



Unbiased and efficient sampling of timeseries reveals redundancy of brain network and gradient structure

Aditya Nanda^{a,*}, Mikail Rubinov^{a,b,c,*}

^a Department of Biomedical Engineering, Vanderbilt University, USA

^b Department of Computer Science, Vanderbilt University, USA

^c Janelia Research Campus, Howard Hughes Medical Institute, USA

ABSTRACT

Many studies in human neuroscience seek to understand the structure of brain networks and gradients. Few studies, however, have tested the redundancy between these outwardly distinct features. Here, we developed methods to directly enable such tests. We built on insights from linear algebra to develop methods for unbiased and efficient sampling of timeseries with network or gradient constraints. We used these methods to show considerable redundancy between popular definitions of network and gradient structure in functional MRI data. On the one hand, we found that network constraints largely accounted for the structure of three major gradients. On the other hand, we found that gradient constraints largely accounted for the structure of seven major networks. Our results imply that some networks and gradients may denote discrete and continuous representations of the same aspects of functional MRI data. We suggest that integrated explanations can reduce redundancy by avoiding the attribution of independent existence or function to these features.

Introduction

Networks and gradients represent two basic features of whole-brain activity. Networks (also known as systems or modules) denote discrete groups of brain regions that, by virtue of similar activity patterns, putatively facilitate specialized brain function (Damoiseaux et al., 2006; Smith et al., 2009; Yeo et al., 2011). Gradients denote spatially continuous variation in anatomy or activity that may reflect the outcome of developmental processes (Dong et al., 2021; Guell et al., 2018; Margulies et al., 2016). Advances in data acquisition (Van Essen et al., 2012) and analysis (Jenkinson et al., 2012; Vos de Wael et al., 2020) have allowed investigators to robustly and noninvasively detect these features in functional MRI data. These advances have enabled an extensive body of work centered on the structure of these features across healthy and diseased brain states (Zhang and Raichle, 2010; Huntenburg et al., 2018).

This body of work has used a diverse group of clustering and dimensionality reduction methods to define networks and gradients. Here, we adopted two popular definitions of these features. First, we used a popular parcellation (clustering) of whole-brain voxel correlation matrices to define networks (Schaefer et al., 2018). Second, we used diffusion embedding (dimensionality reduction) of whole-brain voxel correlations to define gradients (Margulies et al., 2016). Table 1 clarifies the use of these and other technical terms in the article.

Despite this extensive body of work, few studies have tested the statistical redundancy between networks and gradients. In theory, these two features could represent distinct outcomes of selective pressures and developmental constraints (Cembrowski and Menon, 2018). In practice,

however, the structure of networks and gradients typically shows considerable overlap. Specifically, seminal studies (Margulies et al., 2016; Bolt et al., 2022) have used sophisticated dimension-reduction methods to show strong statistical associations between networks and gradients. However, the correlational nature of these studies cannot disambiguate the presence of statistical redundancy between these features. Therefore, some of these studies have considered that networks and gradients are distinct. For example, Margulies et al. (2016) noted that “a principal gradient of cortical organization [...] is anchored at one end by [networks] implicated in perceiving and acting, and at the other end by [...] the default-mode network”. Other studies have suspended judgment on the relationship of networks and gradients. For example, Bolt et al. (2022) noted that “the primary aim of [their] study was descriptive, [and they] have avoided any explanatory or causal explanation”. On this basis, the statistical redundancy of networks and gradients remains an unsettled question.

Some redundancy between networks and gradients may be expected from our knowledge of approximate equivalences between k-means clustering and principal component analysis, canonical methods for clustering and dimensionality reduction (Ding and He, 2004; Drineas et al., 2004). Here, we developed numerical methods to test this redundancy more directly. Our approach is conceptually simple. First, we detected network and gradient structure in functional MRI data. Second, we sampled regional timeseries with network constraints, and evaluated the presence of gradient structure in these data. Third, we sampled regional timeseries with gradient constraints, and evaluated the presence of network structure in these data. This approach resembles a controlled

* Corresponding authors.

E-mail addresses: aditya.nanda@vanderbilt.edu (A. Nanda), mika.rubinov@vanderbilt.edu (M. Rubinov).

<https://doi.org/10.1016/j.neuroimage.2023.120110>.

Received 16 January 2023; Received in revised form 5 April 2023; Accepted 12 April 2023

Available online 5 May 2023.

1053-8119/© 2023 Published by Elsevier Inc. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>)

Table 1
Clarification of terms used in this article

Nullspace sampling	
Feature	Any numeric property of the data (e.g., network and gradient features).
Constraints	A set of features that specify a target data distribution.
Unbiased sampling	Selection of data samples with probability of the target data distribution.
Nullspace	A linear space that maps to the solution space of a linear system.
Network	
Concept	A group of brain regions that have similar activity patterns. This usage is standard in the neuroimaging literature but is distinct from the more general definition of a network.
Structure	Correlations between target network timeseries and all regional timeseries. Network timeseries denote the mean timeseries of all within-network regions.
Constraints	Mean correlations of regional timeseries within and between networks.
Gradient	
Concept	Spatially continuous variation of brain organization.
Structure	Diffusion-embedding components of interregional correlation matrices.
Constraints	Principal components of regional timeseries matrices.

experiment, or a randomized control trial (Siddiqi et al., 2022), that tests the effects of targeted interventions (constraints) on outcomes (features of interest). The approach goes beyond correlations because it can show that one feature is redundant with another (Rubinov, 2016) or that both features are approximately numerically equivalent. Such showings fall within a long tradition that emphasizes the importance of constraints and “spandrels”, nonfunctional or nonadaptive traits, in integrated explanations of biological structure (Gould and Lewontin, 1979).

This conceptually simple approach, however, is practically difficult because it requires the unbiased sampling of data with nontrivial constraints. Unbiased sampling is important because it allows to distinguish the constraints of interest from a wealth of confounding, potentially extraneous, explanations. However, to the best of our knowledge, the field currently lacks methods for unbiased sampling of regional timeseries with network or gradient constraints. By contrast, existing sampling methods come in three main forms. First, methods based on naive shuffling of timeseries are fast but destroy network and gradient structure. Second, methods based on sampling data with spatial autocorrelation constraints (Burt et al., 2020; Markello and Misisic, 2021; Shinn et al., 2023) are considerably more interesting but are not designed to constrain network or gradient structure. Moreover, the heuristic nature of these methods makes them susceptible to sampling bias. Third, methods such as autoregressive randomization (Zivot and Wang, 2006) or phase randomization (Prichard and Theiler, 1994) are perhaps most relevant to our study but have two important limitations. First, both methods trivially constrain all network and gradient structure and in this way cannot be used to test the redundancy between these features. Second, both methods assume a stationary, linear, and Gaussian generative model, and constrain lagged correlations. These additional features reflect properties of real functional MRI data, but also introduce confounding explanations that make it more difficult to perform controlled experiments. Liegeois et al. (2017, 2021) provides a thorough discussion of these issues.

Here, we built on insights from linear algebra to develop two related methods that sample timeseries with network or gradient constraints. In both cases, our main contributions was to first reduce the sampling problems to a sequence of linear systems, and then use the nullspace to sample solutions to these systems. Our methods build on a rich literature for solving linear inverse problems across a wide range of scientific domains (Smith, 1984; Tarantola, 2005; Van den Meersche et al., 2009). The methods are unbiased insofar as they accurately sample timeseries from the target data distribution, and efficient insofar as they scale to multiregional recordings.

In the next sections we describe the details of these methods, and use these methods to show considerable redundancy between popular definitions of network and gradient structure in functional MRI data. We conclude by discussing the implications of these methods and results for future work.

Results

Definition of networks and gradients

We analyzed resting-state functional MRI recordings from subjects in the Human Connectome Project (Van Essen et al., 2013). We used a popular data-driven parcellation (Schaefer et al., 2018) to extract timeseries from 400 cortical regions in these recordings, and performed all our subsequent analyses on regional timeseries matrices. In this section we summarize our definitions of network and gradient structure and constraints.

