejected any component with a spatial correlation of $R^2 > 0.025$ with CSF or of $R^2 < 0.025$ with GM. To supplement this method, we visually inspected all rejected components to verify that they were not of neuronal origin. This method resulted in 5 rejected components: 2 components had high spatial correlations with CSF and 3 components had low spatial correlations with GM.

We correlated the thresholded maps of the remaining components with known network templates and labeled each component based on its highest correlation value to these templates^{7,33}. We selected 7 components (**Figure 3b, Supplementary Table 2**) that showed high correlation with templates of cortical regions involved in executive, motor, visual and motivational networks^{11,12}.

To extract the resting-state time course for each cortical ICA-based component and DTI-based striatal segment, we used an autoregressive AR(1) model (SPM first-level analysis) on the pre-processed data before ICA to treat for serial correlations³⁴. Following the whole-brain modeling, we extracted the time course per voxel per component (SPM VOI extraction), as defined by participant-specific spatial maps thresholded at $z=2.576$ ($p=0.01$). We then applied a 5th order Butterworth band-pass filter, between 0.01 and 0.08 Hz to remove effects of scanner noise and physiological signals (respiration, heart beat)³⁵. In addition, we extracted the first eigenvariate across all voxels in each component to derive a single time course per component for subsequent connectivity analysis.

Graph analysis: To construct a functional connectivity matrix for each participant, we followed the same processing steps as for the extrinsic connectivity analysis. We extracted the first eigenvariate across all voxels in each AAL region (90 areas; excluding Cerebellum and Vermis) and constructed a 90x90 correlation matrix by correlating the time course of each AAL region with every other AAL region. We then standardized the correlation coefficients using

Fisher z-transform and averaged the z-values across the three rs-fMRI runs to derive a single functional connectivity matrix for each participant and session.

To construct a structural connectivity matrix for each participant, we simulated tracts (i.e. probabilistic streamlines) from each AAL area (i.e. seed mask) to any other AAL area (i.e. termination masks; excluding Cerebellum and Vermis) in native space using the *Probabilistic Tracking* algorithm (FSL ProbtrackX)²¹. The parameters we used in *ProbtrackX* are: 5000 samples per voxel, 2000 steps per sample until conversion, 0.5mm step length, 0.2 curvature threshold, 0.01 volume fraction threshold and loopcheck enabled to prevent tracts from forming loops. To control for differences in volume across seeds and participants, we normalized the tract count (i.e. the number of streamlines reaching area j when seeded from areas i) by the total number of tracts started from the seed region³⁶. Finally, we averaged the normalized tract count from area i to area j and from area j to area i to create a symmetric structural connectivity matrix for each participant and session.

We then constructed participant-specific binary graphs based on the connectivity matrices for each modality (i.e. rs-fMRI, DTI). We first generated the Minimum Spanning Tree³⁷ per matrix to create a connected graph for each participant and session. We then iteratively added the strongest edges irrespective of the sign (i.e. using the absolute functional connectivity value), until we reached a certain density level. Previous work in a similar-sized parcellation³⁸ has shown that density lower than 15% may result in sparse graphs and higher than 25% in graphs without small-world topology. Thus, we generated graphs at 20% density and then evaluated the stability of our findings in a range of density levels: from 10 to 30% in increments of 5. We used the Brain Connectivity Toolbox³⁹ to calculate graph metrics per participant and modality.

We note that the DTI and rs-fMRI metrics used in our graph analysis were derived by data pre-processed at native vs. standard space. In particular, DTI tractography is typically

performed in the native space to achieve best performance of the tracking algorithms²¹, whereas rs-fMRI data are typically normalized to a standard space (e.g. MNI) before computing functional connectivity⁵. Following previous studies, we analyzed the DTI data in native space, while the rs-fMRI data in standard space (i.e. data were normalized to MNI), as these data needed to be in a common space for group analysis across participants. While some recent studies recommend performing the rs-fMRI analysis in native space to minimize the effect of interpolation and improve localization^{40,41}, others have found no difference with and without the inclusion of the normalization step⁴². Further, our analysis approach makes it unlikely that these differences in interpolation between data types (i.e. rs-fMRI, DTI) have a significant effect on our results. First, we selected brain regions for both the rs-fMRI and DTI graph analysis based on the AAL parcellation, resulting in larger size brain regions. This makes it unlikely that small differences in the interpolation step would significantly affect the connectivity values estimated across all voxels in each brain region. Second, for the rs-fMRI data we computed the first eigenvariate when we extracted the time course per brain region and computed functional connectivity from these values. This step extracts the most representative time course from all the voxels in each brain region based on their common variance; therefore, it minimizes the effects of noise and interpolation⁴³. Third, for each imaging modality (i.e. rs-fMRI, DTI) we generated binary graphs and compared the connectivity values to select the strongest connections within-modality rather than comparing connectivity across modalities. That is, we created binary graphs at 20% density level by selecting the edges with the top 20% connectivity values, for each modality and session. We computed degree and clustering coefficient from these graphs per modality and used these metrics in the PLS regression to combine data from both modalities.

Partial Least Squares (PLS) modeling: control analyses: Results in the main text are presented for a network density of 20%. Here we show the robustness of these results in a range of densities (10%-30%) typically used in brain network analyses³⁸. We calculated degree and clustering for 10% to 30% density in increments of 5% per session (Pre, Post). We computed the difference between the two curves (Post minus Pre) for each metric (degree, clustering coefficient)⁴⁴ and performed the same PLS regression analysis as before. We tested for model significance using permutation testing (10,000 permutations) and then correlated the estimated PLS components and bootstrapped weights (1,000 samples) with the components and weights estimated for 20% density as shown in the main text. We found that the first PLS component across densities was significant compared to the null ($p=0.05$) and showed a high correlation with the PLS-1 component for 20% density ($r(19)=0.94$, $p<0.001$, $CI=[0.85, 0.98]$). Further, the predictor weights across densities showed a high correlation with the weights for 20% density ($r(46)=0.84$, $p<0.001$, $CI=[0.67, 0.93]$). PLS-2 across densities was not significant in comparison to the null model; however, it showed a high correlation with the PLS-2 component and its weights for 20% density (component: $r(19)=0.89$, $p<0.001$, $CI=[0.75, 0.95]$; weights: $r(46)=0.89$, $p<0.001$, $CI=[0.83, 0.94]$). Similarly, PLS-3 across densities was not significant compared to the null and showed weaker correlations with the PLS-3 component for 20% density (component: $r(19)=0.77$, $p<0.001$, $CI=[0.63, 0.88]$; weights: $r(46)=0.48$, $p<0.001$, $CI=[0.11, 0.71]$). We therefore restricted the main analysis to the first two components. **Supplementary Figure 6** summarizes the weights (combinations of nodes and metrics) for PLS-1 and PLS-2 for the average metrics (10% to 30% density).

Further, to test whether our findings generalize to other parcellation schemes than the AAL atlas, we created graphs at 20% density using the Shen⁴⁵ and Brainnetome⁴⁶ atlases that provide a finer whole brain parcellation. We selected nodes that corresponded to the same anatomical areas as the selected AAL nodes and performed a similar PLS regression analysis.

