

Unsupervised Learning of Functional Network Dynamics in Resting State fMRI

Harini Eavani¹, Theodore D. Satterthwaite², Raquel E. Gur²,
Ruben C. Gur², and Christos Davatzikos¹

¹ Section of Biomedical Image Analysis, Department of Radiology,
University of Pennsylvania

² Brain Behavior Laboratory, Department of Psychiatry, University of Pennsylvania

Abstract. Research in recent years has provided some evidence of temporal non-stationarity of functional connectivity in resting state fMRI. In this paper, we present a novel methodology that can decode connectivity dynamics into a temporal sequence of hidden network “states” for each subject, using a Hidden Markov Modeling (HMM) framework. Each state is characterized by a unique covariance matrix or whole-brain network. Our model generates these covariance matrices from a common but unknown set of sparse basis networks, which capture the range of functional activity co-variations of regions of interest (ROIs). Distinct hidden states arise due to a variation in the strengths of these basis networks. Thus, our generative model combines a HMM framework with sparse basis learning of positive definite matrices. Results on simulated fMRI data show that our method can effectively recover underlying basis networks as well as hidden states. We apply this method on a normative dataset of resting state fMRI scans. Results indicate that the functional activity of a subject at any point during the scan is composed of combinations of overlapping task-positive/negative pairs of networks as revealed by our basis. Distinct hidden temporal states are produced due to a different set of basis networks dominating the covariance pattern in each state.

Keywords: resting state fMRI, functional connectivity, temporal network dynamics.

1 Introduction

Resting state fMRI[1] has emerged as a powerful tool in understanding the effect of mental illnesses on brain function[2]. Functional connectivity or strength of synchronous activity between regions of interest is an important measure that could reveal disease-related changes in brain physiology. Correlation values are widely used as a measure of connectivity but estimation is restricted to a single value obtained from the entire duration of the scan. This could lead to loss of potentially valuable information, since recent exploratory work seems to indicate significant temporal variation in the correlation between regions [3–5]. Using a sliding window framework, the authors reported the presence of repetitive patterns of whole-brain network activity. However sampling the correlation values

may not be very reliable due to high estimation error from the smaller windows. Hence, in this paper, we propose modeling the fMRI time-series directly, avoiding explicit sampling of the correlation values.

In this paper, we present a novel method that uses a HMM framework to discretize the temporal variation into a temporal sequence of hidden states. These hidden states could be cognitive processes like introspection, memory consolidation or arising due to unknown external or internal triggers or stimuli[6]. Within each state, the fMRI time-series data is modeled as observations sampled from a multi-variate Gaussian. Each state is characterized by a mean vector value and a unique covariance matrix or whole-brain network. We assume a relatively small number of underlying regions or processes drive the variation in the fMRI signals, introducing subtle changes in the covariance matrices, from which these hidden states can be identified. We model these underlying co-varying regions as a set of sparse rank-one basis matrices, such that non-negative combinations of these basis matrices act as priors for each of the HMM covariance matrices. These basis matrices are unknown and are learned from the data. Thus, our method is a joint framework that solves for the basis vectors as well as the hidden states simultaneously.

The rest of the paper is organized as follows. In section 2 we discuss our generative model in detail. In section 3 we describe the performance of our method on a simulated dataset. We apply our algorithm to resting state fMRI data, and the results are described in section 4. We wrap up with our conclusions and future work in section 5.

2 Approach

2.1 Hidden Markov Model

We begin by describing the first-order HMM framework. Let $\mathbf{Y}^s = [\mathbf{y}_1^s, \mathbf{y}_2^s, \dots, \mathbf{y}_T^s]$, $\mathbf{y}_t^s \in \mathbf{R}^p$ be the fMRI time-series of a subject s , where p is the number of ROIs and T is the total number of time-points per subject. The superscript s denotes subject index, and the subscript t denotes time-index. Since resting state fMRI is acquired without any control over the subject's stimulus or environment, it is reasonable to assume that the subject wanders in and out of various cognitive states during the duration of the scan. Hence, we will assume that every time-point belongs to one of a finite number N of states. Each state is associated with an occurrence probability δ_i , and every pair of states i, j is associated with a transition probability Π_{ij} of moving from state i to state j . We are interested in describing these states quantitatively, as well finding the optimal sequence of states for each subject. A schematic diagram of an HMM with $N = 3$ states is shown in Fig 1.