Network structure and constraints. We defined network structure as the correlations between a network timeseries and the timeseries of all brain regions, averaged over all subjects. We used a popular division of the cortex into the visual, somatomotor, temporoparietal, dorsal attention, ventral attention, control, default, and limbic networks (Schaefer et al., 2018), as well as a hierarchical subdivision of these 8 networks into 34 (17 bilaterally symmetric) subnetworks. We defined network constraints as the mean interregional correlations within and between these subnetworks. Below, we evaluated the extent to which these constraints accounted for the structure of 7 networks (we did not present results on the limbic network because of known problems with low signal-to-noise ratio in that network).

Gradient structure and constraints. We defined gradient structure as diffusion-embedding components of interregional correlations, averaged over all subjects. We followed the pipeline of Margulies et al. (2016) to first nonlinearly transform the mean interregional correlation matrix to a Markov chain matrix, and then extract diffusion-embedding components as the leading eigenvectors of this matrix. We defined gradient constraints as the principal components of regional timeseries matrices (Hong et al., 2020), and evaluated the extent to which these constraints accounted for gradient structure. Finally, we used Procrustes analysis to align gradients across samples (Langs et al., 2015), and also explored the effect of this alignment on our results.

Overview of nullspace sampling

We built on insights from linear algebra to sample regional timeseries with network or gradient constraints. In this section we summarize the general formulation of our method, and its several variants. In the Methods section we describe our approach in considerable mathematical detail.

Variant 0: General formulation (Box 1). The general formulation of our method samples the solutions to a system of linear equations $\mathbf{Ax} = \mathbf{b}$, where $\mathbf{x}$ is a vector that denotes empirical data, $\mathbf{A}$ is a matrix that defines features of interest, and $\mathbf{b}$ is a vector that denotes empirical values of these features. Our method samples solution vectors $\tilde{\mathbf{x}}$ that satisfy empirical constraints, such that $\mathbf{A}\tilde{\mathbf{x}} = \mathbf{b}$ where $\tilde{\mathbf{x}}$ denotes model (rather than empirical) data. We can express these solutions as $\tilde{\mathbf{x}} = \mathbf{x}^* + d\mathbf{Zq}$,

where $\mathbf{x}^*$ is the unique minimum-norm solution, d is a scaling constant, $\mathbf{q}$ is a uniformly sampled unit-norm weighting vector, and $\mathbf{Z}$ is an orthonormal basis of the nullspace of $\mathbf{A}$, such that $\mathbf{Z}^\top \mathbf{Z} = \mathbf{I}$ and $\mathbf{A}\mathbf{Z} = \mathbf{0}$ (Van den Meersche et al., 2009).

Box 1. General formulation of nullspace sampling.

1. Define a system of linear equations, $\mathbf{A}\mathbf{x} = \mathbf{b}$.
2. Compute the basis of the nullspace $\mathbf{Z}$, the minimum-norm solution $\mathbf{x}^*$, and the scaling parameter d .
3. Uniformly sample weighting vector $\mathbf{q}$ to sample solutions $\tilde{\mathbf{x}} = \mathbf{x}^* + d\mathbf{Z}\mathbf{q}$.

Variant 1: Principal-component constraints. We used a sequential variant of our general formulation to sample correlation matrices constrained by a subset of empirical principal components.

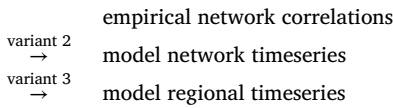
Variant 2: Pairwise-correlation constraints. We used a sequential variant of our general formulation to sample timeseries matrices constrained by empirical pairwise correlations. In the Methods section, we discuss the relationship between this variant and autoregressive or phase randomization.

Variant 3: Global mean and norm constraints. We used a sequential variant of our general formulation to sample timeseries matrices constrained by the mean and norm of empirical timeseries.

Sampling timeseries with network and gradient constraints

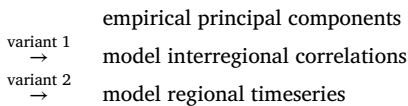
We combined the above variants into pipelines for sampling timeseries with network or gradient constraints. We summarize these pipelines below, and illustrate them in Figure 1.

Network constraints. We defined network constraints as the mean correlations of regional timeseries within and between networks. We constrained these correlations by adopting the following pipeline (Figure 1C):



We considered two models with network constraints. First, the intra-network model constrained only the mean interregional correlations within a network. Second, the all-network model constrained the mean interregional correlations within and between all networks. Therefore, for a system with l networks, the intra-network model had l constraints, while the all-network model had $\frac{1}{2}l(l+1)$ constraints (l intra-network constraints and $\frac{1}{2}l(l-1)$ inter-network constraints).

Gradient constraints. We defined gradient constraints as principal components of regional timeseries matrices (Hong et al., 2020). We sampled regional timeseries with these constraints by adopting the following pipeline (Figure 1C):



We considered two models with gradient constraints. First, the one-gradient model constrained only the first principal component. Second, the two-gradient model constrained the first and second principal components. Both models also constrained all eigenvalues of the empirical correlation matrix. Therefore, for a system with n regions, the k -gradient model had $(k+1)n$ empirical constraints (kn eigenvector constraints and n eigenvalue constraints).

Model complexity. Our regional timeseries comprised 400 regions and 1200 timepoints, and were divided into 34 networks. Therefore, the one-network model had 34 constraints (0.01% of all data points),

the all-network model had 595 constraints (0.12%) constraints, the one-gradient model had 800 constraints (0.17%), and the two-gradient model had 1200 constraints (0.25%). The total number of constraints was most comparable for the all-network models and one-gradient models. More generally, this relatively small number of constraints in all models implied that all model and empirical data were essentially uncorrelated (Figure S1).

Model performance

This section describes our main results, summarized as correlations between network or gradient structure of empirical and model timeseries. We first evaluated the extent to which network constraints accounted for network and gradient structure. We then similarly evaluated the extent to which gradient constraints accounted for gradient and network structure.

Network constraints (Figure 2). Our simplest intra-network model provided a relatively coarse, albeit generally accurate, representation of network structure (with median [95% uncertainty interval] model-data Pearson correlations of 0.74 [0.63, 0.79] across all networks). However, this model largely failed to recapitulate gradient structure (with corresponding model-data correlations of 0.51 [0.41, 0.60] for the first gradient, 0.07 [0.02, 0.20] for the second gradient, and 0.19 [0.07, 0.23] for the third gradient). Overall, these results suggest that intra-network correlations alone were not sufficient to fully account for network or gradient structure.

By contrast, the all-network model provided much better representations of all network structure (with median [95% uncertainty interval] model-data Pearson correlations of 0.89 [0.84, 0.91] across all networks). More interestingly, this model also accurately represented gradient structure (with corresponding model-data correlations of 0.93 [0.93, 0.93] for the first gradient, 0.94 [0.93, 0.94] for the second gradient, and 0.82 [0.82, 0.82] for the third gradient). Overall, these results suggest that inter-network correlations were largely sufficient to account for gradient structure.

Gradient constraints (Figure 3). The one-gradient model had a similar number of constraints to the all-network model, and performed similarly well to that model. Specifically, this model provided accurate representations of first-gradient structure (with median [95% uncertainty interval] model-data Pearson correlations of 0.96 [0.96, 0.96]), and reasonable approximations of second- and third-gradient structure (with corresponding model-data correlations of 0.63 [0.62, 0.65], and 0.73 [0.72, 0.75]). More interestingly, this model provided a highly accurate representation of all network structure (with corresponding model-data correlations of 0.89 [0.79, 0.99] across all networks). Overall, these results suggest that primary-gradient constraints alone were largely sufficient to account for network structure.

Finally, our most complex two-gradient model provided additional improvements in representation of all gradient structure (with median [95% uncertainty interval] model-data Pearson correlations of 0.90 [0.83, 0.98] across all gradients), and all network structure (with corresponding model-data correlations of 0.94 [0.87, 0.99] across all networks).

Control experiments. We evaluated the effects of preprocessing and analysis methods on our results. First, we considered the effects of global signal regression (Figures S2 and S3). As we discuss in the Methods section, we adopted this step to remove the effects of vigilance and non-neuronal physiology, and correspondingly to better align principal components (gradient constraints) with diffusion-embedding components (gradient structure). The exclusion of global signal regression considerably worsened the performance of both gradient models, and the one-network model, but had little effect on the performance of the all-network model.