We found that both atlases yielded significant results (Shen: first three components; Brainnetome: first four components). Moreover, we found that the first two components for these atlases were highly similar to our results when using the AAL atlas (Shen: PLS-1: $r(19)=0.75$, $p<0.001$, $CI=[0.42, 0.92]$, PLS-2: $r(19)=0.83$, $p<0.001$, $CI=[0.53, 0.93]$; Brainnetome: PLS-1: $r(19)=0.73$, $p<0.001$, $CI=[0.44, 0.89]$, PLS-2: $r(19)=0.87$, $p<0.001$, $CI=[0.68, 0.94]$). Note that the Brainnetome atlas provides a parcellation of the striatum (i.e. ventral caudate, dorsal caudate, dorsolateral putamen and ventromedial putamen) that is comparable to our DTI-based segmentation (**Figure 3a**). Further, the significant predictors for PLS-1 were: a) degree change in right ventral caudate (rs-fMRI), left dorsal caudate (rs-fMRI), left ACC (DTI) and left postcentral (rs-fMRI); b) clustering change in right ventral caudate (DTI) and left postcentral (rs-fMRI); whereas for PLS-2 were: a) degree change in right MFG (DTI) and left postcentral (DTI); b) clustering change in left ACC (DTI), right dorsolateral putamen (rs-fMRI) and right ACC (rs-fMRI). Taken together, these findings suggest that our graph analysis is robust across parcellation schemes that segment the striatum at different scales, making it unlikely that our results were confounded by the selected parcellation atlas.

Finally, we tested whether our findings generalize to other graph metrics that relate to global and local integration. In particular, we tested: a) the average shortest path length (i.e. average number of a node's transitions via graph edges to any other node in the network) and betweenness centrality (i.e. number of shortest paths that traverse through a certain node) as measures of global integration^{47,48}, b) the local efficiency (i.e. how efficiently a node's neighbors communicate if this node is removed) as measure of local integration⁴⁹. These measures have been previously shown to relate to learning and brain plasticity⁵⁰⁻⁵². We conducted similar PLS regression analyses as for our main model (i.e. Model-1: degree and clustering coefficient) for the following models based on combinations between global and local integration metrics: a) Model-2: average shortest path length and clustering coefficient,

b) Model-3: average shortest path length and local efficiency, c) Model-4: degree and local efficiency, d) Model-5: betweenness centrality and clustering coefficient, e) Model-6: betweenness centrality and local efficiency. All models showed significant results when tested for 10,000 permutations (Model-2: first component, $p=0.010$; Model-3: first two components, $p=0.044$; Model-4: first three components, $p=0.012$; Model-5: first three components, $p=0.026$; Model-6: first component, $p=0.022$). Further, the first two components for these models were highly correlated to the components of the main model (Model-1) including degree and clustering coefficient (**Supplementary Table 5**). Thus, our findings showing that learning-dependent plasticity in corticostriatal networks predicts individual behaviour (i.e. decision strategy) are not limited only to selected measures of global or local integration.

Further, including all the above graph metrics in the same PLS model (Model-7: degree, average shortest path length, betweenness centrality, clustering coefficient and local efficiency), the model was significant for the first three PLS components compared to a null model ($p=0.045$, 10,000 permutations). In addition, the first two components for this model were highly correlated to the components of Model-1 (**Supplementary Table 5**), generalizing our results to a larger number of metrics that characterize whole-brain network connectivity.

No-training control experiment: Scanning for the no-training control experiment was conducted using a 3T MRI scanner with a 32-channel head coil. T1-weighted anatomical data (175 slices; $1 \times 1 \times 1 \text{ mm}^3$ resolution) were collected during the first scanning session. Resting-state EPI data (gradient echo-pulse sequences) were acquired in both scanning sessions with the same sequence as the one used in the training experiment (whole brain coverage; 180 volumes; TR=2s; TE=30ms; 36 slices; $2.5 \times 2.5 \times 4 \text{ mm}^3$ resolution; GRAPPA). We collected rs-fMRI from three runs that each lasted for 6 minutes. DTI data were also collected in both scanning sessions and the acquisition parameters were matched as closely as possible to the

training group: 60 isotropically-distributed diffusion weighted directions ($b=1500 \text{ smm}^{-2}$; $TR=8.9s$; $TE=91ms$; 72 slices; $2x2x2 \text{ mm}^3$ resolution; GRAPPA) plus a single volume without diffusion weighting ($b=0 \text{ smm}^{-2}$). The DTI sequence was repeated twice during each session, once following the Anterior-to-Posterior phase-encoding direction and once the Posterior-to-Anterior direction.

To ensure that the data quality was similar between the two groups (training vs. no-training control) that were tested using highly similar sequences and scanning parameters, we tested for differences related to a) head movement and b) spikes for the rs-fMRI data, and a) head movement and b) diffusion tensor model fit for the DTI data. For the rs-fMRI data, we calculated the maximum root mean square (rms) movement per run (based on x,y,z motion parameters estimated by SPM realign) and the maximum number of spikes per run (based on the Spike Percentage output of the Brain Wavelet toolbox²⁴). For the DTI data, we calculated the root mean square (rms) movement per session (based on *eddy*'s *restricted_movement_rms* output) and the sum of squared errors (sse) from diffusion tensor model fit¹⁸. No significant differences were observed between groups for head movement (rs-fMRI: $F(1,40)=0.31$, $p=0.578$, $\eta_p^2=0.008$; DTI: $F(1,40)=1.84$, $p=0.182$, $\eta_p^2=0.044$), number of spikes ($F(1,40)=1.19$, $p=0.283$, $\eta_p^2=0.029$) or diffusion tensor model fit for the seed areas, the whole brain and the white-matter ($F(1,40)=0.77$, $p=0.386$, $\eta_p^2=0.019$). Thus, these analyses suggest that it is unlikely that differences in connectivity between groups could be due to differences in data quality.