We model the “emission” probabilities by a Gaussian distribution. Let $S_t^s \in \{1, 2, \dots, N\}$ be a random variable denoting the state assignment for subject s at time t . Then, given the state assignment $S_t^s = i$, we let $\mathbf{y}_t^s \sim \mathcal{N}(\mu_i, \Sigma_i)$, i.e.,

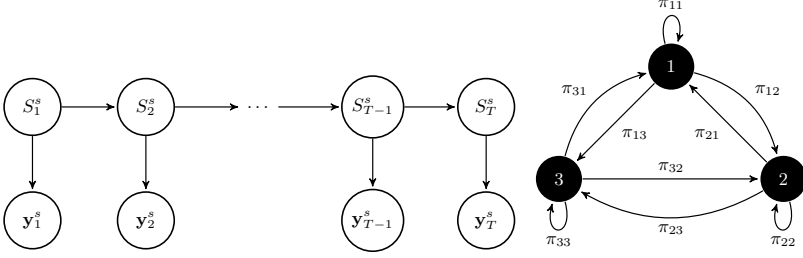


Fig. 1. Schematic of an HMM with $N = 3$. Left shows the temporal sequence $S_1^s, S_2^s, \dots, S_T^s$ associated with a subject s . Each S_t^s could be one of N states (right).

$$p(\mathbf{y}_t^s | S_t^s = i, \mu_i, \Sigma_i) = \frac{1}{(2\pi)^{p/2} |\Sigma_i|^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{y}_t^s - \mu_i)^T (\Sigma_i)^{-1} (\mathbf{y}_t^s - \mu_i) \right\} \quad (1)$$

where $|\cdot|$ denotes determinant. The log-likelihood of the data as a function of the variables is

$$l_1 \left(\{\mu_i\}_1^N, \{\Sigma_i\}_1^N \right) = \sum_{i=1}^N N_i \left\{ \log \det \Sigma_i^{-1} - \text{tr} (\mathbf{S}_i \Sigma_i^{-1}) \right\} \quad (2)$$

where N_i is the number of time-points that exist in state i , and $\mathbf{S}_i$ is the sample covariance matrix computed using all the time-points assigned to state i .

2.2 Sparse Dictionary Learning for Positive Definite Matrices

We hypothesize that a relatively small number of ROIs change their co-variance pattern from state to state. Similar to the basis learning formulation proposed in [7, 8], let $\mathbf{B} = [\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_K]$, $-1 \preceq \mathbf{b}_k \preceq 1$, $\mathbf{b}_k \in \mathbf{R}^p$ be a set of K basis vectors such that each vector $\mathbf{b}_k$ reflects the membership of the ROIs to the basis network k . If $|\mathbf{b}_k(i)| > 0$, ROI i belongs to the basis vector k , and if $\mathbf{b}_k(i) = 0$ it does not. If two ROIs in $\mathbf{b}_k$ have the same sign, then they are positively correlated and opposing sign reflects that they are anti-correlated. Therefore, the rank-one matrix $\mathbf{b}_k \mathbf{b}_k^T$ reflects the covariance behavior of basis k . In addition, we constrain these basis networks to be much smaller than the whole-brain network by restricting their l_1 -norm to a constant value λ .

We would like to approximate the matrices $\{\Sigma_i\}_{i=1}^N$ representing the HMM states by a non-negative combination of these basis networks. Thus, we want

$$\begin{aligned} \Sigma_i &\approx \sum_{k=1}^K c_i(k) \mathbf{b}_k \mathbf{b}_k^T = \mathbf{B} \text{diag}(\mathbf{c}_i) \mathbf{B}^T \triangleq \hat{\Sigma}_i \\ \|\mathbf{b}_k\|_1 &\leq \lambda, \quad -1 \preceq \mathbf{b}_k \preceq 1, \quad \mathbf{c}_i \geq 0 \end{aligned} \quad (3)$$

where $\text{diag}(\mathbf{c}_i)$ denotes a diagonal matrix with values $\mathbf{c}_i$ along the diagonal. In practice, another term $\alpha \mathbf{I}_p$ is added to the $\hat{\Sigma}_i$ matrix to make it positive definite. Here $\mathbf{I}_p$ is the identity matrix of size $p \times p$, and α is a small, fixed ($= 0.1$) value.