Second, we considered the effect of Procrustes alignment on our results (Figures S4 and S5). This alignment tends to increase the observed similarity of gradient structure. The exclusion of Procrustes alignment

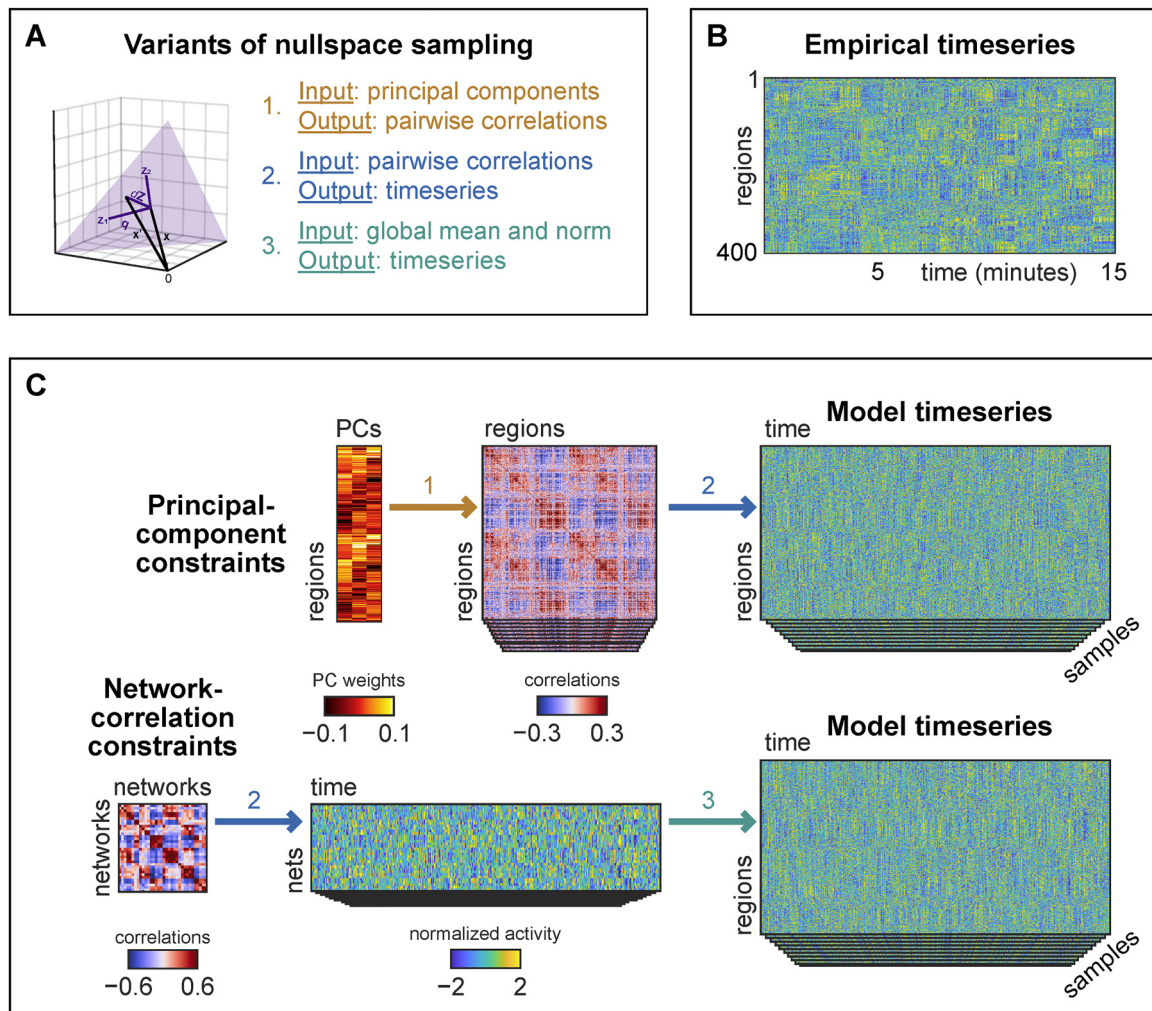


Fig. 1. Overview of nullspace sampling pipelines. (A) Three variants of nullspace sampling. (B) Empirical timeseries matrix. (C) Pipelines for sampling timeseries with network or gradient constraints. Arrows denote distinct variants of nullspace sampling (A). The principal-component pipeline first uses variant 1 to generate interregional correlation matrices constrained by empirical principal components, and then uses variant 2 to generate model timeseries constrained by these interregional correlations. The network-correlation pipeline first uses variant 2 to generate model network timeseries and then uses variant 3 to generate model regional timeseries constrained by these network timeseries.

likewise considerably worsened the performance of both gradient models, and the one-network model, but had little effect on the performance of the all-network model. By definition, the exclusion of this step only affected the representation of gradient structure (i.e., it had no effect on the representation of network structure).

Discussion

Summary. We developed methods for sampling timeseries with network or gradient constraints. We first validated our approach by showing that network constraints largely accounted for a parcellation-based definition of network structure, and gradient constraints largely accounted for a diffusion-embedding based definition of gradient structure. We then noted that network constraints also largely accounted for gradient structure, while gradient constraints also largely accounted for network structure. Specifically, we found that the all-network and one-gradient models had similar complexity, and induced relatively similar network and gradient structure (although the one-gradient model was more sensitive to changes in preprocessing methodology, Figures S2-S5). We also found that a simpler (intra-network) model considerably reduced the similarity between empirical and model data, while a more complicated (two-gradient) model somewhat increased this similarity.

Implications. Our results suggest that these popular definitions of gradient and network structure show considerable redundancy, and may simply denote discrete and continuous representations of the same aspects of functional MRI data. The strong similarity between these distinctly defined features has an intuitive technical explanation. Specifically, both networks and gradients are extracted from interregional-correlation matrices, and both network and gradient constraints represent low-dimensional approximations of these same matrices (Figure 1). It follows, therefore, that accurate representations of interregional-correlation matrices (whether due to network or gradient constraints) will simultaneously recapitulate both network and gradient structure in model data.

Our demonstration of redundancy goes beyond previously reported correlations between networks and gradients (Margulies et al., 2016; Guell et al., 2018; Raut et al., 2021; Vos de Wael et al., 2021; Dong et al., 2021; Bolt et al., 2022). This demonstration implies that the field is not justified to assume the independent importance of all networks and all gradients, much in the same way that one is not justified to assume the importance of redundant regressors in a linear model (Rubinov, 2022). We suggest that future studies need to combine additional data on the functional relevance (Krakauer et al., 2017), evolutionary ancestry (Cisek and Hayden, 2022), and developmental mech-