Supplementary References

1. Brainard, D. H. The Psychophysics Toolbox. *Spat. Vis.* **10**, 433–436 (1997).
2. Pelli, D. G. The VideoToolbox software for visual psychophysics: transforming numbers into movies. *Spat. Vis.* **10**, 437–442 (1997).
3. Wang, R., Shen, Y., Tino, P., Welchman, A. E. & Kourtzi, Z. Learning predictive statistics from temporal sequences: Dynamics and strategies. *J. Vis.* **17**, 1 (2017).
4. Calhoun, V. D., Adali, T., Pearlson, G. D. & Pekar, J. J. Functional neuroanatomy of visuo-spatial working memory in turner syndrome. *Hum. Brain Mapp.* **14**, 96–107 (2001).
5. Vergara, V. M., Mayer, A., Damaraju, E., Hutchison, K. & Calhoun, V. The effect of preprocessing pipelines in subject classification and detection of abnormal resting state functional network connectivity using group ICA. *Neuroimage* (2016). doi:10.1016/j.neuroimage.2016.03.038
6. Calhoun, V. D., Liu, J. & Adali, T. A review of group ICA for fMRI data and ICA for joint inference of imaging, genetic, and ERP data. *Neuroimage* **45**, 163–172 (2009).
7. Shirer, W. R., Ryali, S., Rykhlevskaia, E., Menon, V. & Greicius, M. D. Decoding Subject-Driven Cognitive States with Whole-Brain Connectivity Patterns. *Cereb. Cortex* **22**, 158–165 (2011).
8. Allen, E. A. *et al.* A baseline for the multivariate comparison of resting-state networks. *Front. Syst. Neurosci.* **5**, 2 (2011).
9. Chen, Z. & Calhoun, V. Effect of Spatial Smoothing on Task fMRI ICA and Functional Connectivity. *Front. Neurosci.* **12**, (2018).
10. Andersson, J. L. R., Skare, S. & Ashburner, J. How to correct susceptibility distortions in spin-echo echo-planar images: application to diffusion tensor imaging. *Neuroimage* **20**, 870–888 (2003).
11. Alexander, G. E., DeLong, M. R. & Strick, P. L. Parallel Organization of Functionally Segregated Circuits Linking Basal Ganglia and Cortex. *Annu. Rev. Neurosci.* **9**, 357–381 (1986).
12. Seger, C. A. The Involvement of Corticostriatal Loops in Learning Across Tasks, Species, and Methodologies. in *The basal ganglia IX* 25–39 (Springer, 2009). doi:10.1007/978-1-4419-0340-2_2
13. Tzourio-Mazoyer, N. *et al.* Automated Anatomical Labeling of Activations in SPM Using a Macroscopic Anatomical Parcellation of the MNI MRI Single-Subject Brain. *Neuroimage* **15**, 273–289 (2002).
14. Behrens, T. E. J. *et al.* Non-invasive mapping of connections between human thalamus and cortex using diffusion imaging. *Nat. Neurosci.* **6**, 750–757 (2003).
15. Andersson, J. L. R. & Sotiropoulos, S. N. An integrated approach to correction for off-resonance effects and subject movement in diffusion MR imaging. *Neuroimage* **125**, 1063–1078 (2016).
16. Jones, D. K. & Cercignani, M. Twenty-five pitfalls in the analysis of diffusion MRI data. *NMR Biomed.* **23**, 803–820 (2010).
17. Leemans, A. & Jones, D. K. The B -matrix must be rotated when correcting for subject motion in DTI data. *Magn. Reson. Med.* **61**, 1336–1349 (2009).
18. Behrens, T. E. J. *et al.* Characterization and propagation of uncertainty in diffusion-weighted MR imaging. *Magn. Reson. Med.* **50**, 1077–1088 (2003).
19. Smith, S. M., De Stefano, N., Jenkinson, M. & Matthews, P. M. Normalized accurate measurement of longitudinal brain change. *J. Comput. Assist. Tomogr.* **25**, 466–75 (2001).
20. Thomas, C. & Baker, C. I. Teaching an adult brain new tricks: A critical review of

- evidence for training-dependent structural plasticity in humans. *Neuroimage* **73**, 225–236 (2013).
21. Behrens, T. E. J., Johansen-Berg, H., Jbabdi, S., Rushworth, M. F. S. & Woolrich, M. W. Probabilistic diffusion tractography with multiple fibre orientations: What can we gain? *Neuroimage* **34**, 144–155 (2007).
 22. Tomassini, V. *et al.* Diffusion-Weighted Imaging Tractography-Based Parcellation of the Human Lateral Premotor Cortex Identifies Dorsal and Ventral Subregions with Anatomical and Functional Specializations. *J. Neurosci.* **27**, 10259–10269 (2007).
 23. Rousseeuw, P. J. Silhouettes: A graphical aid to the interpretation and validation of cluster analysis. *J. Comput. Appl. Math.* **20**, 53–65 (1987).
 24. Patel, A. X. *et al.* A wavelet method for modeling and despiking motion artifacts from resting-state fMRI time series. *Neuroimage* **95**, 287–304 (2014).
 25. Friston, K. J., Williams, S., Howard, R., Frackowiak, R. S. J. & Turner, R. Movement-Related effects in fMRI time-series. *Magn. Reson. Med.* **35**, 346–355 (1996).
 26. Van Dijk, K. R. a *et al.* Intrinsic Functional Connectivity As a Tool For Human Connectomics: Theory, Properties, and Optimization. *J. Neurophysiol.* **103**, 297–321 (2010).
 27. Murphy, K. & Fox, M. D. Towards a Consensus Regarding Global Signal Regression for Resting State Functional Connectivity MRI. *Neuroimage* 0–1 (2016). doi:10.1016/j.neuroimage.2016.11.052
 28. McKeown, M. J. *et al.* Analysis of fMRI data by blind separation into independent spatial components. *Hum. Brain Mapp.* **6**, 160–188 (1998).
 29. Rissanen, J. Modeling by shortest data description. *Automatica* **14**, 465–471 (1978).
 30. Himberg, J., Hyvärinen, A. & Esposito, F. Validating the independent components of neuroimaging time series via clustering and visualization. *Neuroimage* **22**, 1214–1222 (2004).
 31. Ma, L., Narayana, S., Robin, D. A., Fox, P. T. & Xiong, J. Changes occur in resting state network of motor system during 4weeks of motor skill learning. *Neuroimage* **58**, 226–233 (2011).
 32. Stevens, M. C., Kiehl, K. A., Pearlson, G. & Calhoun, V. D. Functional neural circuits for mental timekeeping. *Hum. Brain Mapp.* **28**, 394–408 (2007).
 33. Laird, A. R. *et al.* Behavioural interpretations of intrinsic connectivity networks. *J. Cogn. Neurosci.* **23**, 4022–37 (2011).
 34. Arbabshirani, M. R. *et al.* Impact of autocorrelation on functional connectivity. *Neuroimage* **102**, 294–308 (2014).
 35. Murphy, K., Birn, R. M. & Bandettini, P. A. Resting-state fMRI confounds and cleanup. *Neuroimage* **80**, 349–359 (2013).
 36. Johansen-Berg, H. & Rushworth, M. F. S. Using Diffusion Imaging to Study Human Connectional Anatomy. *Annu. Rev. Neurosci.* **32**, 75–94 (2009).
 37. Kruskal, J. B. On the Shortest Spanning Subtree of a Graph and the Traveling Salesman Problem. *Proc. Am. Math. Soc.* **7**, 48 (1956).
 38. Bassett, D. S. *et al.* Hierarchical Organization of Human Cortical Networks in Health and Schizophrenia. *J. Neurosci.* **28**, 9239–9248 (2008).
 39. Rubinov, M. & Sporns, O. Complex network measures of brain connectivity: uses and interpretations. *Neuroimage* **52**, 1059–1069 (2010).
 40. Magalhães, R. *et al.* The impact of normalization and segmentation on resting state brain networks. *Brain Connect.* **5**, 1–28 (2015).
 41. Razlighi, Q. R. *et al.* Unilateral disruptions in the default network with aging in native space. *Brain Behav.* **4**, 143–157 (2014).
 42. van den Heuvel, M. P., Stam, C. J., Boersma, M. & Hulshoff Pol, H. E. Small-world

- and scale-free organization of voxel-based resting-state functional connectivity in the human brain. *Neuroimage* **43**, 528–539 (2008).
43. Friston, K. J., Rotshtein, P., Geng, J. J., Sterzer, P. & Henson, R. N. A critique of functional localisers. *Neuroimage* **30**, 1077–1087 (2006).
 44. Bassett, D. S., Nelson, B. G., Mueller, B. A., Camchong, J. & Lim, K. O. Altered resting state complexity in schizophrenia. *Neuroimage* **59**, 2196–2207 (2012).
 45. Shen, X., Tokoglu, F., Papademetris, X. & Constable, R. T. Groupwise whole-brain parcellation from resting-state fMRI data for network node identification. *Neuroimage* **82**, 403–415 (2013).
 46. Fan, L. *et al.* The Human Brainnetome Atlas: A New Brain Atlas Based on Connectional Architecture. *Cereb. Cortex* **26**, 3508–3526 (2016).
 47. Watts, D. J. & Strogatz, S. H. Collective dynamics of ‘small-world’ networks. *Nature* **393**, 440–442 (1998).
 48. Freeman, L. C. A Set of Measures of Centrality Based on Betweenness. *Sociometry* **40**, 35 (1977).
 49. Latora, V. & Marchiori, M. Efficient Behaviour of Small-World Networks. 3–6 (2001). doi:10.1103/PhysRevLett.87.198701
 50. Román, F. J. *et al.* Enhanced structural connectivity within a brain sub-network supporting working memory and engagement processes after cognitive training. *Neurobiol. Learn. Mem.* **141**, 33–43 (2017).
 51. Heitger, M. H. *et al.* Motor learning-induced changes in functional brain connectivity as revealed by means of graph-theoretical network analysis. *Neuroimage* **61**, 633–650 (2012).
 52. Soto, F. A., Bassett, D. S. & Ashby, F. G. Dissociable changes in functional network topology underlie early category learning and development of automaticity. *Neuroimage* **141**, 220–241 (2016).

