

Quantifying the compressibility of the human brain

Nicholas J. Weaver^{1,2,3}, Joshua I. Faskowitz⁴, Richard F. Betzel^{5,6} & Christopher W. Lynn^{1,2,7,*}

¹*Department of Physics, Yale University, New Haven, CT 06520, USA*

²*Quantitative Biology Institute, Yale University, New Haven, CT 06520, USA*

³*Program in Physical and Engineering Biology, Yale University, New Haven, CT 06520, USA*

⁴*Department of Psychological and Brain Sciences, Indiana University, Bloomington, IN 47405, USA*

⁵*Department of Neuroscience, University of Minnesota, Minneapolis, MN 55455, USA*

⁶*Masonic Institute for the Developing Brain, University of Minnesota, Minneapolis, MN 55414, USA*

⁷*Wu Tsai Institute, Yale University, New Haven, CT 06520, USA*

**Corresponding author: christopher.lynn@yale.edu*

Abstract

In the human brain, the allowed patterns of activity are constrained by the correlations between brain regions. Yet it remains unclear which correlations—and how many—are needed to predict large-scale neural activity. Here, we present an information-theoretic framework to identify the most important correlations, which provide the most accurate predictions of neural states. Applying our framework to cortical activity in humans, we discover that the vast majority of variance in activity is explained by a small number of correlations. This means that the brain is highly compressible: only a sparse network of correlations is needed to predict large-scale activity. We find that this compressibility is strikingly consistent across different individuals and cognitive tasks, and that, counterintuitively, the most important correlations are not necessarily the strongest. Together, these results suggest that nearly all correlations are not needed to predict neural activity, and we provide the tools to uncover the key correlations that are.

Main

Non-invasive neuroimaging has revolutionized our ability to record large-scale activity in the human brain.¹ Individual regions, each reflecting the coarse-grained activity of a large population of neurons, interact to produce brain-wide activity that is responsible for learning, memory, perception, planning, and information processing.²⁻⁶ The functional networks formed by these regions (nodes) and the correlations between them (edges) have provided key insights into human cognition.⁷⁻⁹ Variations in these networks are linked to cognitive disorders, human development, and environmental factors.¹⁰⁻¹⁵ However, despite their crucial role in understanding cognition and disease, there remain basic questions about what correlations reveal about the nature of neural activity itself. Among all pairs of brain regions, which correlations are most important for predicting large-scale activity? And how many correlations do we need to construct an accurate model of the human brain?

Answering these questions is a problem of compression. The human brain is strongly correlated, with activity only exploring a small subset of all possible states.¹⁶⁻¹⁸ This suggests that a dense network of correlations is needed to constrain the neural activity. Alternatively, a small number of correlations may combine to have a large impact on the brain as a whole. In this simplified scenario, one would only need a sparse network of important correlations to predict the rest. If such a network exists, then the brain is highly compressible.

Information theory provides the foundation to make this intuition concrete.^{19,20} Using the maximum entropy principle, one can map a network of correlations to predictions about neural activity.^{19,21} For a given number of correlations, the optimal network (which yields the most accurate predictions) provides the best compression of the system.²² We show that this leads to a direct generalization of maximum entropy known as the minimax entropy principle,²³⁻²⁶ which we solve to identify the optimal networks of correlations. Applying our framework to cortex-wide activity in healthy human subjects, we find that only a small number of correlations are needed to predict large-scale patterns of activity. We therefore discover that the brain is highly compressible.

Constraining activity with a network of correlations

Correlations tell us which regions are likely to activate together; once combined into a network they place complex constraints on the allowed patterns of activity. Consider N brain regions with collective activity defined by the vector $\mathbf{x} = \{x_i\}$, where x_i reflects the activity of region i . Any distribution over the states $P(\mathbf{x})$ defines a space of possible activity patterns. The size of this space—that is, our uncertainty about the neural states $\mathbf{x}$ —is quantified by the entropy $S(P)$. In the simplest case, suppose we only measure the variance in activity σ_i^2 of each region in an experiment. The most unbiased distribution over states is composed of independent Gaussians,^{21,27} and our uncertainty about the neural activity is defined by the independent entropy $S_{\text{ind}} = \frac{1}{2} \sum_i \log \sigma_i^2 + \frac{N}{2} \log(2\pi e)$. At the other extreme, in addition to the individual variances, consider the full set of covariances between regions Σ_{ij} . By constraining the activity, these correlations reduce our uncertainty to the entropy of a fully-connected Gaussian $S_{\text{tot}} = \frac{1}{2} \log |\Sigma| + \frac{N}{2} \log(2\pi e)$, where $|\Sigma|$ is the determinant of the covariance matrix.^{21,27}

Between these two extremes, a subset of the covariances can be represented as a network G , with each edge (ij) reflecting a covariance Σ_{ij} included in our constraints.^{7,8} Given such a network, the most unbiased prediction for the distribution over states is the one with maximum entropy,^{21,27}

$$P_G(\mathbf{x}) = \sqrt{\frac{|J|}{(2\pi)^N}} \exp\left(-\frac{1}{2} \mathbf{x}^T J \mathbf{x}\right), \quad (1)$$

where J is the precision matrix. In practice, for each covariance in the network G , one must compute the precision J_{ij} such that the model covariance $(J^{-1})_{ij}$ matches Σ_{ij} . This is equivalent to a Gaussian graphical model (GGM),^{28–31} for which there exist efficient inference algorithms (Methods).^{32,33} Thus, for a network of covariances G , our uncertainty about the collective activity is defined by the entropy $S_G = -\frac{1}{2} \log |J| + \frac{N}{2} \log(2\pi e)$.

Consider a minimal example of four brain regions (Fig. 1a-b). With no covariances, the regions behave independently with high uncertainty S_{ind} (Fig. 1c, *left*). With all of the covariances, we necessarily capture all of the correlation structure, reducing our uncertainty to S_{tot} (Fig. 1c,

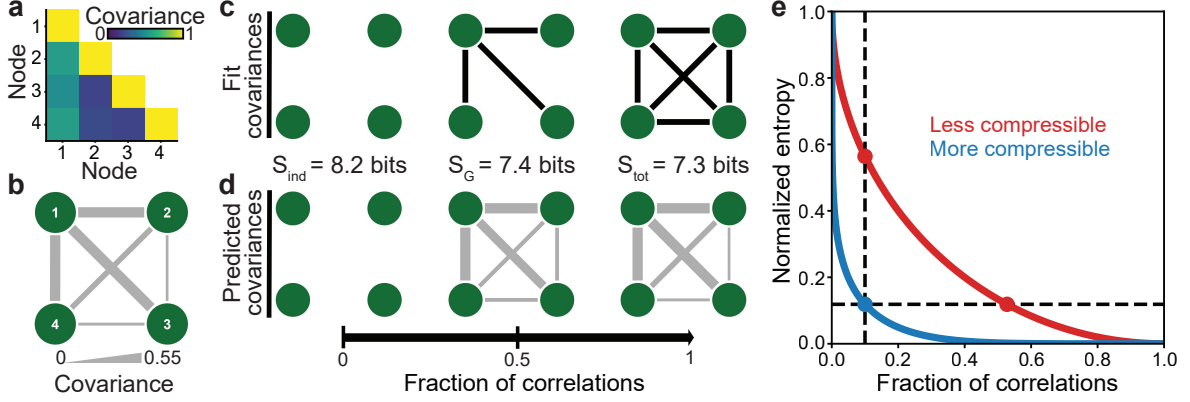


Fig. 1 | Quantifying uncertainty given a network of correlations. **a-b**, Covariance matrix (**a**) and network representation (**b**) for a minimal system of four regions with z-scored activity (Methods). **c-d**, Networks of correlations G (**c**) and the covariances predicted by the corresponding maximum entropy models P_G (**d**). The network formed by the three strongest covariances produces accurate predictions for all covariances (middle). **e**, Illustration of the normalized entropy $\tilde{S}$ versus the fraction of correlations (or the density of the network G) for neural activity that is either more compressible (blue) or less compressible (red). In a more compressible system, one can achieve lower uncertainty for a given number of correlations (vertical line), and one can find a sparser network (with fewer correlations) to achieve a given level of uncertainty (horizontal line).

right). Between these limits, one might hope to find a small subset of covariances that produces a large reduction in uncertainty; that is, a good compression. Indeed, by constraining the three strongest covariances, the maximum entropy model accurately predicts the remaining indirect correlations (Fig. 1d), and the entropy S_G nearly reduces to that of the fully-connected network S_{tot} .

We therefore arrive at the following picture: As one includes more correlations in a network, the uncertainty about neural activity decreases, resulting in a hierarchy of entropy $S_{\text{ind}} \geq S_G \geq S_{\text{tot}}$. To compare across different datasets, we define the normalized entropy,

$$\tilde{S}_G = \frac{S_G - S_{\text{tot}}}{S_{\text{ind}} - S_{\text{tot}}}, \quad (2)$$

which quantifies the progress from no correlations ($\tilde{S}_G = 1$) to a network that explains all of the structure in the neural activity ($\tilde{S}_G = 0$). If the brain is easier to compress, then for a given number of correlations, one should be able to find a network that reduces our uncertainty $\tilde{S}_G$ (Fig. 1e, *vertical line*). Equivalently, to achieve a given level of uncertainty, one should be able to find a network

with fewer correlations (Fig. 1e, *horizontal line*). Thus, in order to quantify the compressibility of the human brain, we must first identify the optimal networks of correlations.

Optimally compressing neural activity

For a given number of correlations, we seek the network G that minimizes the entropy S_G . This is an instance of the “minimax entropy” principle, which has recently been proposed to construct optimal models of complex systems, but has yet to be applied to human neural activity.^{22–26} Notably, since P_G is a maximum entropy model [Eq. (1)], we show that minimizing the entropy S_G is equivalent to minimizing the KL divergence between the model and the data (Methods). Thus, the optimal network G , which minimizes our uncertainty, also provides the most accurate description of the observed activity.