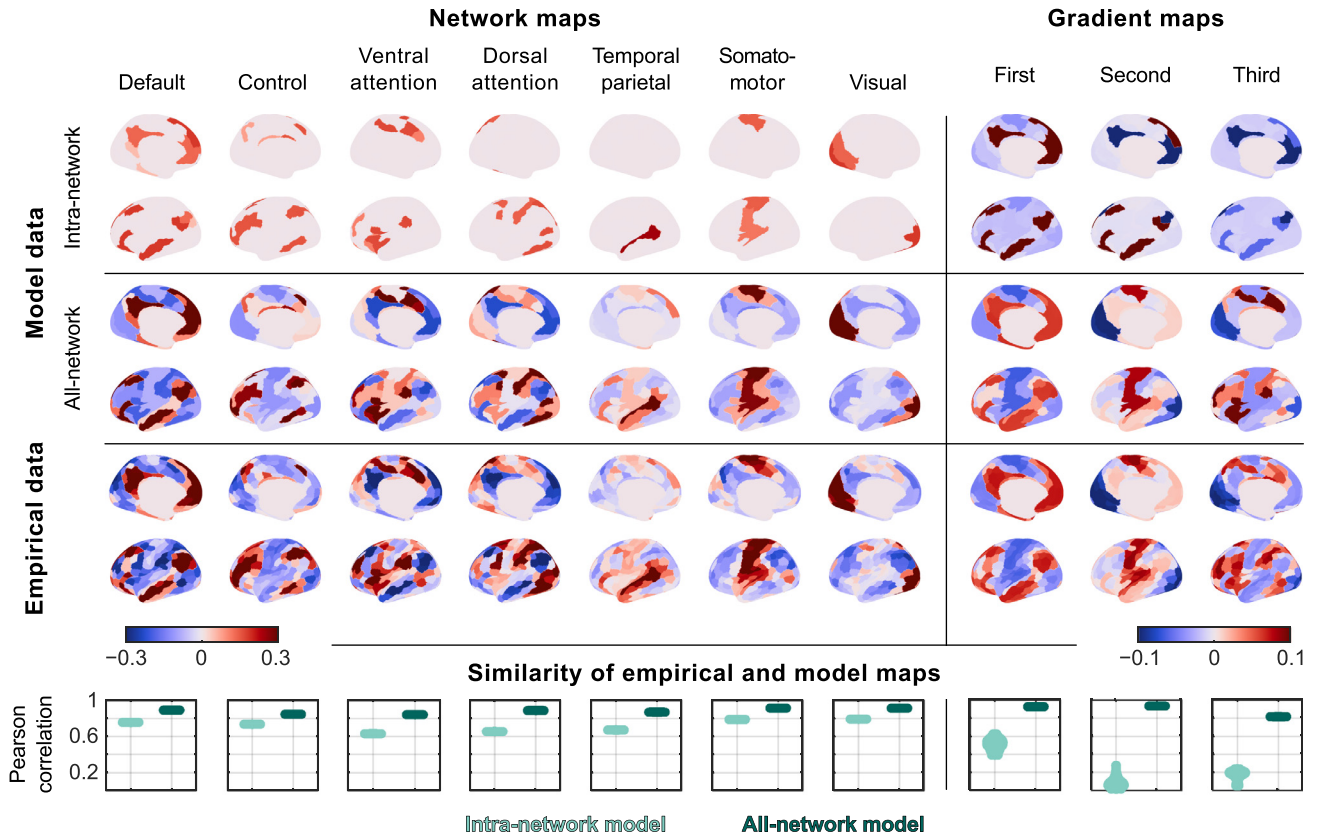


Fig. 2. Effects of network constraints. Left: maps of network structure (correlations between the timeseries of a target network and all brain regions). Right: maps of gradient structure (diffusion-embedding components of regional timeseries matrices). First row: maps of intra-network model data. Second row: maps of all-network model data. Third row: Maps of empirical data. Bottom row: violin plots of Pearson correlation coefficients between empirical and model data. Network and gradient structure were computed from correlation matrices averaged over all 100 subjects. Model data comprised 100 such mean correlation matrices (for a total of 10,000 sampled timeseries matrices).

anisms (Cembrowski and Menon, 2018) of specific networks and gradients, in order to rigorously validate the inclusion of these features in unified explanations of brain organization.

Limitations and future work. Our methods have two main limitations. First, the methods can be slow and memory-intensive. This is not a practical limitation for parcel-resolution timeseries, but can be a practical limitation for voxel-resolution timeseries. Second, the methods do not admit additional constraints, such as spatial or temporal autocorrelations, or diffusion-embedding components. This second set of limitations is rather subtle, and worth additional consideration.

Our methods do not admit spatial or temporal autocorrelation constraints. This may be an important limitation because neuroimaging data are dominated by autocorrelation structure. However, in theory, the unbiasedness of our methods guarantees the sampling of data with all empirical structure, including empirical autocorrelations. In other words, a sufficiently large number of data samples is bound to include samples with empirical autocorrelations. Nonetheless, in practice, our methods do not sample data with empirical autocorrelations because these data form a negligible fraction of our target distributions. It is possible, therefore, that our results on these data may differ from our main results.

Separately from these considerations, our methods do not directly constrain diffusion-embedding components. This limitation may be important if one is interested in exploring the subtle differences between gradients defined with diffusion embedding, and gradients defined with principal component analysis.

Our future work will focus on resolving some of these limitations, through improvements in scalability, and through expansion of sampling variants to admit other constraints. More generally, as neuroimaging data continue to increase in size and complexity, the ability to sam-

ple these data with a rich set of spatial, temporal, correlational, spectral and other structure will become increasingly important for delineating common principles of brain organization. We hope that our methods will help fulfill an important part of this increasing need.

Methods

Nullspace sampling methodology

We sampled regional timeseries with network-correlation constraints, or with principal-component constraints. In this section we describe the details of our nullspace sampling methodology. Our MATLAB software implements these methods and is freely available at: <https://github.com/AdityaNanda/Networks-Gradients-Sampling-Toolbox>.

Variant 0: General formulation

The general formulation of our method leverages insights from linear algebra to uniformly sample data constrained by sets of predetermined features (Figure 4). Consider a linear system

$$\mathbf{A}\mathbf{x} = \mathbf{b}, \quad (1)$$

where $\mathbf{x} \in \mathbb{R}^{I \times 1}$ is some empirical data vector, $\mathbf{A} \in \mathbb{R}^{m \times I}$ is a matrix that encodes m features of interest, and $\mathbf{b} \in \mathbb{R}^{m \times 1}$ denotes empirical values of these features. Let us assume, without loss of generality, that the m features of interest are linearly independent or, equivalently, that the matrix $\mathbf{A}$ has rank m .

Our method uniformly samples vectors $\tilde{\mathbf{x}}$ that match empirical features of interest, such that