Supplementary Tables

Supplementary Table 1: Striatal segments. Four striatal segments for each hemisphere were estimated by a DTI connectivity-based and hypothesis-free classification method. The size of the segments and the MNI coordinates of their center of gravity are shown.

Hemisphere	Name	voxels	Center of gravity		
			x	y	z
Left	ventral striatum	102	-13	13	-9
	caudate head, anterior putamen	117	-16	14	-1
	caudate body/tail	120	-16	7	13
	posterior putamen	208	-27	-1	5
Right	ventral striatum	99	14	13	-8
	caudate head, anterior putamen	126	17	15	-1
	caudate body/tail	129	14	6	15
	posterior putamen	197	27	1	4

Supplementary Table 2: ICA components. Clusters within the 7 selected components are extracted from the group maps ($z=1.96$, $p=0.05$) and are organized into known functional groups^{7,33}. The table shows the number of voxels within each cluster (clusters smaller than 20 voxels are not included), the MNI coordinates, the label of the corresponding AAL area and the t-statistic of the peak voxel.

Network	Component	Cluster	voxels	x	y	z	t-value
Executive	CP_9 (RCEN)	R MFG	718	39	23	50	3.87
		R IPL	477	48	-49	54	4.64
		L Cerebellum	39	-36	-70	-42	2.61
		R Cingulate	38	3	35	38	3.01
		R MTG	27	66	-25	-10	2.23
	CP_14 (LCEN)	L IFG triangular	510	-51	17	30	4.55
		L IPL	413	-33	-70	50	3.81
		L MFG	55	-27	17	58	2.8
		L MTG	47	-60	-49	-10	2.46
		L SFG medial	25	-3	29	42	2.71
Motor	CP_4 (Sensorimotor)	R SMA	853	0	-22	58	3.92
	CP_5 (Lateral Motor)	R Postcentral	368	51	-25	54	3.55
		L Postcentral	330	-51	-31	54	3.8
Visual	CP_2 (Secondary)	R MOG	726	33	-82	22	3.42
		L MOG	406	-24	-88	22	2.88
	CP_12 (Early)	R Calcarine	606	12	-97	-2	3.39
Motivational	CP_15 (ACN)	R ACC	620	0	44	-2	4.38

Supplementary Table 3: Intrinsic and extrinsic connectivity correlations with strategy index. Semipartial Pearson skipped correlations are reported for (a) intrinsic connectivity change (post minus pre-training) and (b) extrinsic connectivity change with strategy index for frequency and context-based statistics. Significant correlations are determined based on bootstrapped confidence intervals (CI) and denoted in bold. The r-value and 95% CI are shown for each statistical test (n=21).

a. Intrinsic connectivity analysis

Network	frequency statistics		context-based statistics	
	r	CI	r	CI
ACN	0.12	[-0.32, 0.51]	0.35	[0.04, 0.63]
RCEN	-0.17	[-0.61, 0.33]	-0.16	[-0.57, 0.33]
LCEN	-0.01	[-0.39, 0.41]	0.42	[0.01, 0.68]
Secondary Visual	-0.09	[-0.43, 0.29]	-0.49	[-0.74, -0.10]
Early Visual	-0.32	[-0.73, 0.16]	-0.03	[-0.44, 0.40]
Sensorimotor	0.20	[-0.13, 0.53]	0.23	[-0.22, 0.59]
Lateral Motor	0.77	[0.60, 0.89]	-0.07	[-0.50, 0.39]

b. Extrinsic connectivity analysis

Corticostriatal pathways	frequency statistics		context-based statistics	
	r	CI	r	CI
ACN - right ventral striatum	-0.09	[-0.45, 0.28]	-0.15	[-0.43, 0.12]
ACN - left ventral striatum	-0.31	[-0.65, 0.12]	-0.14	[-0.53, 0.27]
RCEN - right caudate head, anterior putamen	-0.05	[-0.40, 0.36]	0.13	[-0.26, 0.42]
RCEN - left caudate head, anterior putamen	0.34	[-0.03, 0.66]	-0.14	[-0.41, 0.10]
LCEN - right caudate head, anterior putamen	0.17	[-0.31, 0.52]	0.22	[-0.19, 0.52]
LCEN - left caudate head, anterior putamen	0.03	[-0.34, 0.40]	0.01	[-0.35, 0.33]
Secondary Visual - right caudate body/tail	0.15	[-0.38, 0.57]	0.38	[-0.09, 0.72]
Secondary Visual - left caudate body/tail	0.19	[-0.25, 0.56]	0.21	[-0.28, 0.58]
Early Visual - right caudate body/tail	-0.04	[-0.50, 0.41]	0.05	[-0.41, 0.45]
Early Visual - left caudate body/tail	-0.19	[-0.60, 0.25]	-0.46	[-0.83, -0.13]
Sensorimotor - right posterior putamen	-0.14	[-0.49, 0.26]	0	[-0.35, 0.35]
Sensorimotor - left posterior putamen	0.01	[-0.55, 0.45]	0.03	[-0.37, 0.43]
Lateral Motor - right posterior putamen	0.51	[0.20, 0.74]	-0.19	[-0.59, 0.29]
Lateral Motor - left posterior putamen	0.13	[-0.41, 0.65]	0.03	[-0.50, 0.46]

Supplementary Table 4: PLS weights of the first two components: for (a) predictors and (b) response variables. Asterisks denote significant weights ($|z| > 2.576$, $p = 0.01$).