Given a desired number of correlations, identifying the optimal network poses two distinct challenges. First, for a given network G , one must compute the model P_G that matches the covariances Σ_{ij} in G ; this can be accomplished using efficient GGM algorithms (Methods).^{32,33} Second, one must search over all networks G to find the one that provides the best compression, minimizing the entropy S_G . However, the number of possible networks explodes combinatorially, making brute force search impossible. Instead, as is common in discrete optimization, we decompose the search into a sequence of local optimization problems.^{22,34} Beginning with the empty network, we iteratively add the correlation that produces the largest decrease in entropy S_G , thus maximally reducing our uncertainty. When these drops in entropy become small, we can even expand in the weak-correlation limit analytically to further increase efficiency (Methods). We repeat this greedy algorithm until all correlations have been added. In doing so, we identify not only a single (locally) optimal network, but the sequence of all optimal compressions across all numbers of correlations. The result is a compression curve $\tilde{S}(f)$, which defines the minimum uncertainty that can be achieved with a given fraction of the correlations f (Fig. 1e).

We apply our framework to cortex-wide activity from 99 healthy adults both at rest and

across a suite of seven cognitive tasks, recorded using functional magnetic resonance imaging (fMRI) as part of the Human Connectome Project.³⁵ The collective activity consists of blood-oxygen-level-dependent (BOLD) fMRI signals from $N = 100$ cortical parcels.³⁶ Concatenating this activity across all subjects and tasks, we apply our minimax entropy algorithm to compute optimal compressions across all numbers of correlations. Strikingly, we find that with only a small number of correlations, the entropy drops significantly; only 1.4% of the correlations are needed to achieve a 50% reduction in entropy, and a 90% reduction only requires 9% of the correlations (Fig. 2a). Intuitively, one might expect the strongest correlations between regions to provide the tightest constraints on neural activity. To test this hypothesis, we compare against networks consisting of the strongest covariances. While these networks provide significantly better compressions than random correlations, the optimal networks still achieve up to orders of magnitude lower uncertainty (Fig. 2a). Together, these results indicate that the human brain is highly compressible, and that the optimal compressions themselves do not simply consist of the strongest correlations.

Given this compressibility, with only a small number of important correlations, we should be able to predict the entire correlation structure (Fig. 2b). For the optimal networks, as we increase the number of correlations, we see that the full structure quickly comes into focus (Fig. 2c). In fact, by fitting only 10% of the correlations between regions, the optimal model P_G quantitatively predicts the remaining 90% (Fig. 2e). By contrast, with the same number of random correlations, most of the structure is lost (Fig. 2d), and the model significantly underestimates the true strengths of correlations (Fig. 2f). Thus, while the brain is strongly correlated (Fig. 2b), only a sparse network of these correlations is needed to explain the rest.

Quantifying the compressibility of neural activity

We are now prepared to quantify the compressibility of the brain. Rather than choosing a specific number of correlations, we would like compressibility to be a property of the neural activity itself. We therefore define compressibility to be the amount of uncertainty that can be removed via

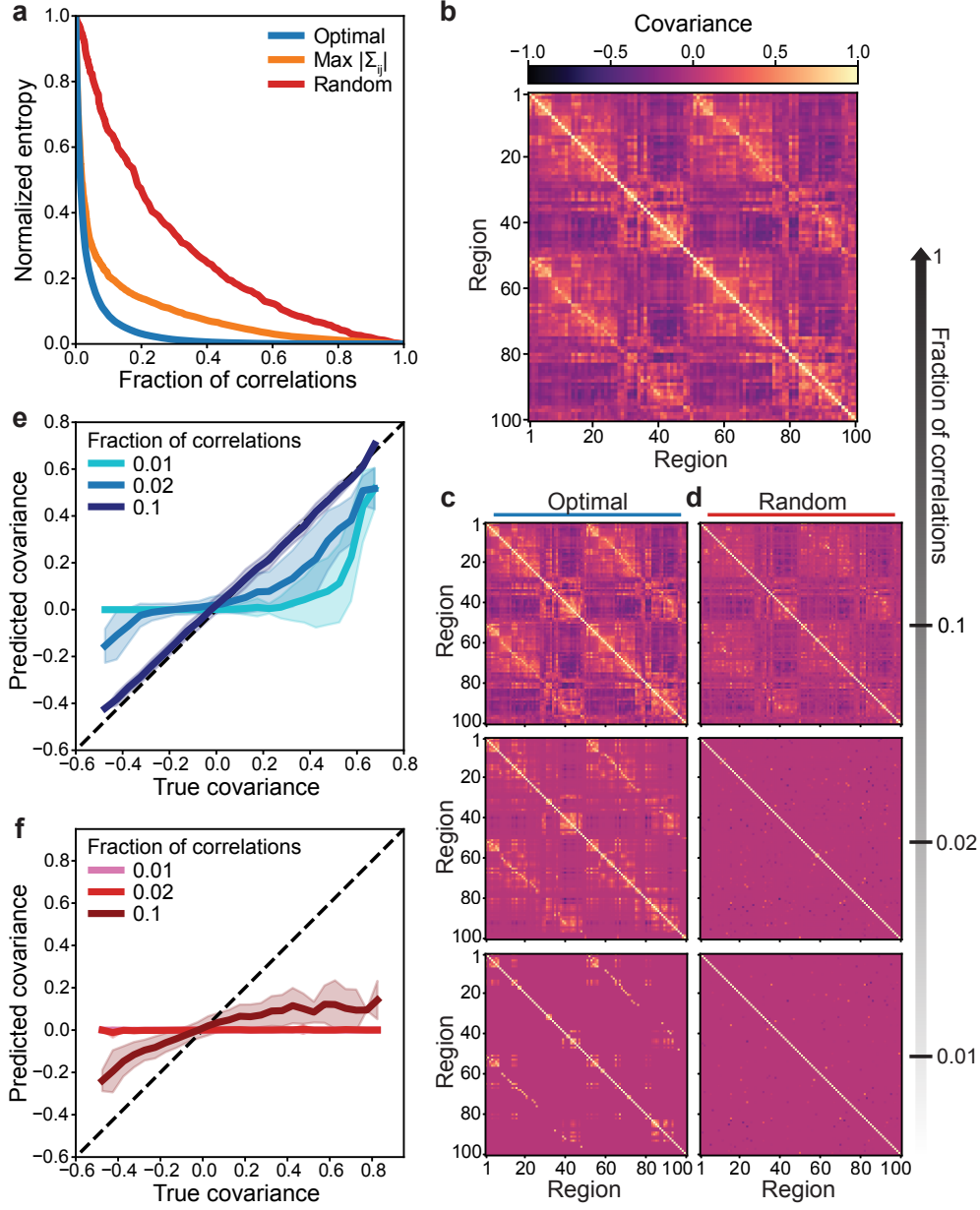


Fig. 2 | Human neural activity is highly compressible. **a**, Normalized entropy $\tilde{S}_G$ as a function of the fraction of correlations in G for the optimal networks (blue), strongest correlations $|\Sigma_{ij}|$ (orange), and random networks (red). **b**, Covariances between brain regions measured in activity concatenated across all subjects and tasks. **c-d**, Predicted covariances based on the optimal networks (**c**) and random networks (**d**) with different fractions of the correlations. **e-f**, Predicted covariances versus their true values for models constructed from optimal networks (**e**) and random networks (**f**). Lines and shaded regions indicate means and standard deviations across all covariances that are not included in each model.

compression, averaged across all numbers of correlations,

$$C = 1 - \int_0^1 \tilde{S}(f) df. \quad (3)$$

Visually, the compressibility represents the area above the compression curve (Fig. 3a). This approaches $C = 1$ for perfectly compressible activity (with all structure concentrated on one correlation) and $C = 0$ for perfectly incompressible activity (with all correlations needed to explain the structure). For the fMRI data combined across subjects and tasks, we find a compressibility $C = 0.96$ near the theoretical maximum (Fig. 3a); meanwhile, sub-optimal networks (composed of random or strong correlations) yield lower compressibilities. This means that, on average across all numbers of correlations, one can achieve a 96% reduction in uncertainty about the neural activity.

One might be concerned that concatenating activity across subjects and tasks introduces spurious correlations that lead to good compressions. Instead, we can investigate the compressibility of the brain within a given cognitive task (combined across subjects) or for a specific subject (combined across tasks, including rest). In both cases, our compression framework still discovers sparse networks of correlations that produce steep drops in entropy (Fig. 3b-c). Averaging across all numbers of correlations, we find that the brain is highly compressible both within tasks ($C = 0.96$) and for individual subjects ($C = 0.94$). Moreover, we see that the compression curves are surprisingly consistent across different tasks (Fig. 3b) and subjects (Fig. 3c), suggesting that the optimal networks themselves may be consistent. Indeed, for the top 10% of correlations, we find that the optimal networks for specific subjects overlap 53 times more strongly than random, and this overlap increases to 62 times more than random for different cognitive tasks (Fig. 3d). Given this consistency, and given that they capture the vast majority of correlations, what is the network structure of optimal compressions?