$$\mathbf{A}\tilde{\mathbf{x}} = \mathbf{b},$$

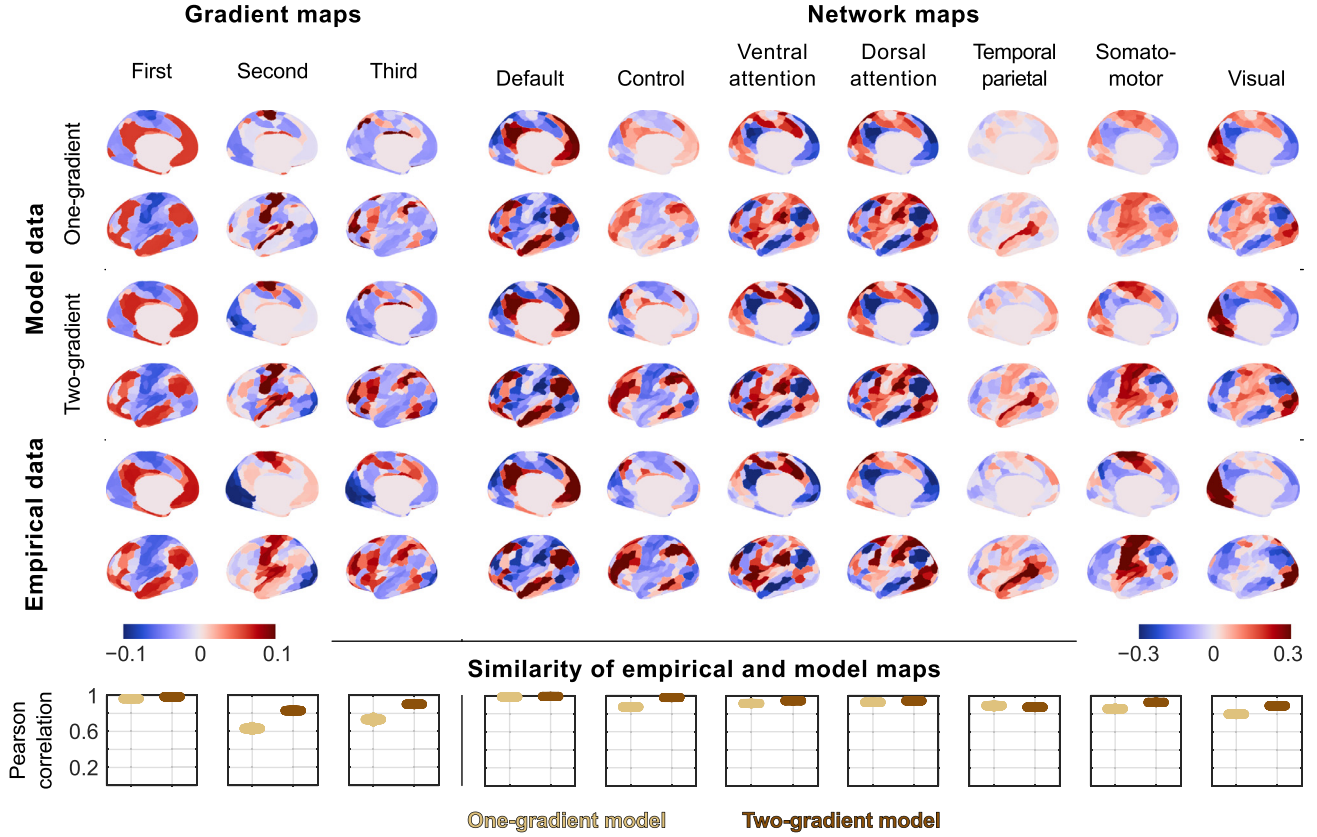


Fig. 3. Effects of gradient constraints. Left: maps of gradient structure (correlations between the timeseries of a target network and all brain regions). Right: maps of network structure (diffusion-embedding components of regional timeseries matrices). First row: maps of one-gradient model data. Second row: maps of two-gradient model data. Third row: Maps of empirical data. Bottom row: violin plots of Pearson correlation coefficients between empirical and model data. Gradient and network structure were computed from correlation matrices averaged over all 100 subjects. Model data comprised 100 such mean correlation matrices (for a total of 10,000 sampled timeseries matrices).

where the $\tilde{\cdot}$ operator denotes model (rather than empirical) vectors.

Note that the number of features m will in general be much smaller than the number of data elements t . This implies that the system in Equation 1 has infinitely many solutions. The solution space of this system maps to the nullspace of the matrix $\mathbf{A}$, $\mathbf{Z} \in \mathcal{R}^{t \times (t-m)}$, such that $\mathbf{Z}^T \mathbf{Z} = \mathbf{I}$ and $\mathbf{A}\mathbf{Z} = \mathbf{0}$, where $\mathbf{0}$ is a matrix of zeros (Laub, 2005; Van den Meersche et al., 2009). Let us now define $\tilde{\mathbf{x}}$ as points on this solution space,

$$\tilde{\mathbf{x}} = \mathbf{x}^* + d\mathbf{Z}\mathbf{q}, \quad (2)$$

where $\mathbf{x}^*$ is the unique minimum-norm solution, d is a scaling constant, and $\mathbf{q} \in \mathcal{R}^{(t-m) \times 1}$ is a unit-norm weighting vector. Geometrically, each row in $\mathbf{A}$ represents an unbounded hyperplane in t dimensions, and the solution space is an $t - m$ vector space formed by the intersection of these hyperplanes.

We sampled $\tilde{\mathbf{x}}$ in two steps. Figure 4 summarizes these steps, and here we discuss each step in more detail. First, we computed the nullspace matrix $\mathbf{Z}$, and the minimum-norm solution $\mathbf{x}^* = \mathbf{A}^\dagger \mathbf{b}$, where $\dagger$ denotes the Moore-Penrose pseudoinverse (Laub, 2005). We also computed the scaling parameter d to enforce additional, problem-specific constraints. In practice, we set d to restrict our sampling to all $\tilde{\mathbf{x}}$ with some pre-determined Euclidean norm e , $\|\tilde{\mathbf{x}}\| = e$. Note that, by the fundamental theorem of linear algebra, the minimum norm solution $\mathbf{x}^*$ is orthogonal to all column vectors in $\mathbf{Z}$. Because the additive components in Equation 2 are orthogonal, and because $\|\mathbf{q}\| = 1$, we used Pythagoras theorem to set $d = \sqrt{e^2 - \|\mathbf{x}^*\|^2}$ and thereby ensure that $\|\tilde{\mathbf{x}}\| = e$.

Second, we uniformly sampled weighting vectors $\mathbf{q}$ from the $n - m$ dimensional standard normal distribution, and rescaled these vectors to have unit norm. This sampling approach guarantees to produce

uniformly distributed random samples of $\mathbf{q}$ (Smith, 1984). Geometrically, this procedure is equivalent to sampling from the surface of the $t - m$ dimensional hypersphere that has unit radius and is centered at the origin.

Let us now show that our sampling of $\tilde{\mathbf{x}}$ is uniform. First, let us express the solution space in Equation 2 solely as a function of the weighting vector, such that $\tilde{\mathbf{x}} = f(\mathbf{q})$. Then, let $P(\tilde{\mathbf{x}})$ and $P(\mathbf{q})$ denote probability distributions over $\tilde{\mathbf{x}}$ and $\mathbf{q}$ respectively. Following Papoulis (1965) we can write

$$P(\tilde{\mathbf{x}}) = \frac{P(\mathbf{q})}{\left\| \frac{\partial \tilde{\mathbf{x}}}{\partial \mathbf{q}} \right\|} = \frac{P(\mathbf{q})}{d \|\mathbf{Z}\|} = \frac{1}{d} P(\mathbf{q}). \quad (3)$$

This equation demonstrates that the probability density functions $P(\mathbf{q})$ and $P(\tilde{\mathbf{x}})$ are related by the scaling constant d and implies that a uniform sampling of $\mathbf{q}$ guarantees a uniform sampling of $\tilde{\mathbf{x}}$.

This completes the general formulation of our method. In the following sections, we describe three variants of this formulation. In the Results section, we used combinations of these variants to sample regional timeseries with network or gradient constraints (Figure 1).

Variant 1: Correlation matrices constrained by principal-component structure

We first built on our basic formulation to sample correlation matrices constrained by k principal components. Let us denote a matrix of normalized (zero-mean, unit-norm) regional timeseries by $\mathbf{X} \in \mathcal{R}^{t \times n}$, and the corresponding correlation matrix by $\mathbf{C} \in \mathcal{R}^{n \times n}$. Note that $\mathbf{C} = \mathbf{X}^T \mathbf{X} = \mathbf{V} \mathbf{\Lambda} \mathbf{V}^T$ where $\mathbf{\Lambda} \in \mathcal{R}^{n \times n}$ denotes the eigenvalue matrices of $\mathbf{C}$, and $\mathbf{V} \in \mathcal{R}^{n \times n}$ denotes the eigenvectors of $\mathbf{C}$ or, equivalently, the