a. Weights for predictors

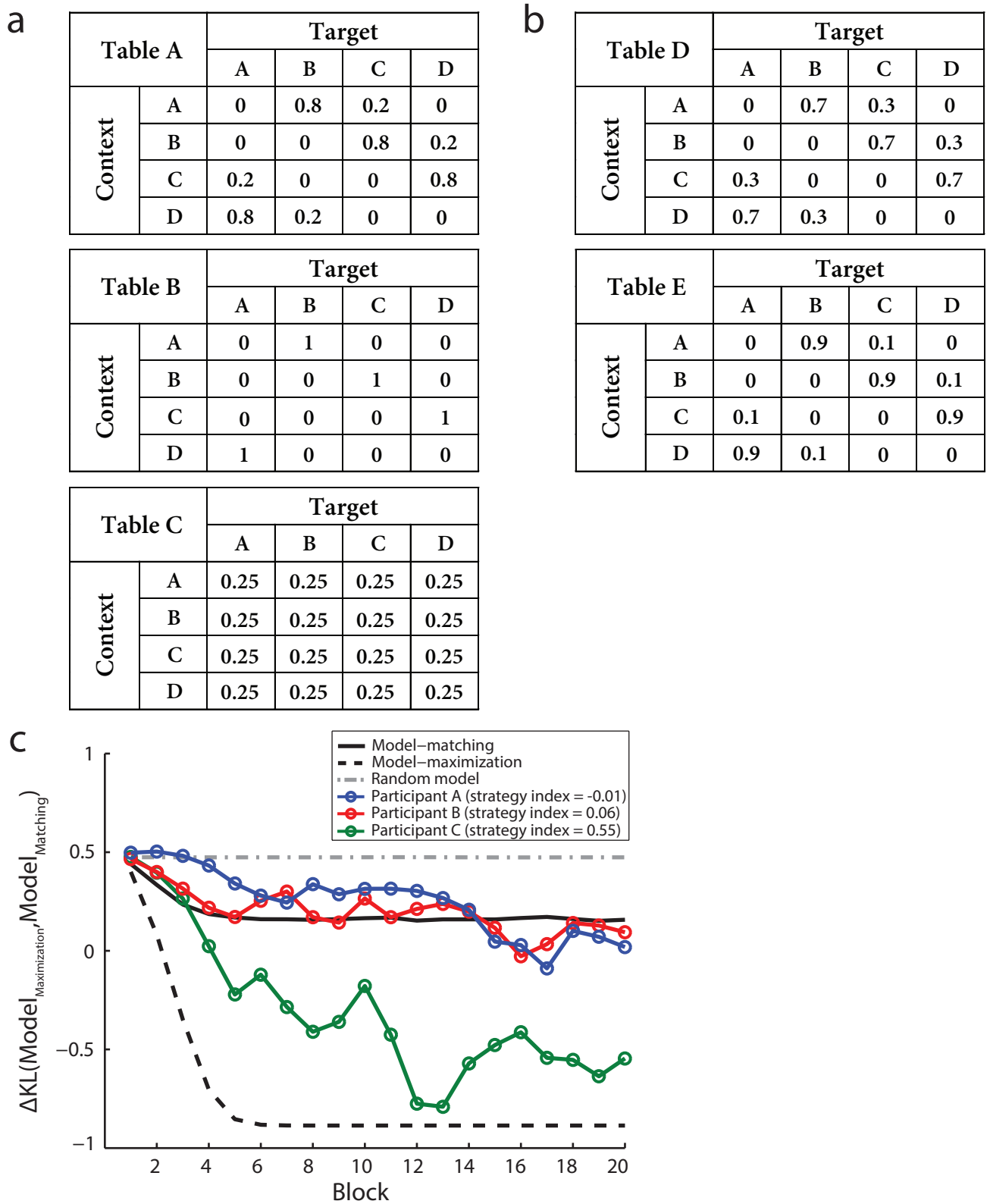
Node	Graph metric	PLS-1		PLS-2	
		rs-fMRI	DTI	rs-fMRI	DTI
L Caudate	Degree	1.79	-0.97	0.64	-2.84*
L Caudate	Clustering	1.18	1.05	-0.22	3.99*
R Caudate	Degree	2.30	-0.89	0.77	3.21*
R Caudate	Clustering	2.07	-0.10	0.03	-0.66
L Putamen	Degree	1.78	4.60*	1.38	-0.67
L Putamen	Clustering	0.29	-2.13	0.96	1.37
R Putamen	Degree	1.35	-2.06	0.31	0.34
R Putamen	Clustering	-0.40	-0.03	1.24	-0.27
R MFG	Degree	0.41	-0.22	0.39	2.67*
R MFG	Clustering	-1.92	-1.94	-0.49	-0.49
L IFG triangular	Degree	2.83*	1.50	0.11	1.24
L IFG triangular	Clustering	1.72	2.05	-0.57	1.32
L Postcentral	Degree	-1.86	-2.01	-1.69	-0.90
L Postcentral	Clustering	0.20	2.66*	-1.38	-0.44
R Postcentral	Degree	-0.74	0.15	-1.11	-0.69
R Postcentral	Clustering	-1.15	-1.71	-1.24	0.65
L Calcarine	Degree	-0.39	1.46	-0.23	-1.64
L Calcarine	Clustering	0.95	0.50	1.96	0.64
R Calcarine	Degree	0.40	3.58*	-0.67	0.02
R Calcarine	Clustering	-1.04	-1.67	2.18	-0.95
L ACC	Degree	0.39	-0.27	1.38	3.67*
L ACC	Clustering	0.34	-0.52	2.84*	1.12
R ACC	Degree	-0.18	2.16	2.55	1.21
R ACC	Clustering	-0.56	-3.45*	1.44	-0.30

b. Weights for response variables

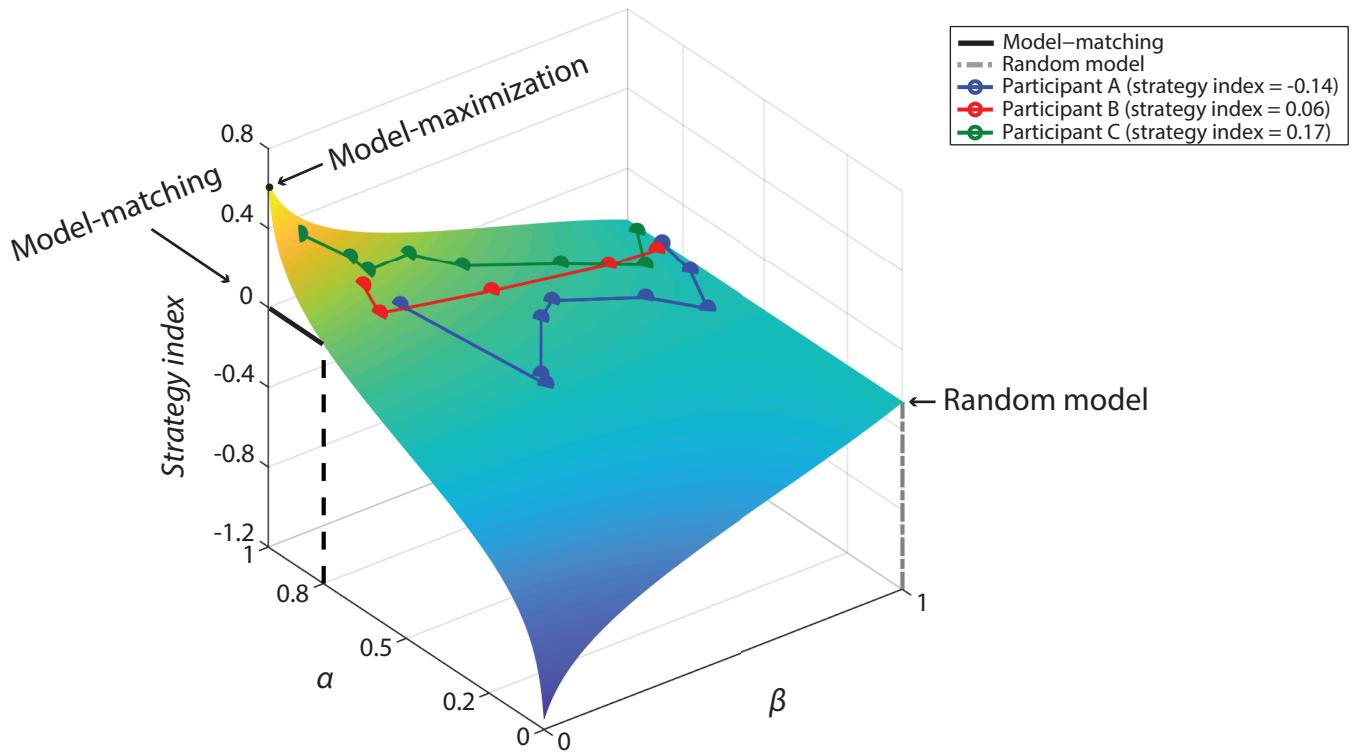
Behaviour	PLS-1	PLS-2
Strategy 0	-2.85*	2.01
Strategy 1&2	3.28*	2.47

Supplementary Table 5: PLS results across graph metrics. Pearson correlation of the first two PLS components between models (Model-1 is the reference model for the comparisons).