Structure of optimal correlations

We have seen that the optimal correlations between regions, which minimize our uncertainty about neural activity, are not simply the strongest (Figs. 2a and 3b-c). To understand how this is possible,

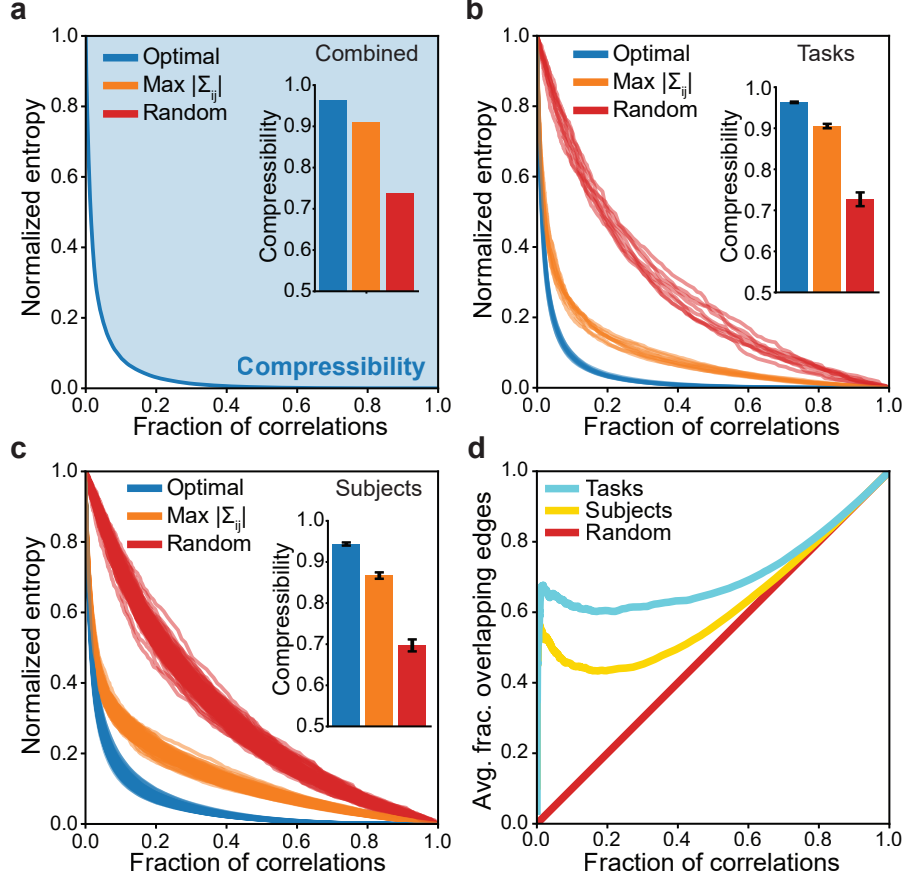


Fig. 3 | Quantifying compressibility across subjects and cognitive tasks. **a**, Optimal normalized entropy $\tilde{S}(f)$ versus the fraction of correlations f for data combined across all tasks and subjects. Compressibility is the shaded area above the compression curve. Inset displays compressibility values for optimal, maximum correlation, and random networks. **b-c**, Normalized entropy versus fraction of correlations in optimal (blue), maximum correlation (orange), and random (red) networks for data within specific tasks (**b**) and within specific subjects (**c**). Insets display compressibility values averaged across tasks (**b**) and across subjects (**c**) with one-standard-deviation error bars. **d**, Average overlap (fraction of shared correlations) between optimal networks for pairs of tasks (blue) and pairs of subjects (yellow) versus the fraction of correlations. Red line illustrates the overlap between random networks.

consider the minimal system in Fig. 4a. With knowledge of the two strongest covariances Σ_{12} and Σ_{13} , one can accurately predict the next strongest covariance Σ_{23} . Thus, when selecting the third correlation to include in the model, rather than choosing the strongest option Σ_{23} , one can gain more information by selecting a weaker covariance (in this case Σ_{14}). In general, while strong

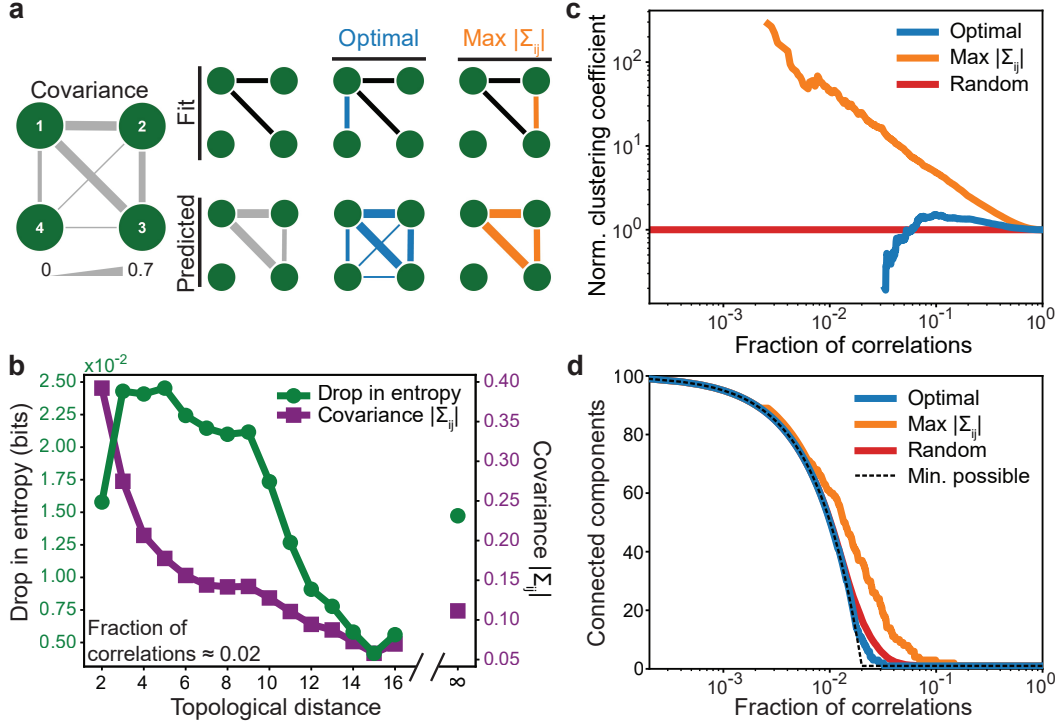


Fig. 4 | Network structure of optimal correlations. **a**, For a minimal system of four regions, the first two optimal covariances are the strongest, yielding an accurate prediction for the third strongest (*left*). Thus, the third strongest covariance is redundant, providing little improvement in accuracy (*right*). The third optimal covariance is weaker, but provides more accurate predictions (*middle*). **b**, For the optimal model with 100 correlations, we plot the average drop in entropy from fitting covariances and average strength of covariances as a function of the topological distance between brain regions in the network (that is, the length of the shortest path). Topological distance is infinity if there does not exist a path between regions. **c**, Global clustering coefficient (normalized by the density of the network) versus the fraction of correlations. Note that zero values are not shown. **d**, Number of connected components versus the fraction of correlations, where each random value is averaged over 100 random networks. Dashed line indicates the minimum possible value. In **a-c**, neural data is combined across all tasks and subjects.

correlations tend to form tight loops,³⁷ some of these may be predicted indirectly from the others; these strong but redundant correlations do not reduce our uncertainty about the neural activity.

This phenomenon plays a key role in shaping the network structure of optimal compressions. For example, for the optimal network with 100 correlations, we find that the strength of correlations decreases with increasing distance in the network (Fig. 4b). This means that selecting the strongest

correlations leads to networks with many short loops, a property known as clustering.⁷ Meanwhile, the largest drops in entropy are not achieved by these strong, short-range correlations, but instead by correlations of intermediate strength and distance in the network (Fig. 4b). As a result, while the strongest correlations form networks with high clustering, the optimal networks exhibit orders of magnitude lower clustering (Fig. 4c). Rather than concentrating connectivity within tight clusters of regions, we find that optimal networks spread connectivity across distant parts of the network. This leads to the formation of a single giant component—wherein each region connects at least indirectly to every other region—with nearly as few correlations as possible (Fig. 4d).^{38,39}

Highly correlated brain regions are known to form functionally specific subnetworks known as cognitive systems.^{7,8,40} Our above results suggest that, rather than focusing on strong correlations within these systems, one might construct better models of neural activity by including weaker correlations between different cognitive systems. To test this hypothesis, we sort the 100 regions into eight subnetworks^{36,41} Correlations within each of these cognitive systems are much more likely to appear in optimal networks than random, indicating that they are important for constraining neural activity (Fig. 5a). However, relative to the strongest correlations, optimal networks tend to connect regions spanning different cognitive systems, suggesting that many of the strong correlations within systems are redundant (Fig. 5b). In fact, while the majority of the strongest correlations lie within specific systems, most of the optimal correlations span separate subnetworks (Fig. 5c).

We observe subtle deviations from this pattern in different cognitive tasks (Fig. 5d): Correlations between the visual (VIS), default mode (DMN), and dorsal attention (DAT) systems, which are highly engaged during visual and motor response tasks,^{42–44} provide less information about cortex-wide activity than one would expect from their strength alone. Meanwhile, correlations within the limbic (LIM) and temporoparietal (TEP) systems, which are associated with emotional and social processing,^{45,46} are weak but important for constraining activity. Yet across all tasks, we consistently find that optimal compressions tend to spread connectivity between different cognitive systems (Fig. 5d). Together, these results demonstrate how, by avoiding tight clusters of strong

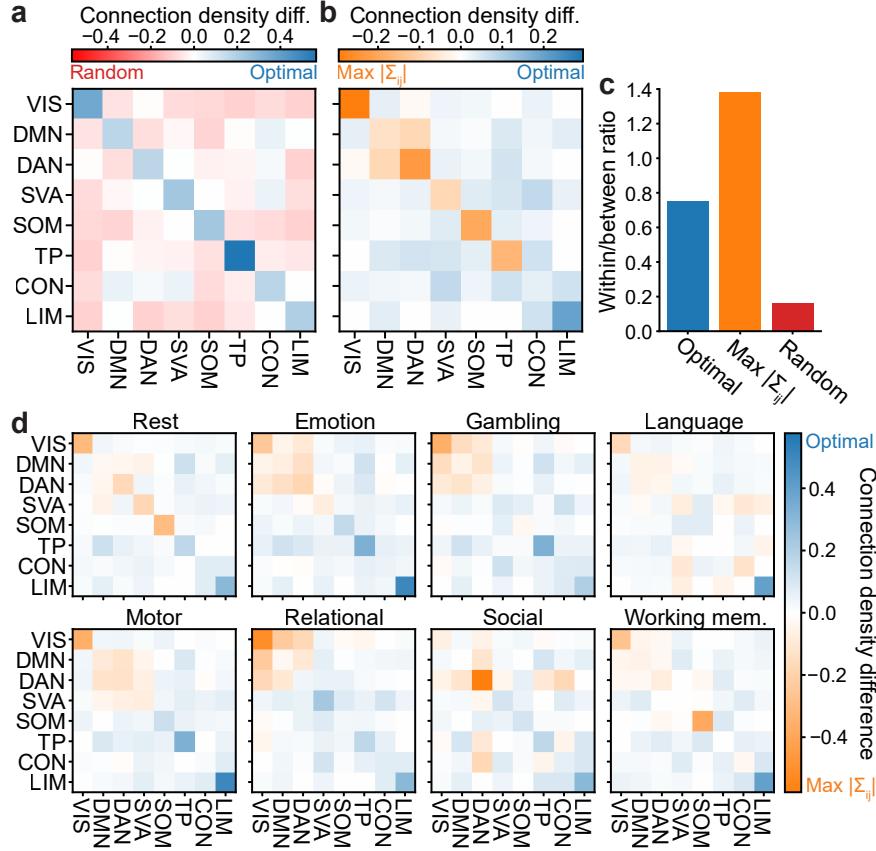


Fig. 5 | Systems-level structure of optimal compressions. **a**, Difference in the density of connections between the optimal network (blue) and random networks (red) for pairs of regions in different cognitive systems. Diagonal (off-diagonal) elements indicate connections within (between) systems. **b**, Difference in the density of connections between the optimal network (blue) and the network of strongest correlations (orange). **c**, Ratio of correlations within versus between cognitive systems. In **a-c**, neural data is combined across all tasks and subjects. **d**, Difference in the density of connections between optimal and maximum correlation networks during specific cognitive tasks including rest (data combined across subjects). Across all panels, networks contain 10% of correlations. We coarse-grain regions into eight systems: visual (VIS), default mode (DMN), dorsal attention (DAN), salience/ventral attention (SVA), somatomotor (SOM), temporoparietal (TP), control (CON), and limbic (LIM).

but redundant correlations, one can achieve a much more effective compression of the human brain.