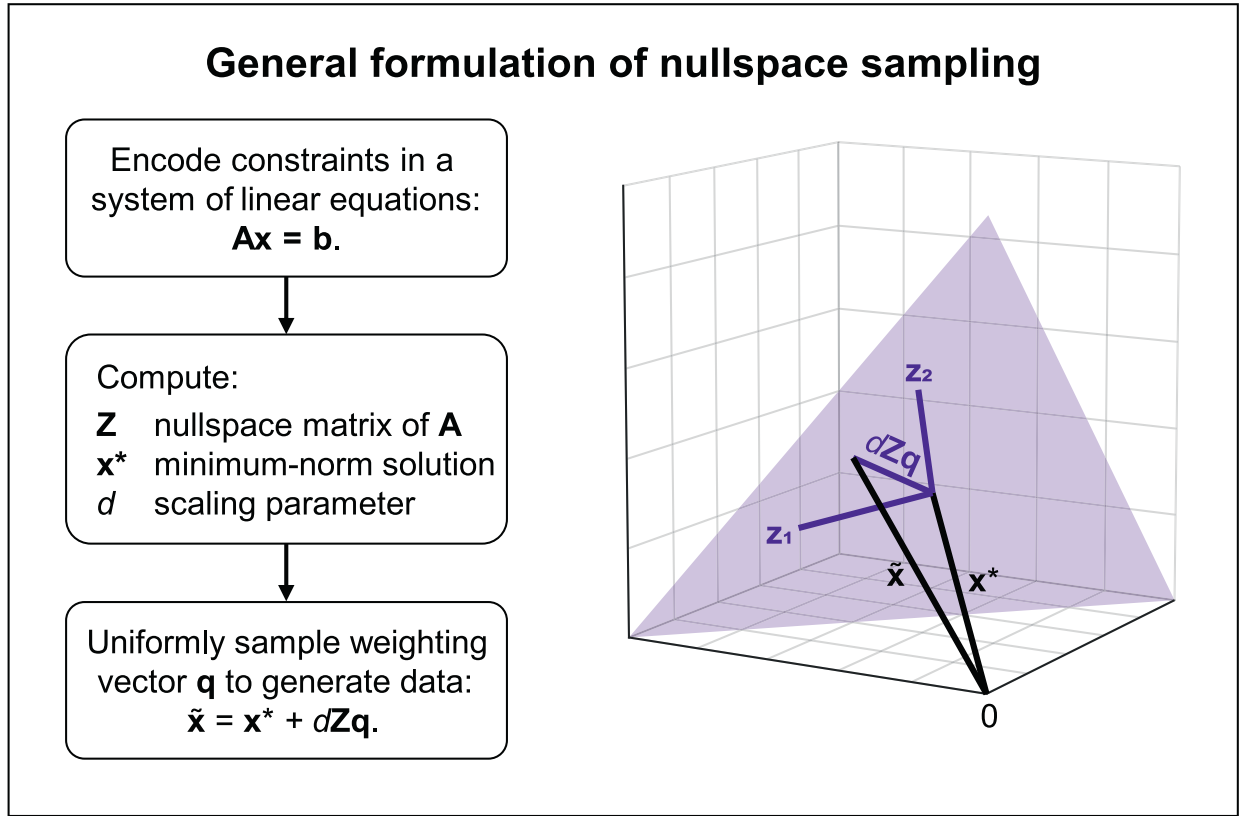


Fig. 4. General formulation of nullspace sampling. Left: a flowchart of the main steps. Right: a geometric representation of the method for a three-dimensional linear system with one feature of interest ($Ax = b$ where $A \in \mathbb{R}^{1 \times 3}$, $x \in \mathbb{R}^{3 \times 1}$, and $b \in \mathbb{R}^{1 \times 1}$). The purple plane represents the solution space of this system. This plane is affinely spanned by orthonormal vectors z_1 and z_2 that form the nullspace matrix $Z = [z_1 \ z_2] \in \mathbb{R}^{3 \times 2}$. Our method uniformly samples points $\tilde{x}$ on this plane by expressing each point as a sum of vectors x^* and dZq . The vector x^* denotes the minimum norm solution, a point on the plane with the shortest Euclidean distance to the origin. The vector dZq denotes a random linear combination of the nullspace vectors. See the text for details.

principal components of X . Here, we sought to sample matrices $\tilde{V}$ that shared the k leading empirical eigenvectors with V .

We sampled these matrices using a sequential variant of our general formulation. Specifically, we sequentially sampled $n - k$ unit vectors orthogonal to the k empirical eigenvectors, $v_1, \dots, v_k$, as well as to any previously sampled eigenvectors. We first sampled $\tilde{v}_1$ such that $[v_1 \ \dots \ v_k]^T \tilde{v}_1 = 0$ and $\|\tilde{v}_1\| = 1$. We then sampled $\tilde{v}_i$ ($i = 2, \dots, n - k$), such that

$$[v_1 \ \dots \ v_k \ \tilde{v}_1 \ \dots \ \tilde{v}_{i-1}]^T \tilde{v}_i = 0 \text{ and } \|\tilde{v}_i\| = 1 \quad (4)$$

where $v_1, \dots, v_k$ are the empirical k leading eigenvectors. The structure of the coefficient matrix in Equation 4 guarantees that $\tilde{V} = [v_1 \ \dots \ v_k \ \tilde{v}_1 \ \dots \ \tilde{v}_{n-k}]$ forms a set of unit eigenvectors. We then used these vectors to define a correlation matrix $\tilde{C} = \tilde{V}\Lambda\tilde{V}^T$. The uniqueness of eigendecomposition guarantees that $\tilde{C}$ is constrained to have the desired principal component structure.

Variant 2: Timeseries matrices constrained by pairwise-correlation structure

We next sought to sample matrices of normalized (zero-mean, unit-norm) regional timeseries $\tilde{X} = [\tilde{x}_1 \ \dots \ \tilde{x}_n]$ such that $\tilde{C} = \tilde{X}^T \tilde{X} = X^T X = C$. We sampled these timeseries using a similar sequential variant of our general formulation. Specifically, we first sampled $\tilde{x}_1$ to have zero mean and unit norm, and then sampled $\tilde{x}_i$ ($i = 2, \dots, n$), such that

$$[\mathbf{1} \ \tilde{x}_1 \ \dots \ \tilde{x}_{i-1}]^T \tilde{x}_i = [0 \ c_{i,1} \ \dots \ c_{i,i-1}] \text{ and } \|\tilde{x}_i\| = 1 \quad (5)$$

where $\mathbf{1}$ is a vector of ones. The structure of the coefficient matrix in Equation 5 guarantees that $\tilde{X} = [\tilde{x}_1 \ \dots \ \tilde{x}_n]$ will have the desired correlation matrix C .

In practice, we implemented this variant by leveraging the eigendecomposition $C = V\Lambda V^T$. Specifically, we first sampled timeseries $\tilde{X}'$ constrained to have $(\tilde{X}')^T \tilde{X}' = \Lambda$, and then obtained our timeseries of interest via a rotation, $\tilde{X} = \tilde{X}' V^T$. Note that $\tilde{X}$ will have the desired correlation matrix C because

$$\tilde{C} = \tilde{X}^T \tilde{X} = (\tilde{X}' V^T)^T (\tilde{X}' V^T) = V \tilde{X}'^T \tilde{X}' V^T = V \Lambda V^T = C.$$

This approach guarantees that the minimum-norm solution will be $x^* = 0$, and therefore does not require the computation of this solution at every step. Moreover, the orthonormal transformation V^T preserves probability distributions (Equation 3), and in this way guarantees that the uniform sampling of $\tilde{X}'$ leads to the uniform sampling of $\tilde{X} V^T$.

Note that this variant is broadly similar to autoregressive or phase randomization. However, unlike these methods, the variant does not assume a stationary, linear, and Gaussian generative model, and does not constrain lagged correlations. The lack of these assumptions allows us to test the effect of instantaneous correlations in a controlled way. By contrast, the presence of these assumptions introduces confounding or extraneous explanations, even if it increases the similarity between model and empirical data.

Variant 3: Timeseries matrices constrained by global mean and norm structure

We finally sought to sample timeseries matrices constrained by global mean and norm structure. We first considered empirical activity vectors at time i , $y_i \in \mathbb{R}^{n \times 1}$, and then sequentially sampled model activity vectors $\tilde{y}_i$ such that

$$\mathbf{1}^T \tilde{y}_i = \mathbf{1}^T y_i \text{ and } \|\tilde{y}_i\| = \|y_i\|. \quad (6)$$

Efficient computation of nullspace matrices

We optimized our sampling by replacing the slow direct computation of nullspace matrices $\mathbf{Z}$ with fast sequential or analytical computation of these matrices. Let us now describe the details of this computation.

Sequential computation of nullspace matrix. We developed a sequential computation of the nullspace matrix for our sampling variants 1 and 2 Equations 4 and (5). Let $\mathbf{A}_i$ denote the coefficient matrix at step i . A sequential variant of our method allowed us to express $\mathbf{A}_{i+1}$ as

$$\mathbf{A}_{i+1} = \begin{bmatrix} \mathbf{A}_i \\ \tilde{\mathbf{x}}_i^\top \end{bmatrix} = \begin{bmatrix} \mathbf{A}_i \\ (\mathbf{x}_i^* + d_i \mathbf{Z}_i \mathbf{q}_i)^\top \end{bmatrix},$$

where $\tilde{\mathbf{x}}_i$ denotes our sampled vector at step i . This, in turn, allowed us to express $\mathbf{Z}_{i+1}$ as

$$\mathbf{Z}_{i+1} = \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) \quad (7)$$

where $\text{null}(\mathbf{q}_i^\top)$ is the nullspace matrix of the vector $\mathbf{q}_i^\top$. We then efficiently computed $\text{null}(\mathbf{q}_i)$ using the Householder transformation (Golub and Van Loan, 2013; Trefethen and Bau, 2022).