Model comparison	PLS-1	PLS-2
Model-2 vs. Model-1	$r=0.94$, CI=[0.81, 0.98]	$r=0.89$, CI=[0.75, 0.95]
Model-3 vs. Model-1	$r=0.88$, CI=[0.58, 0.97]	$r=0.86$, CI=[0.66, 0.96]
Model-4 vs. Model-1	$r=0.99$, CI=[0.96, 0.99]	$r=0.98$, CI=[0.94, 0.99]
Model-5 vs. Model-1	$r=0.95$, CI=[0.90, 0.98]	$r=0.93$, CI=[0.82, 0.97]
Model-6 vs. Model-1	$r=0.92$, CI=[0.80, 0.97]	$r=0.89$, CI=[0.73, 0.97]
Model-7 vs. Model-1	$r=0.98$, CI=[0.92, 0.99]	$r=0.97$, CI=[0.90, 0.99]

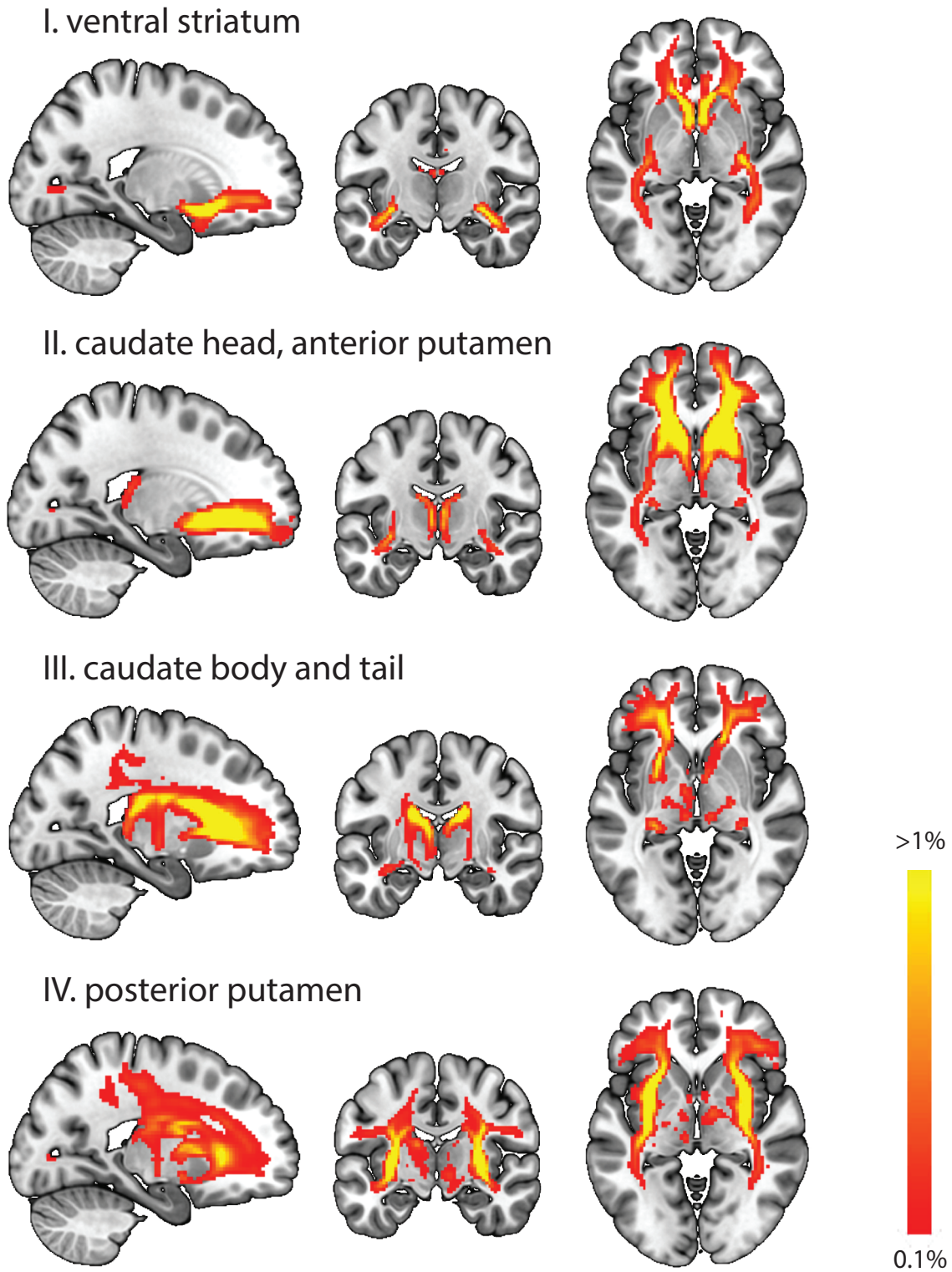


Supplementary Figure 1: Examples of participant responses for level-1 sequences. (a) Response tables for model-matching (Table A), model-maximization (Table B) and a random model (i.e. equal responses to all context-target contingencies; Table C). (b) Table D shows example responses for underestimating the probability of the most probable contingency (i.e. responses between random and model-matching). Table E shows example responses for overestimating the probability of the most probable contingency (i.e. responses between model-matching and model-maximization). (c) Participant strategy choice across training blocks for three representative participants (blue: participant A; red: participant B; green: participant C) against the three models (solid black line: model-matching; dashed black line: model-maximization; dashed grey line: random model). We computed the strategy index as the integral between the values of participant strategy choice and the model-matching across blocks.



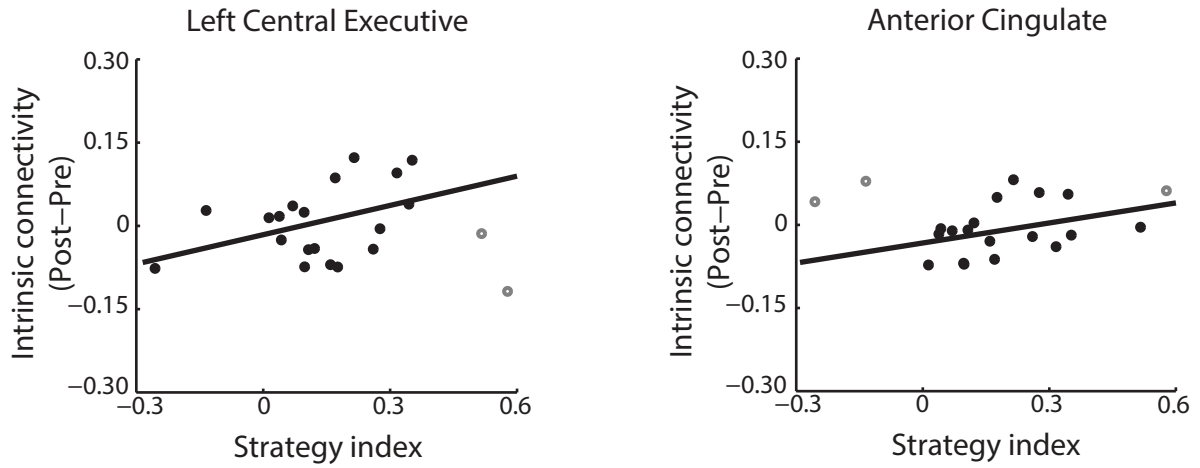
Supplementary Figure 2: Two-parameter model of participant response distribution.