Discussion

In the human brain, the correlations between brain regions have provided foundational insights into cognition and disease.^{1–15} But how many correlations—and which correlations—are needed to explain brain-wide activity remains poorly understood. To identify the most important correlations between brain regions, which generate the most accurate predictions of neural states, we show that one must solve the minimax entropy problem, which has roots in information theory and statistical physics.^{22–26} In fMRI data from a large cohort of subjects across several cognitive tasks,³⁵ we discover that only a small number of correlations are needed to explain the observed cortex-wide activity (Fig. 2). This reveals that the brain is highly compressible: one can ignore nearly all of the correlations between regions while still accurately predicting neural activity. In fact, we demonstrate that the compressibility of the human brain is near the theoretical maximum, and that this high compressibility is consistent across different subjects and cognitive tasks (Fig. 3). Finally, we find that the most important correlations for predicting neural activity are not simply the strongest, a fact that has profound implications for the structure of optimal compressions (Figs. 4 and 5).

These results demonstrate that the constraints on neural activity in the human brain are distributed highly non-uniformly. A small fraction of the correlations between regions have orders of magnitude more influence on brain-wide activity than the vast majority of correlations. In turn, this suggests that the human brain may be dominated by an extremely sparse backbone of abnormally strong interactions. Indeed, such heavy-tailed distributions of connection strength are observed in the structural wiring between brain regions.^{47–49} Identifying the interactions that dominate neural activity—which we have shown differ systematically from the strongest correlations—may have important implications for understanding neurological disorders and developing targeted treatments.^{7,50–52}

Moreover, with the tools to quantify the compressibility of the human brain, a number of questions immediately arise. How does the brain’s high compressibility emerge during development^{53,54}? And how is it altered in diseases such as Alzheimer’s,⁵⁵ Parkinson’s,⁵⁶ and Schizophrenia⁵⁷? Moreover, how does the network and systems-level structure of optimal compressions vary across de-

velopment and disease? With the explosion of available fMRI data, one can directly apply our compression framework to begin answering these questions. In fact, our methods can be used to study any continuous-valued recordings of neural activity, opening the door for future investigations of neural activity across different imaging modalities (including EEG and MEG^{58,59}) and even distinct species.^{60,61} Given the tendency of direct interactions to produce indirect correlations (Fig. 4a), we anticipate that many strong correlations in the brain may in fact be redundant, the natural consequences of neighboring interactions. In turn, we hypothesize that many neural systems will be amenable to simplified descriptions with only a small number of important correlations; that is, highly compressible.

Methods

Maximum entropy. Consider the matrix of covariances Σ between N brain regions. A subset of these covariances can be represented as a network G (Fig. 1). Given a network of correlations, the most unbiased prediction for the remaining covariances is provided by the maximum entropy model in Eq. (1). To compute the maximum entropy model, for each edge (ij) in the network G , one must compute the entry of the precision matrix J_{ij} such that the corresponding model covariance $(J^{-1})_{ij}$ matches the experimental value Σ_{ij} . There exist efficient algorithms for solving this problem, which are guaranteed to converge to the correct solution^{32,33,62} Here, we use Algorithms 1 and 2 in Ref. 32. Algorithm 1 iterates over entries of the precision matrix corresponding to edges in the network G , while Algorithm 2 iterates over elements of the model covariance corresponding to edges not in G . Thus, Algorithm 1 (Algorithm 2) is more efficient for networks with fewer (more) than half the available edges. When calculating maximum entropy models, we switch from Algorithm 1 to 2 at 1500 correlations, and we check that solutions converge within experimental errors (Supplementary Information).

Minimax entropy. Given a specified number of correlations, there are many different subsets—or networks—one could choose. Here we show that the optimal network, which provides the most accurate model of the neural activity, also generates the best compression of the data. Letting $Q(x)$ denote the experimental distribution over states x , the Kullback-Leibler (KL) divergence²⁰ with a maximum entropy model $P_G(x)$ simplifies to a difference in entropy,²²

$$D_{\text{KL}}(Q||P_G) = \left\langle \log \frac{Q}{P_G} \right\rangle_Q \quad (4)$$

$$= \langle \log Q \rangle_Q - \langle \log P_G \rangle_Q \quad (5)$$

$$= -S(Q) + \frac{N}{2} \log(2\pi) - \frac{1}{2} \log(|J|) + \sum_{(ij) \in G} J_{ij} \langle x_i x_j \rangle_Q \quad (6)$$

$$= -S(Q) + \frac{N}{2} \log(2\pi) - \frac{1}{2} \log(|J|) + \sum_{(ij) \in G} J_{ij} \langle x_i x_j \rangle_{P_G} \quad (7)$$

$$= -S(Q) - \langle \log P_G \rangle_{P_G} \quad (8)$$

$$= S_G - S(Q), \quad (9)$$

where $S(Q)$ is the data entropy. We therefore find that minimizing the KL divergence is equivalent to minimizing the entropy S_G of the maximum entropy model P_G . This is the minimax entropy principle tells us that the optimal network G , which provides the most accurate description of the activity, also produces the best compression of the data.

For a given number of correlations, the number of possible networks explodes combinatorially, making a brute force search for the optimal network impossible. Mathematically, this optimization is equivalent to constructing a Gaussian graphical model (GGM) with an L_0 regularization on the precision matrix, which is notoriously difficult. One heuristic

for overcoming this challenge is using the graphical lasso, with an L_1 norm on the inverse covariance to penalize a lack of sparsity.^{63–66} Here, we instead decompose the problem into a sequence of local optimizations that we can solve exactly. Specifically, we develop a greedy algorithm for constructing the optimal network G . Starting from the independent model (with no correlations), we iteratively add the correlation to our model that produces the largest drop in entropy S_G . Using this greedy approach, we can identify the locally optimal network not only for a one number of correlations, but across all numbers of correlations, thus generating the entire entropy curve (Fig. 2). We confirm that this greedy algorithm outperforms many other heuristics (Supplementary Information).

Improving the efficiency of compression. At each step of the greedy algorithm, we must fit $O(N^2)$ different maximum entropy models: one for each of the possible correlations that could be added to the model. This process must then be repeated for each of the $O(N^2)$ steps of the greedy algorithm. To scale our compression to large-scale data, we therefore need strategies for improving efficiency. To begin, we note that if the network has no loops, then each correlation produces a drop in entropy equal to the mutual information between the two brain regions.^{24,25,67} Thus, for sparse networks with no loops, one can simply select the correlations corresponding to the largest mutual information.

Once loops are included in the network, we require new strategies for efficiency. For this purpose, we derive an analytic approximation to the drop in entropy that does not require fitting a new maximum entropy model (Supplementary Information). We use perturbation theory to estimate entropy drops in the limit of small errors on the predicted correlations (or, equivalently, in the limit of small precision entries). As the greedy algorithm progresses, and the network becomes denser, the entropy drops become very small (Fig. 2a), making our approximation accurate. In practice, we compute entropy drops exactly for the first 50 correlations before switching to our approximation. We confirm the accuracy of our approximation, even in sparse networks with only 2% of correlations (Supplementary Information).

Data. We use fMRI data from the Human Connectome Project (HCP) from Ref. 35 for 99 human subjects. The HCP study was approved by the Washington University Institutional Review Board and informed consent was obtained from all subjects. Participants were not compensated. The subjects were scanned at rest and while performing seven cognitive tasks: emotion, gambling, language, motor, relational, social, and working memory. When we refer to “tasks”, we include rest for convenience. A comprehensive description of the imaging parameters and image preprocessing can be found in Ref. 68. Using a 3T Siemens Connectome Skyra with a 32-channel head coil, gradient-echo EPI images were acquired during eight conditions with the following parameters: TR = 720 ms; TE = 33.1 ms; 2-mm isotropic voxel resolution; flip angle = 52°; multiband acceleration factor = 8. See Ref. 69 for details about the details and timing of each administered task. Minimally preprocessed data, which included slice timing, motion, and distortion correction, normalized to a common surface space template were downloaded. In addition, time series were linearly detrended, temporally filtered (0.008-0.08 Hz), and had 24 head motion parameters, eight mean signals from white matter and cerebrospinal fluid, and four global signals regressed out of the time series data, corresponding to strategy

six in Ref. 70. Denoised vertex-wise time series were averaged within each area of the Schaefer 100 parcellation³⁶ at each time step, to create parcel-wise time series.

We concatenate scans for each combination of subject and task and z-score them, subtracting the mean and dividing by the standard deviation. Data was combined into larger timeseries in three ways: across all tasks and subjects (yielding combined data), across tasks for a given subject (yielding subject data), and across subjects for a given task (yielding task data). Note that all activity was mapped to the same atlas, making it meaningful to combine the time series from the same region in different subjects.^{36,41} Scans for different tasks have different lengths, with the shortest (emotion) having just 176 frames. There were two scans per task per subject, and a maximum of four for rest. For task and subject data, two subjects were excluded due to having too few rest scans. These two subjects still had their data included in the combined data. For task data, to avoid differences due to sampling, we only use the beginning of each scan with the same length as that of the shortest task. For each task, two scans (including rest) were included per subject. Subject data included two scans per non-rest task and three rest scans per subject. This ensured the same amount of data for all subjects.