Let us now prove that $\mathbf{Z}_{i+1}$ is indeed the nullspace of matrix $\mathbf{A}_{i+1}$. We can do this by showing that $\mathbf{Z}_{i+1}$ is an orthonormal matrix and that $\mathbf{A}_{i+1} \mathbf{Z}_{i+1} = \mathbf{0}$. First, let us note that $\mathbf{Z}_{i+1}$ is the product of two orthonormal matrices, $\mathbf{Z}_i$ and $\text{null}(\mathbf{q}_i^\top)$, and is therefore itself an orthonormal matrix (Laub, 2005). Second, let us note that

$$\begin{aligned} \mathbf{A}_{i+1} \mathbf{Z}_{i+1} &= \begin{bmatrix} \mathbf{A}_i \\ (\mathbf{x}_i^* + d_i \mathbf{Z}_i \mathbf{q}_i)^\top \end{bmatrix} [\mathbf{Z}_i \text{null}(\mathbf{q}_i^\top)] \\ &= \begin{bmatrix} \mathbf{A}_i \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) \\ (\mathbf{x}_i^*)^\top \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) + d_i \mathbf{q}_i^\top \mathbf{Z}_i^\top \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) \end{bmatrix}. \end{aligned}$$

Now, $\mathbf{A}_i \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) = \mathbf{0}$ because $\mathbf{A}_i \mathbf{Z}_i = \mathbf{0}$. Likewise, $(\mathbf{x}_i^*)^\top \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) = \mathbf{0}$ because the minimum norm solution is orthogonal to the nullspace. Finally, $d_i \mathbf{q}_i^\top \mathbf{Z}_i^\top \mathbf{Z}_i \text{null}(\mathbf{q}_i^\top) = \mathbf{0}$ because $\mathbf{Z}_i$ is orthonormal and $\mathbf{q}_i^\top$ is orthogonal to its nullspace. Putting this together, we find that $\mathbf{A}_{i+1} \mathbf{Z}_{i+1} = \mathbf{0}$. This completes our proof.

Analytical solution of nullspace matrix. We derived an analytical expression for the nullspace matrix $\mathbf{Z} \in \mathcal{R}^{n \times (n-1)}$ in variant 3 (Equation 6). Specifically, the nullspace matrix of $\mathbf{A} = \mathbf{1}^\top$ is given by

$$\mathbf{Z} = \begin{bmatrix} -\gamma & -\gamma & -\gamma & \dots & -\gamma \\ \beta & -\alpha & -\alpha & \dots & -\alpha \\ -\alpha & \beta & -\alpha & \dots & -\alpha \\ -\alpha & -\alpha & \beta & \dots & -\alpha \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -\alpha & -\alpha & \dots & -\alpha & \beta \end{bmatrix},$$

where α , β and γ denote solutions to the following set of equations:

$$\begin{aligned} -(n-2)\alpha + \beta - \gamma &= 0, \\ (n-2)\alpha^2 + \beta^2 + \gamma^2 - 1 &= 0, \quad \text{and} \\ (n-3)\alpha^2 - 2\alpha\beta + \gamma^2 &= 0. \end{aligned}$$

One can directly verify that this linear system has a unique solution which ensures that $\mathbf{Z}$ is an orthonormal nullspace of $\mathbf{A} = \mathbf{1}^\top$. Specifically, the first equation above ensures that $\mathbf{AZ} = \mathbf{0}$, while the second and third equations ensure that $\mathbf{Z}^\top \mathbf{Z} = \mathbf{I}$.

Data selection and preprocessing

We analyzed 15-minute resting-state functional MRI recordings from 100 healthy subjects of the Human Connectome Project. We chose to analyze 100 recordings with the lowest-available head movement. All our recordings had relative root-mean-square head movements of less than 0.2 at all timepoints. We used data processed with standard Human Connectome Project methods: the minimal preprocessing pipeline (Glasser et al., 2013), MSM-All registration (Robinson et al., 2014), and

ICA-FIX denoising (Salimi-Khorshidi et al., 2014). In addition, we regressed out the six motion parameters, their derivatives, as well as CSF, white matter, and global signals from these timeseries.

Global signal regression is a controversial step (Murphy and Fox, 2017) that warrants additional discussion. The global signal can reflect a mix of structured artifact (Power et al., 2017), but also information about vigilance and non-neuronal physiology (Liu et al., 2017). While the latter information may be of interest in some studies, here we sought to remove this signal to better align principal components and gradients. In functional MRI data, the first principal component tends to strongly reflect the global signal (He and Liu, 2012), while the second principal component tends to reflect the primary gradient (Bolt et al., 2022). These considerations suggest that global signal regression removed information about vigilance and non-neuronal physiology, and in this way aligned the leading principal component and the primary gradient. We directly checked the effect of this step by repeating our analysis without global signal regression.

Credit author statement

AN and MR conceived and planned the project. AN developed nullspace-sampling variants and performed analyses with assistance from MR. AN and MR interpreted results and wrote the manuscript. MR supervised the project.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data and code availability statement

All data are openly available from the Human Connectome Project. MATLAB software that implements these methods is freely available on GitHub, at <https://github.com/AdityaNanda/Networks-Gradients-Sampling-Toolbox>.

Acknowledgments

MR was supported by NIH 1RF1MH125933 and NSF 2207891. Data were provided by the Human Connectome Project, WU-Minn Consortium (Principal Investigators: David Van Essen and Kamil Ugurbil; 1U54MH091657) funded by the 16 NIH Institutes and Centers that support the NIH Blueprint for Neuroscience Research; and by the McDonnell Center for Systems Neuroscience at Washington University.

Supplementary material

Supplementary material associated with this article can be found, in the online version, at doi:[10.1016/j.neuroimage.2023.120110](https://doi.org/10.1016/j.neuroimage.2023.120110).

References

- Bolt, T., Nomi, J.S., Bzdok, D., Salas, J.A., Chang, C., Thomas Yeo, B., Uddin, L.Q., Keilholz, S.D., 2022. A parsimonious description of global functional brain organization in three spatiotemporal patterns. *Nature Neurosci.* 25 (8), 1093–1103.
- Burt, J.B., Helmer, M., Shinn, M., Anticevic, A., Murray, J.D., 2020. Generative modeling of brain maps with spatial autocorrelation. *NeuroImage* 220, 117038.
- Cembrowski, M.S., Menon, V., 2018. Continuous variation within cell types of the nervous system. *Trends Neurosci.* 41 (6), 337–348.
- Damoiseaux, J.S., Rombouts, S., Barkhof, F., Scheltens, P., Stam, C.J., Smith, S.M., Beckmann, C.F., 2006. Consistent resting-state networks across healthy subjects. *Proc. Natl. Acad. Sci.* 103 (37), 13848–13853.
- Ding, C., He, X., 2004. K-means clustering via principal component analysis. *Proc. 21st Int. Conf. Mach. Learn.*, p. 29.
- Dong, H.-M., Margulies, D.S., Zuo, X.-N., Holmes, A.J., 2021. Shifting gradients of macroscale cortical organization mark the transition from childhood to adolescence. *Proc. Natl. Acad. Sci.* 118 (28), e2024448118.
- Drineas, P., Frieze, A., Kannan, R., Vempala, S., Vinay, V., 2004. Clustering large graphs via the singular value decomposition. *Mach. Learn.* 56, 9–33.