The surface of a two-parameter model depicted here describes the strategy index of a virtual participant as a function of α and β ($P = \beta [0.25, 0.25, 0.25, 0.25] + (1-\beta) [1-\alpha, \alpha, 0, 0]$). α describes participant preference for the more over the less probable target: (i) $\alpha=1$ indicates maximization, (ii) $\alpha=0.8$ indicates matching, (iii) $\alpha=0.5$ indicates equal responses to the two possible targets, (iv) $\alpha<0.5$ indicates participant preference of the less probable target. β describes participant preference for the random model: (i) $\beta=1$ indicates random model of responses (i.e. equal responses for all targets), (ii) $\beta=0$ indicates no random responses (i.e. the model is described by the probabilities of the two probable targets). Colder colours (e.g. blue) denote lower strategy index values, whereas warmer colors (e.g. yellow) denote higher strategy index values. Individual data of three representative participants are displayed as walks on the surface (blue: participant A; red: participant B; green: participant C). Individual data points start from the right (i.e. $\beta \approx 1$) and deviate towards the left of the surface (i.e. $\beta \approx 0$) showing three distinct behaviours: participant A underestimates the highly probable targets (i.e. negative strategy index close to matching), participant B matches the target distribution (i.e. zero strategy index close to matching) and participant C overestimates the highly probable targets (i.e. positive strategy index close to maximization).



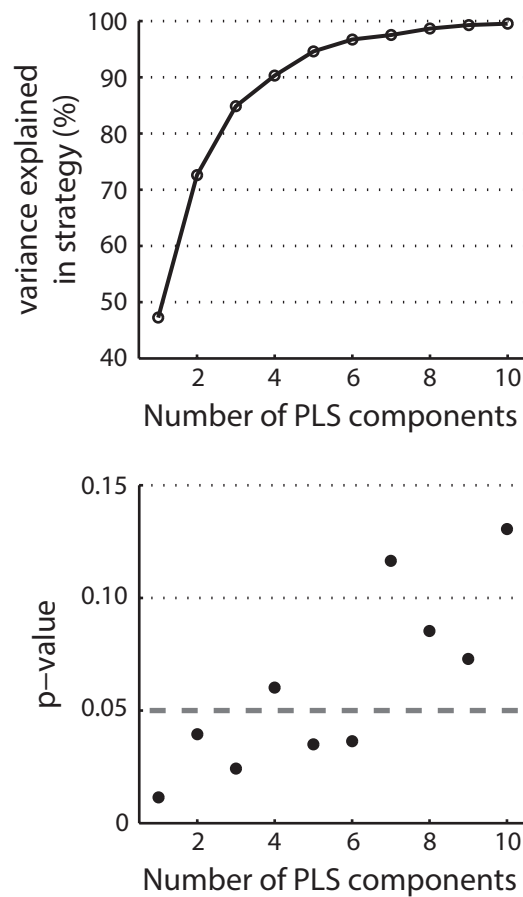
Supplementary Figure 3: DTI tractography for striatal segmentation. Striatal segments were estimated using a DTI connectivity-based and hypothesis-free classification method. Connection probability maps are displayed for each segment on the MNI template (neurological convention: left is left). Maps are thresholded at 0.1% of total tracts and averaged across groups and sessions. Whole brain tractography was computed separately for the left and right hemisphere and the maps were combined for visualization purposes ($x=-20$, $y=-12$, $z=-2$).

Context-based statistics



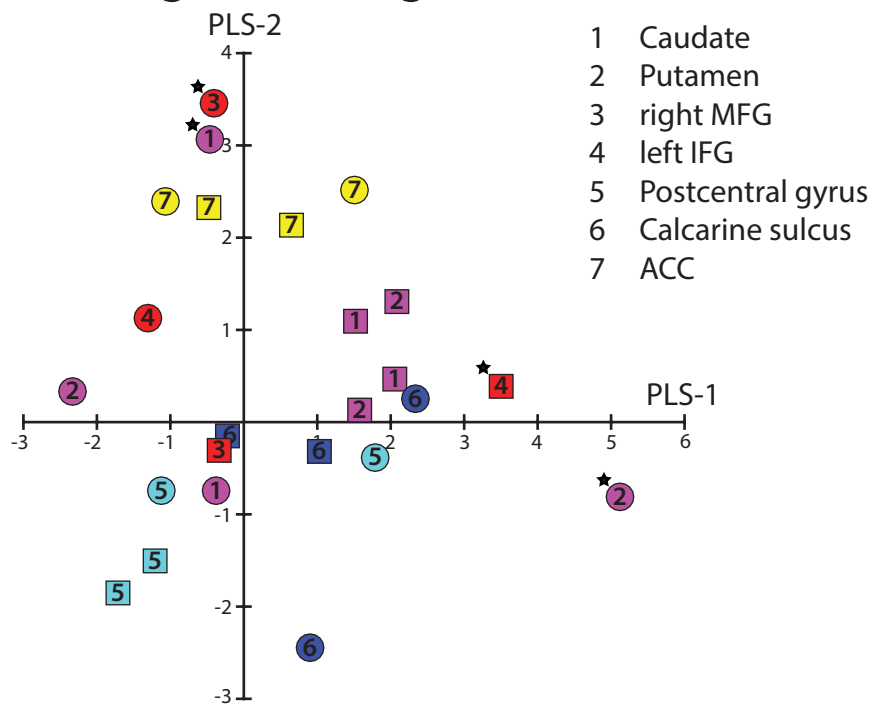
Supplementary Figure 4: Intrinsic connectivity analysis – supplementary results.

Skipped Pearson correlations (two-sided, $n=21$) showed a marginally significant relationship of intrinsic connectivity change (post- minus pre-training) in the Left Central Executive (LCEN) and Anterior Cingulate (ACN) networks with strategy index for frequency statistics (LCEN: $r(19)=0.42$, $p=0.059$, $CI=[0.01, 0.68]$; ACN: $r(19)=0.35$, $p=0.121$, $CI=[0.04, 0.63]$). Open circles in the correlation plots denote outliers as detected by the Robust Correlation Toolbox. Intrinsic connectivity was positive for all participants and sessions (pre-training, post-training); therefore, the sign of the change (Post minus Pre) indicates an increase (if positive) or a decrease (if negative) in the connectivity. In all but 5 cases (3 for posterior putamen - Lateral Motor connectivity; 2 for caudate body/tail - Early Visual connectivity) extrinsic connectivity change (Post minus Pre) had the same sign as the absolute connectivity change ($|Post|$ minus $|Pre|$). Therefore, we interpret these correlations based on the change of the actual connectivity values (that is, $Post > Pre$ is interpreted as increased connectivity). Performing the extrinsic connectivity analysis using the absolute connectivity change ($|Post|$ minus $|Pre|$) showed similar results. That is, we found a) increased connectivity between the right posterior putamen and the Lateral Motor network correlated positively with strategy index for frequency statistics ($r(16)=0.62$, $p=0.006$, $CI=[0.38, 0.79]$), b) increased connectivity between the left body/tail of caudate and the Early Visual network correlated negatively with strategy index for context-based statistics ($r(16)=-0.38$, $p=0.120$, $CI=[-0.74, -0.02]$).