Data Availability

The data analyzed in this paper are openly available at:

<https://github.com/NicholasJWeaver/BrainCompressibility2025/>

Code Availability

The code used to perform the analyses in this paper is openly available at:

<https://github.com/NicholasJWeaver/BrainCompressibility2025/>

Supplementary Information. Supplementary text and figures accompany this paper.

Acknowledgements. We thank Carlton Smith, David Carcamo, Jose Betancourt, Matt Leighton, Benjamin Machta, and Damon Clark for enlightening discussions and comments on earlier versions of the paper. N.J.W. and C.W.L. acknowledge support from the National Institutes of Health (NIH/NIGMS R35GM160188), as well as the Department of Physics, the Quantitative Biology Institute, and the Physical and Engineering Biology Program at Yale University. R.F.B. acknowledges support from the National Science Foundation (2023985), National Institute of Aging (AG075044), the National Institute on Drug Abuse (NS125026), and MNDrive Brain Conditions.

Author Contributions. N.J.W. and C.W.L. conceived the project, designed the models, and wrote the paper. N.J.W. performed the analysis. J.I.F. and R.F.B. processed the data and provided input on analyses and writing.

Competing Interests. The authors declare no competing financial interests.

Corresponding Author. Correspondence and requests for materials should be addressed to C.W.L. (christopher.lynn@yale.edu).

Supplementary Materials

1 Speeding up the minimax entropy procedure, with additional details.

We need to find ways to speed up the greedy algorithm for building minimax entropy models. We can first use the result that if adding an edge (ij) does not create any loops in the graph, the entropy drop from adding that edge is simply the mutual information between the nodes in that edge:^{24,25,67}

$$MI_{ij} = -\frac{1}{2} \log \left[1 - \frac{\Sigma_{ij}\Sigma_{ji}}{\Sigma_{ii}\Sigma_{jj}} \right] \quad (10)$$

This can be computed quickly and does not require us to find the maxent model itself.

As we grow the network of constrained correlations, the algorithm used to compute a particular maximum entropy model takes more and more time to run. We can speed up the algorithm by taking advantage of the fact that it iterates over one edge at a time, allowing us to rewrite the algorithm using results for low rank matrix updates. Even with this speedup, adding edges to the graph takes longer and longer. As we create more opportunities to add loops, we have fewer opportunities to use the mutual information trick and have to run the full maxent algorithm more and more.

To make progress, we developed a method for predicting the entropy drop that would result from adding any edge to the network without actually running the maxent algorithm (See section 2). Past a certain number of edges added, in this case 50, we switch from exactly finding the entropy drop from adding each candidate edge to predicting it. This allows us to then add the edge with the largest predicted entropy drop and run the maxent algorithm only once, to find the model corresponding to that new network of constraints. This model is then used in the process of predicting all of the entropy drops at the next step, and the process repeats. We can even increase the number of edges we add at a given time as the network of constraints becomes more dense. We add one edge per step until 200 edges added, then two per step until 400 edges, then five per step until 600 edges, then 10 per step until 800 edges, then 20 per step until 1000 edges, then 50 per step for the remainder of the process.

We use one final speedup. Maxent Algorithm 2,³² which we discussed in the main text, runs quickly for sparse networks of edges but takes longer as the network becomes more dense. Algorithm 1³² is slower for sparse networks but becomes faster for denser and denser networks. At some point for medium network density, in our case chosen to be 1500 edges added out of 4950, we switch from Algorithm 2 to Algorithm 1. With all of these steps, we are able to build minimax networks all the way from the independent model with no edges constrained to the full model with all of them constrained in a reasonable amount of time, on the order of hours.

2 Predicting the entropy drop from adding an edge.

We imagine wiggling one element of the covariance and seeing how the model entropy changes, while enforcing the constraints imposed by the entropy maximization procedure. We can perform a series expansion of the entropy change in terms of the error in the corresponding element of the covariance matrix. Here we denote the model covariance as K , with model precision $J = K^{-1}$, and the data covariance as Σ . We fit all diagonal elements of the covariance from the start, so we can assume $i \neq j$

$$\Delta S = \frac{dS}{dK_{ij}}(K_{ij} - \Sigma_{ij}) + \frac{1}{2} \frac{d^2 S}{dK_{ij}^2}(K_{ij} - \Sigma_{ij})^2 + \dots \quad (11)$$

The entropy of a Gaussian probability distribution has the following form:²⁰

$$S = \frac{N}{2} \ln(2\pi e) + \frac{1}{2} \ln(|K|) \quad (12)$$

Starting with the first order term:

$$\frac{dS}{dK_{ij}} = \frac{\partial S}{\partial K_{ij}} + \sum_{(kl) \notin E^*, (kl) \neq (ij)} \frac{\partial S}{\partial K_{kl}} \frac{dK_{kl}}{dK_{ij}} \quad (13)$$

Where E^* is the network of covariances that are being fit, including the diagonal (meaning the variances). The elements of K that are being fit are not free parameters as they are constrained to match their experimental values, whereas the edges not in E^* are able to change, which is why we sum over them above. Let's look at the partial derivatives:

$$\frac{\partial S}{\partial K_{ij}} = \frac{1}{2|K|} \frac{\partial |K|}{\partial K_{ij}} \quad (14)$$

To obtain the derivative on the right hand side, we use the following result for the partial derivative of the determinant:⁷¹

$$\frac{\partial |K|}{\partial K_{ij}} = |K| \operatorname{Tr} \left(J \frac{\partial K}{\partial K_{ij}} \right) \quad (15)$$

The matrix $\frac{\partial K}{\partial K_{ij}}$ is equal to one for both the ij and ji elements, as those elements are really the same parameter due to the enforced symmetry of the covariance matrix. We can denote the mn -th element of this as:

$$\left(\frac{\partial K}{\partial K_{ij}} \right)_{mn} = \delta_{mi} \delta_{nj} + \delta_{mj} \delta_{ni} \quad (16)$$

With this, we have:

$$\operatorname{Tr} \left(J \frac{\partial K}{\partial K_{ij}} \right) = \sum_m \sum_k J_{mk} (\delta_{ki} \delta_{mj} + \delta_{kj} \delta_{mi}) = J_{ji} + J_{ij} = 2J_{ij} \quad (17)$$

Using this, we have:

$$\frac{\partial S}{\partial K_{ij}} = \frac{1}{2|K|} \frac{\partial |K|}{\partial K_{ij}} = \frac{1}{2|K|} |K| 2J_{ij} = J_{ij} \quad (18)$$

Plugging this into our total derivative expression, while noting that if $(kl) \notin E^*$, then $k \neq l$ as we fit the variances from the start:

$$\frac{dS}{dK_{ij}} = J_{ij} + \sum_{(kl) \notin E^*, (kl) \neq (ij)} J_{kl} \frac{dK_{kl}}{dK_{ij}} \quad (19)$$

Entropy maximization requires that $J_{kl} = 0$ for $(kl) \notin E^*$, so we see that the sum on the right hand side vanishes (assuming the total derivative term is well behaved), and we are left with:

$$\frac{dS}{dK_{ij}} = J_{ij} \quad (20)$$

We can now move on to the second derivative term in our expansion. We see now that

$$\frac{d^2 S}{dK_{ij}^2} = \frac{dJ_{ij}}{dK_{ij}} \quad (21)$$

We will proceed by calculating $\frac{dK_{ij}}{dJ_{ij}}$ then inverting. We can rewrite the above as follows:

$$\frac{dK_{ij}}{dJ_{ij}} = \frac{\partial K_{ij}}{\partial J_{ij}} + \sum_{(kl) \in E^*} \frac{\partial K_{ij}}{\partial J_{kl}} \frac{dJ_{kl}}{dJ_{ij}} \quad (22)$$

For the partial derivative of an inverse matrix with respect to an element of the original matrix, the following result is known for a general matrix M :⁷¹

$$\frac{\partial M_{ij}^{-1}}{\partial M_{kl}} = - \left(M^{-1} \frac{\partial M}{\partial M_{kl}} M^{-1} \right)_{ij} \quad (23)$$

Using this, and setting $M = J$ and $M^{-1} = K$, we have:

$$\frac{\partial K_{ij}}{\partial J_{kl}} = - \left(K \frac{\partial J}{\partial J_{kl}} K \right)_{ij} \quad (24)$$

Using similar notation as in Eqn. (16) we have, for $k \neq l$:

$$\frac{\partial K_{ij}}{\partial J_{kl}} = - \sum_m K_{im} \sum_n (\delta_{mk} \delta_{nl} + \delta_{ml} \delta_{nk}) K_{nj} = -K_{ik} K_{lj} - K_{il} K_{kj} \quad (25)$$

For $k = l$ we have:

$$\frac{\partial K_{ij}}{\partial J_{kl}} = - \sum_m K_{im} \sum_n \delta_{mk} \delta_{nk} K_{nj} = -K_{ik} K_{kj} \quad (26)$$

Now we just need to find $\frac{dJ_{kl}}{dJ_{ij}}$, then we will have the full expression for the second derivative term.