- Glasser, M.F., Sotiropoulos, S.N., Wilson, J.A., Coalson, T.S., Fischl, B., Andersson, J.L., Xu, J., Jbabdi, S., Webster, M., Polimeni, J.R., et al., 2013. The minimal preprocessing pipelines for the Human Connectome Project. *NeuroImage* 80, 105–124.
- Golub, G.H., Van Loan, C.F., 2013. *Matrix computations*. JHU press.
- Gould, S., Lewontin, R., 1979. The spandrels of San Marco and the Panglossian paradigm: a critique of the adaptationist programme. *Proc. R. Soc. Lond. B Biol. Sci.* 205 (1161), 581–598.
- Guell, X., Schmahmann, J.D., Gabrieli, J.D., Ghosh, S.S., 2018. Functional gradients of the cerebellum. *eLife* 7, e36652.
- He, H., Liu, T.T., 2012. A geometric view of global signal confounds in resting-state functional MRI. *NeuroImage* 59, 2339–2348.
- Hong, S.-J., Xu, T., Nikolaidis, A., Smallwood, J., Margulies, D.S., Bernhardt, B., Vogelstein, J., Milham, M.P., 2020. Toward a connectivity gradient-based framework for reproducible biomarker discovery. *NeuroImage* 223, 117322.
- Huntenburg, J.M., Bazin, P.-L., Margulies, D.S., 2018. Large-scale gradients in human cortical organization. *Trends Cogn. Sci.* 22 (1), 21–31.
- Jenkinson, M., Beckmann, C.F., Behrens, T.E., Woolrich, M.W., Smith, S.M., 2012. FSL. *NeuroImage* 62 (2), 782–790.
- Krakauer, J.W., Ghazanfar, A.A., Gomez-Marín, A., MacIver, M.A., Poeppel, D., 2017. Neuroscience needs behavior: correcting a reductionist bias. *Neuron* 93 (3), 480–490.
- Langs, G., Golland, P., Ghosh, S.S., 2015. Predicting activation across individuals with resting-state functional connectivity based multi-atlas label fusion. In: *Int. Conf. Med. Image Comput. Comput.-Assist. Interv.* Springer, pp. 313–320.
- Laub, A.J., 2005. *Matrix Analysis for Scientists and Engineers*, Vol. 91. SIAM.
- Liegeois, R., Laumann, T.O., Snyder, A.Z., Zhou, J., Yeo, B.T., 2017. Interpreting temporal fluctuations in resting-state functional connectivity MRI. *NeuroImage* 163, 437–455.
- Liégeois, R., Yeo, B.T., Van De Ville, D., 2021. Interpreting null models of resting-state functional MRI dynamics: not throwing the model out with the hypothesis. *NeuroImage* 243, 118518.
- Liu, T.T., Nalci, A., Falahpour, M., 2017. The global signal in fMRI: Nuisance or information? *NeuroImage* 150, 213–229.
- Margulies, D.S., Ghosh, S.S., Goulas, A., Falkiewicz, M., Huntenburg, J.M., Langs, G., Bezgin, G., Eickhoff, S.B., Castellanos, F.X., Petrides, M., et al., 2016. Situating the default-mode network along a principal gradient of macroscale cortical organization. *Proc. Natl. Acad. Sci.* 113 (44), 12574–12579.
- Markello, R.D., Misic, B., 2021. Comparing spatial null models for brain maps. *NeuroImage* 236, 118052.
- Murphy, K., Fox, M.D., 2017. Towards a consensus regarding global signal regression for resting state functional connectivity MRI. *NeuroImage* 154, 169–173.
- Papoulis, A., 1965. *Probability, Random Variables, and Stochastic Processes*. McGraw-Hill.
- Power, J.D., Plitt, M., Laumann, T.O., Martin, A., 2017. Sources and implications of whole-brain fMRI signals in humans. *NeuroImage* 146, 609–625.
- Prichard, D., Theiler, J., 1994. Generating surrogate data for time series with several simultaneously measured variables. *Phys. Rev. Lett.* 73 (7), 951.
- Raut, R.V., Snyder, A.Z., Mitra, A., Yellin, D., Fujii, N., Malach, R., Raichle, M.E., 2021. Global waves synchronize the brain's functional systems with fluctuating arousal. *Sci. Adv.* 7 (30), eabf2709.
- Robinson, E.C., Jbabdi, S., Glasser, M.F., Andersson, J., Burgess, G.C., Harms, M.P., Smith, S.M., Van Essen, D.C., Jenkinson, M., 2014. MSM: a new flexible framework for multimodal surface matching. *NeuroImage* 100, 414–426.
- Rubinov, M., 2016. Constraints and spandrels of interareal connectomes. *Nature Commun.* 7 (1), 1–11.
- Cisek, P., Hayden, B. Y., 2022. Neuroscience needs evolution. *Philos. Trans. R. Soc. Lond. B Biol. Sci.* 377 (1844), 20200518.
- Rubinov, M. *Circular and unified analysis in network neuroscience*. OSF Preprints. <https://doi.org/10.31219/osf.io/mdqak>.
- Salimi-Khorshidi, G., Douaud, G., Beckmann, C.F., Glasser, M.F., Griffanti, L., Smith, S.M., 2014. Automatic denoising of functional MRI data: combining independent component analysis and hierarchical fusion of classifiers. *NeuroImage* 90, 449–468.
- Schaefer, A., Kong, R., Gordon, E.M., Laumann, T.O., Zuo, X.-N., Holmes, A.J., Eickhoff, S.B., Yeo, B.T., 2018. Local-global parcellation of the human cerebral cortex from intrinsic functional connectivity MRI. *Cereb. Cortex* 28 (9), 3095–3114.
- Shinn, M., Hu, A., Turner, L., Noble, S., Preller, K.H., Ji, J.L., Moujaes, F., Achard, S., Scheinost, D., Constable, R.T., Krystal, J.H., Vollenweider, F.X., Lee, D., Anticevic, A., Bullmore, E.T., Murray, J.D., 2023. Functional brain networks reflect spatial and temporal autocorrelation. *Nature Neurosci* doi:10.1038/s41593-023-01299-3.
- Siddiqi, S.H., Kording, K.P., Parvizi, J., Fox, M.D., 2022. Causal mapping of human brain function. *Nature Rev. Neurosci.* 23 (6), 361–375.
- Smith, R.L., 1984. Efficient Monte Carlo procedures for generating points uniformly distributed over bounded regions. *Oper. Res.* 32 (6), 1296–1308.
- Smith, S.M., Fox, P.T., Miller, K.L., Glahn, D.C., Fox, P.M., Mackay, C.E., Filippini, N., Watkins, K.E., Toro, R., Laird, A.R., et al., 2009. Correspondence of the brain's functional architecture during activation and rest. *Proc. Natl. Acad. Sci.* 106 (31), 13040–13045.
- Tarantola, A., 2005. 1. The General Discrete Inverse Problem. SIAM, pp. 1–40.
- Trefethen, L.N., Bau, D., 2022. *Numerical Linear Algebra*, Vol. 181. SIAM.
- Van den Meersche, K., Soetaert, K., Van Oevelen, D., 2009. xsample: an R function for sampling linear inverse problems. *J. Stat. Soft.* 30, 1–15.
- Van Essen, D.C., Smith, S.M., Barch, D.M., Behrens, T.E., Yacoub, E., Ugurbil, K., Consortium, W.-M.H., et al., 2013. The WU-Minn Human Connectome Project: an overview. *NeuroImage* 80, 62–79.
- Van Essen, D.C., Ugurbil, K., Auerbach, E., Barch, D., Behrens, T.E., Bucholz, R., Chang, A., Chen, L., Corbetta, M., Curtiss, S.W., et al., 2012. The Human Connectome Project: a data acquisition perspective. *NeuroImage* 62 (4), 2222–2231.
- Vos de Wael, R., Benkarim, O., Paquola, C., Larivière, S., Royer, J., Tavakol, S., Xu, T., Hong, S.-J., Langs, G., Valk, S., et al., 2020. Brainspace: a toolbox for the analysis of macroscale gradients in neuroimaging and connectomics datasets. *Commun. Biol.* 3 (1), 1–10.
- Vos de Wael, R., Royer, J., Tavakol, S., Wang, Y., Paquola, C., Benkarim, O., Eichert, N., Larivière, S., Xu, T., Misic, B., et al., 2021. Structural connectivity gradients of the temporal lobe serve as multiscale axes of brain organization and cortical evolution. *Cereb. Cortex* 31 (11), 5151–5164.
- Yeo, B.T., Krienen, F.M., Sepulcre, J., Sabuncu, M.R., Lashkari, D., Hollinshead, M., Roffman, J.L., Smoller, J.W., Zöllei, L., Polimeni, J.R., et al., 2011. The organization of the human cerebral cortex estimated by intrinsic functional connectivity. *J. Neurophysiol.* 106 (3), 1125–1165.
- Zhang, D., Raichle, M.E., 2010. Disease and the brain's dark energy. *Nature Rev. Neurol.* 6 (1), 15–28.
- Zivot, E., Wang, J., 2006. Vector autoregressive models for multivariate time series. *Model. Financ. Time Ser. S-PLUS®* 385–429.