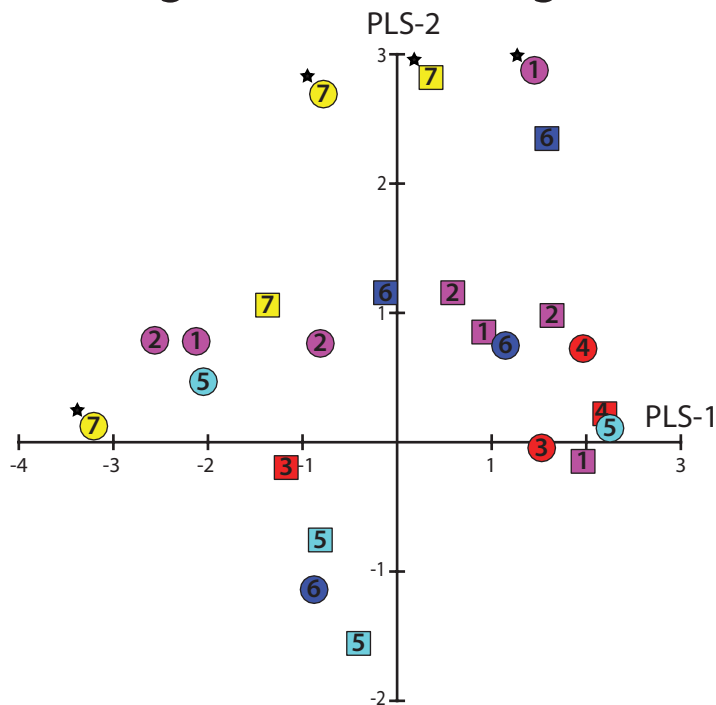


Supplementary Figure 5: Goodness of fit of PLS regression. Top panel shows variance explained in the response variables as a function of PLS components. Bottom panel shows the significance of the PLS model as a function of PLS components. Significance was determined by permutation testing (10,000 permutations); p-values below 0.05 indicate significant results.

a. PLS weights for degree

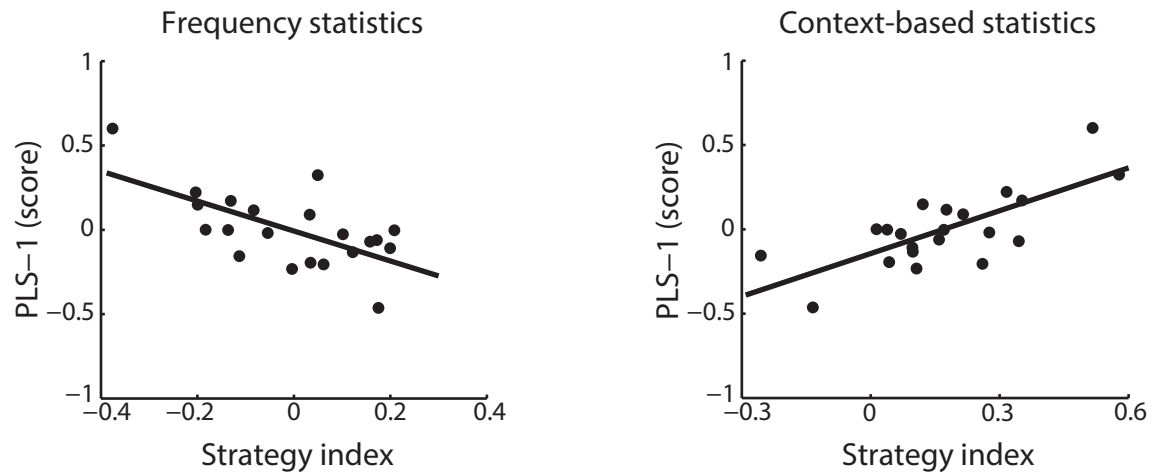


b. PLS weights for clustering coefficient

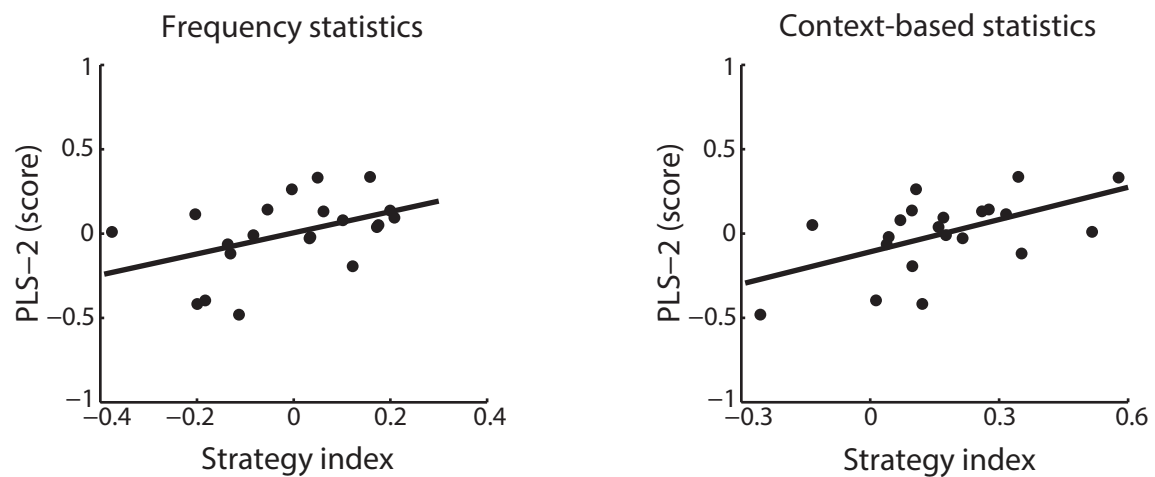


Supplementary Figure 6: PLS results across a range of density levels (from 10% to 30%). Scatterplot of PLS-1 and PLS-2 weights for change (i.e. post- minus pre-training) in (a) degree and (b) clustering coefficient. PLS predictor weights for each selected node are indicated by symbols separately for DTI (circles) and rs-fMRI (squares) data. The colour of the symbols corresponds to nodes in cortico-striatal circuits (**Figure 5**): caudate and putamen (magenta), right MFG and left IFG (red), postcentral gyrus (cyan), calcarine sulcus (blue), and ACC (yellow). PLS predictor weights with $|z| > 2.576$ ($p = 0.01$) are marked by an asterisk to denote significant predictors for the respective PLS component.

a. PLS-1 component



b. PLS-2 component



Supplementary Figure 7: PLS components related to strategy index. Illustration of the first two PLS components in relation to strategy index for frequency and context-based statistics (n=21). (a) Scatterplot of PLS-1 score with strategy index showing opposite patterns for frequency vs. context-based statistics. (b) Scatterplot of PLS-2 score with strategy index showing a similar pattern for frequency and context-based statistics. Note that the scatterplots between PLS components and strategy index are shown here for illustration purposes only. No further statistics were conducted to avoid circularity, as these two PLS components were shown to be significant predictors of the strategy index (**Figure 7a, Supplementary Table 4b**).