To make progress, we use the following trick: For an edge $(kl) \in E^*$, and $(ij) \notin E^*$ we have that K_{kl} cannot change, as it is being fit, and:

$$\frac{dK_{kl}}{dJ_{ij}} = 0 = \frac{\partial K_{kl}}{\partial J_{ij}} + \sum_{(mn) \in E^*} \frac{\partial K_{kl}}{\partial J_{mn}} \frac{dJ_{mn}}{dJ_{ij}} \quad (27)$$

We can solve the equation above to find $\frac{dJ_{mn}}{dJ_{ij}}$. We do this by treating the above as a matrix equation $A\vec{x} = \vec{b}$, where:

$$\vec{b}_{e_1} = -\frac{\partial K_{e_1}}{\partial J_{ij}}, \vec{x}_{e_2} = \frac{\partial J_{e_2}}{\partial J_{ij}}, A_{e_1 e_2} = \frac{\partial K_{e_1}}{\partial J_{e_2}} \quad (28)$$

and e_1 and e_2 are edges in E^* . We can obtain all of the elements in A and b using (25). Assuming A is invertible, we have:

$$\frac{dJ_{kl}}{dJ_{ij}} = - \sum_{(mn) \in E^*} A_{(kl), (mn)}^{-1} \frac{\partial K_{mn}}{\partial J_{ij}} \quad (29)$$

We can summarize

$$\frac{dK_{ij}}{dJ_{ij}} = \frac{\partial K_{ij}}{\partial J_{ij}} - \sum_{(kl) \in E^*} \frac{\partial K_{ij}}{\partial J_{kl}} \sum_{(mn) \in E^*} A_{(kl), (mn)}^{-1} \frac{\partial K_{mn}}{\partial J_{ij}} \quad (30)$$

We now have all that we need. Eqn. 11 becomes:

$$\Delta S \approx J_{ij}(K_{ij} - \Sigma_{ij}) + \frac{1}{2} \left[\frac{\partial K_{ij}}{\partial J_{ij}} - \sum_{(kl) \in E^*} \frac{\partial K_{ij}}{\partial J_{kl}} \sum_{(mn) \in E^*} A_{(kl),(mn)}^{-1} \frac{\partial K_{mn}}{\partial J_{ij}} \right]^{-1} (K_{ij} - \Sigma_{ij})^2 \quad (31)$$

We are evaluating this for an edge (ij) that has not yet been added to the network, so $J_{ij} = 0$, meaning that the first order term in the error vanishes. With the elements of the matrix A and the partial derivatives given previously, everything in this expression now depends only on the elements of the current model covariance K (that is, the one found at the prior step without the edge (ij) constrained), and the data covariance Σ . While the final result may not look pretty, it can be numerically calculated much faster than it would take to actually find the true maxent model with the same edge added. A comparison of predicted versus actual entropy drops can be found in Fig. S1 and we find good agreement.

3 Predicting the information gained from adding a constraint

Here we compare our predictions for the true entropy drops, obtained by finding the corresponding maxent model and calculating its entropy, for each of the possible edges that could be added to the optimal network with 100 edges added for the combined fMRI data. In Fig. S1a we see the predicted and true information gains versus the error on the prediction of the corresponding element of the covariance before that edge was added. Notice the approximately quadratic shape, which is assuring given that the expansion in our prediction goes to second order in the error, with the linear term being zero. In Fig. S1b we see the predicted information gains plotted versus the true gains, finding good agreement.

4 Comparison of our method of selecting edges to heuristics

One could try to choose which correlations to add to the model as we build it using some heuristic, for example picking the largest correlations in magnitude, as we did in the main text. There are many other heuristics one could use. In figure S2 we compare the optimal, random, and $\max(|\Sigma_{ij}|)$ methods from the main text to a list of other methods, all computed on the full data: $\max(\Sigma_{ij})$, $\max$

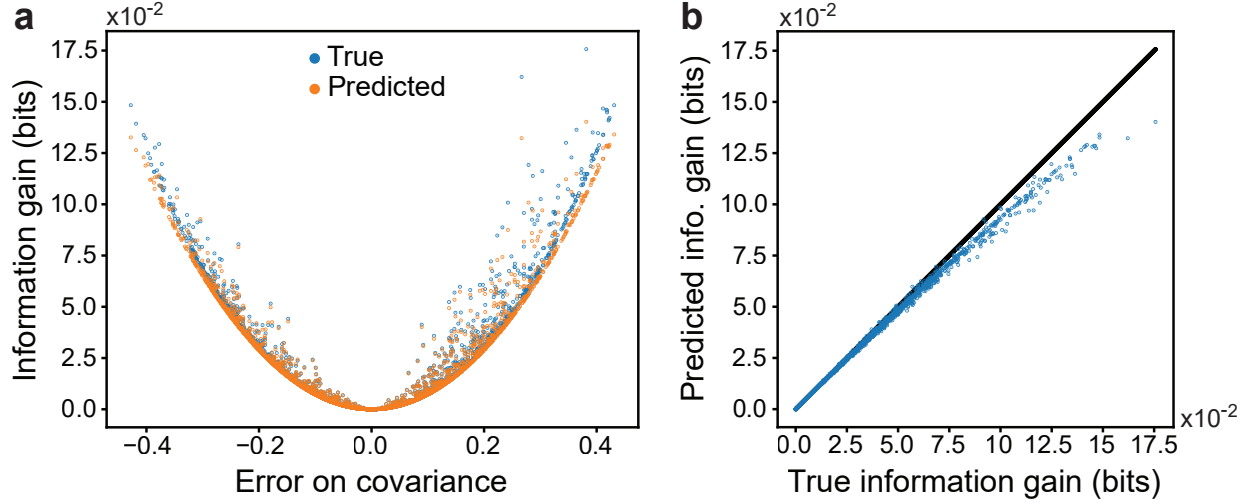


Fig. S1 | Information gained from adding a constraint can be predicted accurately. **a**, Predicted and true information gains from adding each possible next constraint to the optimal network with 100 correlations fit, for the combined fMRI data, versus the error on the prediction of the corresponding covariance before it is fit. For each of true and predicted, each point represents one of the 4850 possible correlations that could be added to become the 101st edge in the network. **b**, Predicted information gain from adding each possible next constraint to the optimal network with 100 correlations fit, for the combined fMRI data, versus the true information gain. Each point represents one of the 4850 possible correlations that could be constrained. The diagonal line represents a slope of one.

partial correlation in magnitude, max signed partial correlation, $\max(|\Sigma_{ij}^{-1}|)$, $\max(\Sigma_{ij}^{-1})$, $\min(\Sigma_{ij}^{-1})$, and largest mutual informations, where Σ^{-1} is the data precision matrix. For Gaussians, the partial correlation is $\frac{-\Sigma_{ij}^{-1}}{\sqrt{\Sigma_{ii}^{-1}\Sigma_{jj}^{-1}}}$, which measures the correlation between regions i and j when controlling for the indirect dependencies between them through all of the other regions in the brain. We see varying degrees of success, with the largest absolute partial correlations and the $\max(|\Sigma_{ij}^{-1}|)$ working the best.

5 Additional brain-system level structure analysis

In main text Fig. 5a, we see the brain-system level structure of optimal models compared to the random expectation for all tasks. Specifically we examine the difference in connection density between/within systems for the optimal model and random expectation for networks with approx-

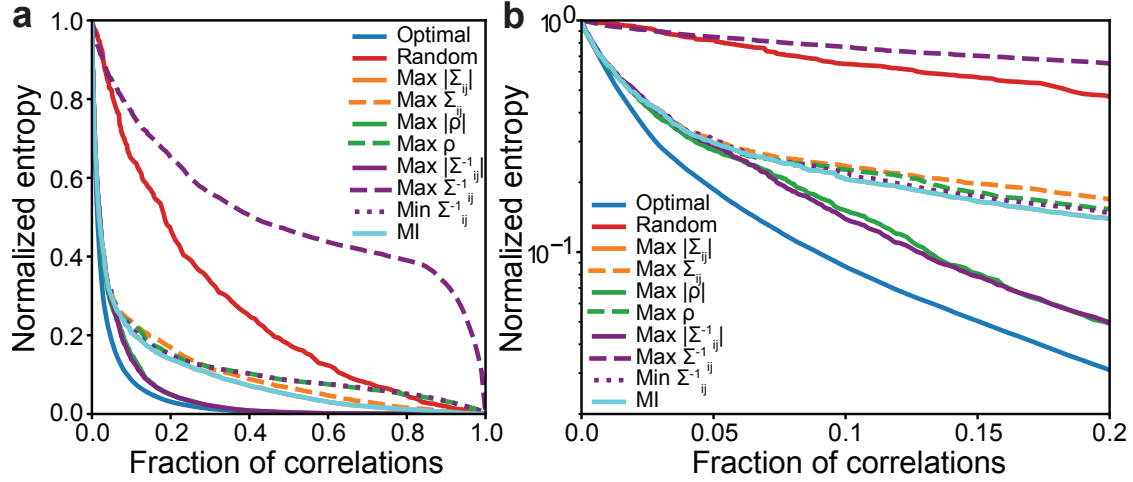


Fig. S2 | Model entropy for heuristic methods. **a**, Normalized model entropy of models of the combined fMRI data for various heuristics for selecting correlations to add to the model, compared to optimal and random methods. Here ρ stands for partial correlation, Σ is the data covariance, Σ^{-1} is the data precision matrix, and MI stands for mutual information. **b**, Same as **a** but with log scale y-axis and for fraction of correlations between 0 and 0.2

imately 10% of correlations. In Fig. S3a we look at the same but for each task including rest. One consistent feature is that there are more edges within systems than would be expected at random for all systems, which indicates that the coarse-graining of regions into systems was reasonable. One might also be interested to see which regions are more or less likely to be informative for a certain task as compared to rest. In Fig. S3b we see the difference of connection density between each task and rest, again for 10% of correlations. We see some unique differences for each task, although we will not attempt to interpret them here.

6 Model predictions relative to bootstrap errors

The high compressibilities seen in main text Fig. 3 suggest that with only a small fraction of correlations constrained, we should be able to accurately predict the rest. In order to get a sense for how accurate the predictions really are, we should measure the error relative to the inherent scale of uncertainties in the correlations themselves. To do this, we use the standard deviation of a correlation across bootstrap samples of the data (see Methods). We calculate the maximum

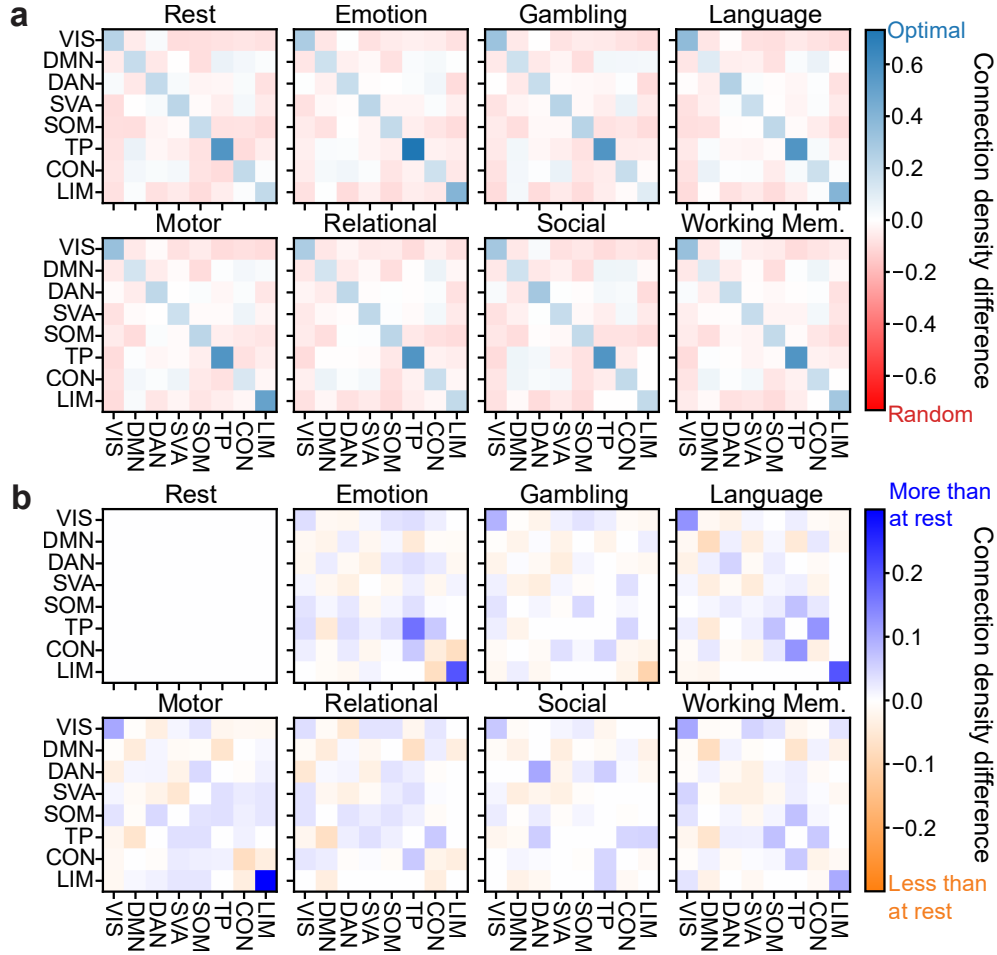


Fig. S3 | System level structure of optimal models. **a**, Difference in connection density between/within systems for the optimal model (blue) and random expectation (red) for tasks, with approximately 10% of correlations **b**, Difference of connection density between each task and rest, with approximately 10% of correlations. Blue corresponds to higher density in the task than at rest, and orange corresponds to less.

error of our models for the combined data, for tasks, and for subjects on correlations not being fit in the models, divided by the bootstrap standard deviation of the corresponding correlation, with the results shown in Fig. S 4. We see that for the combined data, we need around 85.86% of the edges to predict the remaining correlations to within two bootstrap standard deviations. This is perhaps not surprising given that each pair of regions may exhibit a significant correlation in at least one subject or task. If we focus on specific tasks, it takes a significantly smaller fraction, between 10.2% and 12.73% of the correlations added, to predict the others within errors. For

individual subjects, it takes between 47.47% and 60.61% of the correlations. In all cases, we see a large decrease in the max error over bootstrap standard deviation for a relatively small fraction of edges, meaning that the sparse networks of covariances are able to give the majority of the reduction in prediction error. We would expect that once the predictions are within errors, the model is capturing most of the available information, and we see that this is indeed the case in Fig. S 4e. In summary, we find a large decrease in normalized prediction errors for a small fraction of correlations across subjects, tasks, and combined data, with the point at which the predictions are within errors changing significantly for different ways of slicing the data.

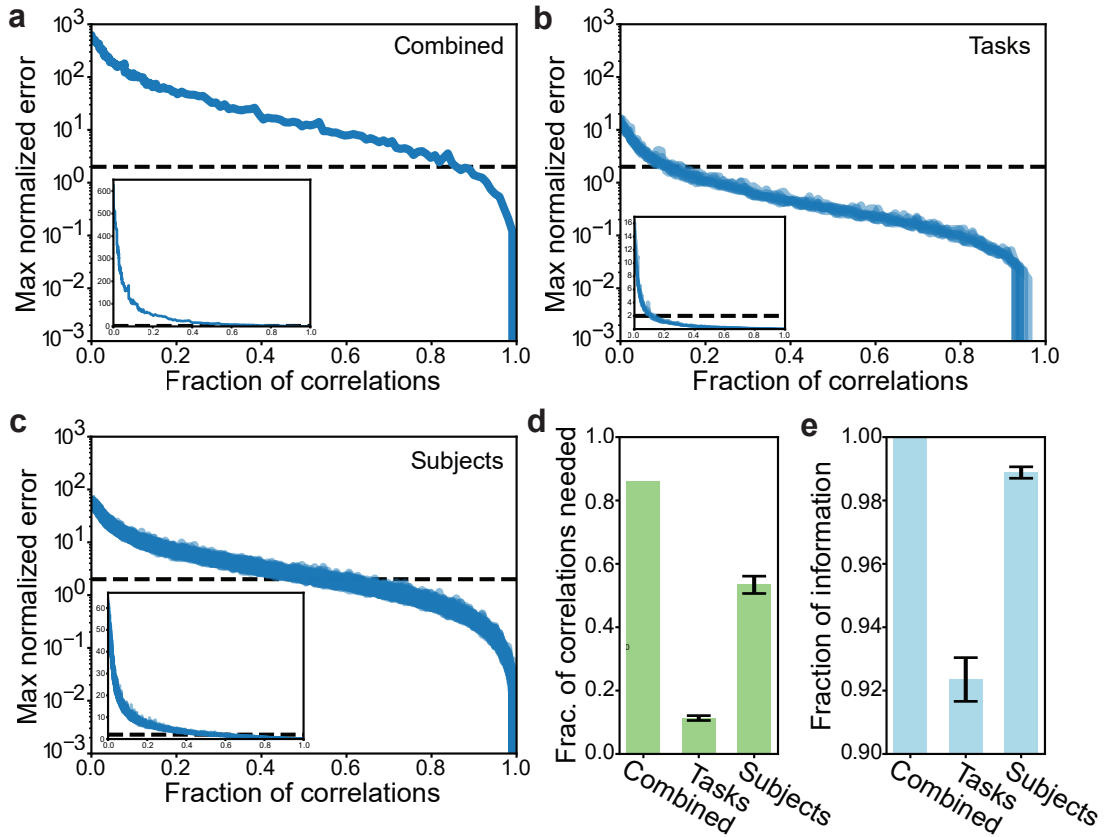


Fig. S4 | Maximum error on unconstrained covariances relative to bootstrap standard deviation. **a**, Maximum of errors on unconstrained covariances divided by corresponding bootstrap standard deviation for data combined across both subjects and tasks. Dashed lines in a-c represent value of 2. Inset shows linear-linear plot of same results. **b**, Same as in (a) for data combined across subjects for a given task. **c**, Same as in (a) for data combined across tasks for a given subject. **d**, Fraction of edges added at which threshold of 2 is crossed. **e**, Fraction of possible entropy drop at the time when threshold is crossed.

7 Tolerances and Bootstrapping

To find reasonable error tolerances for the maxent algorithms above, we bootstrap sample the data 100 times and calculate the model covariance and precision matrix for each of those bootstrap sample timeseries. We then find the standard deviation of the elements of each matrix. We set the error tolerance as the smallest standard deviation divided by 100. For task models we used the smallest tolerances across tasks for all of the tasks due to a reasonably large difference in tolerances between tasks. For the subject models, we used each individual's tolerances.

References

1. Bassett, D. S. & Bullmore, E. T. Human brain networks in health and disease. *Curr. Opin. Neurol.* **22**, 340–347 (2009).
2. Bressler, S. L. & Menon, V. Large-scale brain networks in cognition: emerging methods and principles. *Trends Cogn. Sci.* **14**, 277–290 (2010).
3. Park, H.-J. & Friston, K. Structural and functional brain networks: from connections to cognition. *Science* **342**, 1238411 (2013).
4. Rugg, M. D. & Vilberg, K. L. Brain networks underlying episodic memory retrieval. *Current opinion in neurobiology* **23**, 255–260 (2013).
5. Bassett, D. S. *et al.* Dynamic reconfiguration of human brain networks during learning. *Proc. Natl. Acad. Sci. U.S.A.* **108**, 7641–7646 (2011).
6. Fincham, J. M., Carter, C. S., van Veen, V., Stenger, V. A. & Anderson, J. R. Neural mechanisms of planning: a computational analysis using event-related fMRI. *Proc. Natl. Acad. Sci. U.S.A.* **99**, 3346–3351 (2002).
7. Lynn, C. W. & Bassett, D. S. The physics of brain network structure, function and control. *Nat. Rev. Phys.* **1**, 318–332 (2019).
8. Bassett, D. S. & Sporns, O. Network neuroscience. *Nat. Neurosci.* **20**, 353–364 (2017).
9. van den Heuvel, M. P. & Hulshoff Pol, H. E. Exploring the brain network: A review on resting-state fMRI functional connectivity. *Eur. Neuropsychopharmacol.* **20**, 519–534 (2010).
10. Castellanos, F. X., Di Martino, A., Craddock, R. C., Mehta, A. D. & Milham, M. P. Clinical applications of the functional connectome. *NeuroImage* **80**, 527–540 (2013).
11. Dennis, E. L. & Thompson, P. M. Functional brain connectivity using fMRI in aging and Alzheimer’s disease. *Neuropsychology review* **24**, 49–62 (2014).

12. Baggio, H.-C. *et al.* Functional brain networks and cognitive deficits in Parkinson’s disease. *Human brain mapping* **35**, 4620–4634 (2014).
13. Uddin, L. Q., Supekar, K. & Menon, V. Typical and atypical development of functional human brain networks: insights from resting-state fMRI. *Frontiers in systems neuroscience* **4**, 1447 (2010).
14. Chan, M. Y. *et al.* Socioeconomic status moderates age-related differences in the brain’s functional network organization and anatomy across the adult lifespan. *Proc. Natl. Acad. Sci. U.S.A.* **115**, E5144–E5153 (2018).
15. Yang, Z. *et al.* Genetic and environmental contributions to functional connectivity architecture of the human brain. *Cerebral cortex* **26**, 2341–2352 (2016).
16. Hipp, J. F., Hawellek, D. J., Corbetta, M., Siegel, M. & Engel, A. K. Large-scale cortical correlation structure of spontaneous oscillatory activity. *Nat. Neurosci.* **15**, 884–890 (2012).
17. Jazayeri, M. & Ostojic, S. Interpreting neural computations by examining intrinsic and embedding dimensionality of neural activity. *Curr. Opin. Neurobiol.* **70**, 113–120 (2021).
18. Mehrkanoon, S., Breakspear, M. & Boonstra, T. W. Low-dimensional dynamics of resting-state cortical activity. *Brain Topogr.* **27**, 338–352 (2014).
19. Shannon, C. E. A mathematical theory of communication. *Bell Syst. Tech. J.* **27**, 379–423 (1948).
20. Cover, T. M. & Thomas, J. A. *Elements of Information Theory* (Wiley, 2006), 2nd edn.
21. Jaynes, E. T. Information theory and statistical mechanics. *Phys. Rev.* **106**, 620 (1957).
22. Carcamo, D. P., Weaver, N. J., Dixit, P. D. & Lynn, C. W. Minimax entropy: The statistical physics of optimal models (2025). URL <https://arxiv.org/abs/2505.01607>. 2505.01607.

23. Zhu, S. C., Wu, Y. N. & Mumford, D. Minimax entropy principle and its application to texture modeling. *Neural Comput.* **9**, 1627–1660 (1997).
24. Lynn, C. W., Yu, Q., Pang, R., Bialek, W. & Palmer, S. E. Exactly solvable statistical physics models for large neuronal populations. *Phys. Rev. Res.* **7**, L022039 (2025).
25. Lynn, C. W., Yu, Q., Pang, R., Palmer, S. E. & Bialek, W. Exact minimax entropy models of large-scale neuronal activity. *Phys. Rev. E* **111**, 054411 (2025).
26. Carcamo, D. P. & Lynn, C. W. Statistical physics of large-scale neural activity with loops (2024). URL <https://arxiv.org/abs/2412.18115>. 2412.18115.
27. Dowson, D. & Wragg, A. Maximum-entropy distributions having prescribed first and second moments. *IEEE Trans. Inf. Theory* **19**, 689–693 (1973).
28. Lauritzen, S. *Graphical Models* (Oxford University Press, 1996).
29. Dyrba, M., Mohammadi, R., Grothe, M., Kirste, T. & Teipel, S. Gaussian graphical models reveal inter-modal and inter-regional conditional dependencies of brain alterations in Alzheimer’s disease. *Front. Aging Neurosci.* **12** (2020).
30. Zhang, A., Fang, J., Liang, F., Calhoun, V. D. & Wang, Y.-P. Aberrant brain connectivity in schizophrenia detected via a fast Gaussian graphical model. *IEEE J. Biomed. Health Inform.* **23**, 1479–1489 (2019).
31. Greenewald, K., Park, S., Zhou, S. & Giessing, A. Time-dependent spatially varying graphical models, with application to brain fMRI data analysis. In *Adv. Neural Inf. Process.*, vol. 30 (2017). URL https://proceedings.neurips.cc/paper_files/paper/2017/file/769675d7c11f336ae6573e7e533570ec-Paper.pdf.
32. Uhler, C. Gaussian graphical models: An algebraic and geometric perspective (2017). URL <https://arxiv.org/abs/1707.04345>. 1707.04345.
33. Wermuth, N. & Scheidt, E. Fitting a covariance selection model to a matrix. *J. R. Stat. Soc., C: Appl. Stat.* **26**, 88–92 (1977).

34. Cormen, T., Leiserson, C., Rivest, R. & Stein, C. *Introduction to Algorithms* (MIT Press, 2009), 3rd edn.
35. Van Essen, D. C. *et al.* The WU-Minn human connectome project: An overview. *NeuroImage* **80**, 62–79 (2013).
36. Schaefer, A. *et al.* Local-global parcellation of the human cerebral cortex from intrinsic functional connectivity MRI. *Cereb. Cortex* **28**, 3095–3114 (2018).
37. Masuda, N., Sakaki, M., Ezaki, T. & Watanabe, T. Clustering coefficients for correlation networks. *Front. Neuroinform.* **12**, 7 (2018).
38. Newman, M. *Networks* (Oxford university press, 2018).
39. Molloy, M. & Reed, B. A critical point for random graphs with a given degree sequence. *Random Struct. Algorithms* **6**, 161–180 (1995).
40. Crossley, N. A. *et al.* Cognitive relevance of the community structure of the human brain functional coactivation network. *Proc. Natl. Acad. Sci. U.S.A.* **110**, 11583–11588 (2013).
41. Yeo, B. T. *et al.* The organization of the human cerebral cortex estimated by intrinsic functional connectivity. *J. Neurophysiol.* **106**, 1125–65 (2011).
42. Raichle, M. E. *et al.* A default mode of brain function. *Proc. Natl. Acad. Sci. U.S.A.* **98**, 676–682 (2001).
43. Szczepanski, S. M., Pinsk, M. A., Douglas, M. M., Kastner, S. & Saalmann, Y. B. Functional and structural architecture of the human dorsal frontoparietal attention network. *Proc. Natl. Acad. Sci. U.S.A.* **110**, 15806–15811 (2013).
44. Desimone, R., Duncan, J. *et al.* Neural mechanisms of selective visual attention. *Annu. Rev. Neurosci.* **18**, 193–222 (1995).
45. Catani, M., Dell’Acqua, F. & De Schotten, M. T. A revised limbic system model for memory, emotion and behaviour. *Neurosci. Biobehav. Rev.* **37**, 1724–1737 (2013).

46. Krall, S. C. *et al.* The role of the right temporoparietal junction in attention and social interaction as revealed by ALE meta-analysis. *Brain Struct. Funct.* **220**, 587–604 (2015).
47. Roberts, J. A. *et al.* The contribution of geometry to the human connectome. *Neuroimage* **124**, 379–393 (2016).
48. Gast, H. & Assaf, Y. Weighting the structural connectome: Exploring its impact on network properties and predicting cognitive performance in the human brain. *Netw. Neurosci.* **8**, 119–137 (2024).
49. Cirunay, M., Ódor, G., Papp, I. & Deco, G. Scale-free behavior of weight distributions of connectomes. *Phys. Rev. Res.* **7**, 013134 (2025).
50. Breit, S., Schulz, J. B. & Benabid, A.-L. Deep brain stimulation. *Cell Tissue Res.* **318**, 275–288 (2004).
51. Kobayashi, M. & Pascual-Leone, A. Transcranial magnetic stimulation in neurology. *Lancet Neurol.* **2**, 145–156 (2003).
52. Betzel, R. F., Gu, S., Medaglia, J. D., Pasqualetti, F. & Bassett, D. S. Optimally controlling the human connectome: the role of network topology. *Sci. Rep.* **6**, 30770 (2016).
53. Luna, B., Padmanabhan, A. & O’Hearn, K. What has fMRI told us about the development of cognitive control through adolescence? *Brain Cogn.* **72**, 101–113 (2010).
54. Dosenbach, N. U. *et al.* Prediction of individual brain maturity using fMRI. *Science* **329**, 1358–1361 (2010).
55. Wang, K. *et al.* Altered functional connectivity in early Alzheimer’s disease: A resting-state fMRI study. *Hum. Brain Mapp.* **28**, 967–978 (2007).
56. Wolters, A. F. *et al.* Resting-state fMRI in Parkinson’s disease patients with cognitive impairment: a meta-analysis. *Parkinsonism & related disorders* **62**, 16–27 (2019).

57. 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).
58. Chang, C., Liu, Z., Chen, M. C., Liu, X. & Duyn, J. H. EEG correlates of time-varying BOLD functional connectivity. *Neuroimage* **72**, 227–236 (2013).
59. Brookes, M. J. *et al.* Measuring functional connectivity using MEG: methodology and comparison with fcMRI. *Neuroimage* **56**, 1082–1104 (2011).
60. Stephan, K. E. *et al.* Computational analysis of functional connectivity between areas of primate cerebral cortex. *Philos. Trans. R. Soc. Lond. B Biol. Sci.* **355**, 111–126 (2000).
61. White, B. R. *et al.* Imaging of functional connectivity in the mouse brain. *PLoS One* **6**, e16322 (2011).
62. Speed, T. & Kiiveri, H. Gaussian Markov distributions over finite graphs. *Ann. Stat.* **14**, 138–150 (1986).
63. Friedman, J., Hastie, T. & Tibshirani, R. Sparse inverse covariance estimation with the graphical lasso. *Biostat.* **9**, 432–441 (2007).
64. Varoquaux, G. & Craddock, R. C. Learning and comparing functional connectomes across subjects. *NeuroImage* **80**, 405–415 (2013).
65. Varoquaux, G., Gramfort, A., Poline, J.-b. & Thirion, B. Brain covariance selection: better individual functional connectivity models using population prior. In *Adv. Neural Inf. Process.*, vol. 23 (2010). URL https://proceedings.neurips.cc/paper_files/paper/2010/file/db576a7d2453575f29eab4bac787b919-Paper.pdf.
66. Banerjee, O., El Ghaoui, L., d’Aspremont, A. & Natsoulis, G. Convex optimization techniques for fitting sparse Gaussian graphical models. In *Proceedings of the 23rd International Conference on Machine Learning* (2006).
67. Nguyen, H. C., Zecchina, R. & Berg, J. Inverse statistical problems: from the inverse Ising problem to data science. *Adv. Phys.* **66**, 197–261 (2017).

68. Glasser, M. F. *et al.* The minimal preprocessing pipelines for the human connectome project. *Neuroimage* **80**, 105–124 (2013).
69. Barch, D. M. *et al.* Function in the human connectome: task-fMRI and individual differences in behavior. *Neuroimage* **80**, 169–189 (2013).
70. Parkes, L., Fulcher, B., Yücel, M. & Fornito, A. An evaluation of the efficacy, reliability, and sensitivity of motion correction strategies for resting-state functional MRI. *Neuroimage* **171**, 415–436 (2018).
71. Magnus, J. R. & Neudecker, H. *Matrix Differential Calculus with Applications in Statistics and Econometrics* (Wiley and Sons, 2019), 3rd edn.