We quantify the approximation between Σ_i and $\hat{\Sigma}_i$ using the Kullback-Liebler divergence:

$$\text{KL}(\hat{\Sigma}_i \parallel \Sigma_i) = \text{tr}(\hat{\Sigma}_i \Sigma_i^{-1}) - \log \det(\hat{\Sigma}_i \Sigma_i^{-1}) \quad (4)$$

The KL divergence described above quantifies the amount of information *lost* when a Gaussian distribution with covariance matrix $\hat{\Sigma}_i$ is instead modeled by covariance matrix Σ_i . A low value of $\text{KL}(\hat{\Sigma}_i \parallel \Sigma_i)$ indicates a good approximation between the two matrices. Therefore, we are interested in minimizing the KL-divergence for all the pairs $(\hat{\Sigma}_i, \Sigma_i)$. This amounts to maximizing the function

$$l_2(\mathbf{B}, \mathbf{C}) = \sum_{i=1}^N N_i \left\{ \log \det(\hat{\Sigma}_i \Sigma_i^{-1}) - \text{tr}(\hat{\Sigma}_i \Sigma_i^{-1}) \right\} \quad (5)$$

where $\mathbf{C} = [\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_N]$ and N_i is defined as before in equation 2.

We are interested in finding HMM states that are distinct primarily due to differences in their covariance matrices Σ_i . Clustering data solely based on covariances is a challenging problem. However, if we assume that the matrices are generated from a common underlying basis $\mathbf{B}$ (as described above), we may be able to separate the clusters by forcing the coefficients $\mathbf{c}_i$ to be distinct, or equivalently, requiring that the inner product $\langle \mathbf{c}_i, \mathbf{c}_j \rangle$, $i \neq j$ be small. This constraint can be imposed by making the term $C^T C$ resemble the identity matrix. Thus, we would like to maximize

$$l_3(\mathbf{C}) = - \left(\sum_{i=1}^N N_i \right) \text{KL}(C^T C \parallel \mathbf{I}_N) = \left(\sum_{i=1}^N N_i \right) \left\{ \log \det(C^T C) - \text{tr}(C^T C) \right\} \quad (6)$$

where $\mathbf{I}_N$ is the identity matrix of size $N \times N$.

2.3 Joint Framework: HMM + Sparse Dictionary Learning

As mentioned earlier, a joint framework causes the HMM to converge to hidden states with distinct covariance matrices. Combining the three objectives in equations 2, 5 and 6 amounts to imposing a prior on the covariance matrices Σ_i with prior variables $\mathbf{B}$ and $\mathbf{C}$. We would like to maximize the joint log-likelihood

$$l \left(\{\mu_i\}_1^N, \{\Sigma_i\}_1^N, \mathbf{B}, \mathbf{C} \right) = l_1(\{\mu_i\}, \{\Sigma_i\}) + v_1 l_2(\mathbf{B}, \mathbf{C}) + v_2 l_3(\mathbf{C}) \quad (7)$$

where v_1 and v_2 are user-defined scalar parameters that control the amount of coupling between the HMM and priors. The constraints on the variables are given by

$$\Sigma_i \succ 0, \quad \|\mathbf{b}_k\|_1 \leq \lambda, \quad -1 \preceq \mathbf{b}_k \preceq 1, \quad \mathbf{c}_i \succeq 0, \quad i = 1, 2, \dots, N, \quad k = 1, 2, \dots, K \quad (8)$$

Algorithm 1. Optimization strategy for maximizing log-likelihood in Eqn. 7**Input:** Data $\mathbf{Y}$, Parameters v_1, v_2, λ, K, N **Initialize:** $\Pi, \delta, \mathbf{B}, \mathbf{C}, \{\mu_i\}_1^N, \{\Sigma_i\}_1^N$ **while** Log-likelihood in Eqn. 7 is increasing **do****E-step**Compute weights $q_1(S_t^s = i)$ and $q_2(S_t^s = i, S_{t-1}^s = j)$ using Baum-Welch[9]**M-step:**Update $\{\mathbf{S}_i\}_1^N, \{\mu_i\}_1^N$ and $\{\Sigma_i\}_1^N$ using Equations 9 and 10Solve for $\mathbf{B}$ and $\mathbf{C}$ using SPG solver, similar to [10]Update Π and δ using standard HMM update formulae[11]**end while**

Compute optimal state sequences using Viterbi algorithm [12]

2.4 Joint Optimization Strategy

We use Expectation-Maximization[11] to obtain a local maximum. In the Expectation step the posterior marginals $q_1(S_t^s = i) = p(S_t^s = i|\mathbf{Y})$ and $q_2(S_t^s = i, S_{t-1}^s = j) = p(S_t^s = i, S_{t-1}^s = j|\mathbf{Y})$ are efficiently computed using the Baum-Welch[9] algorithm. These values are used as weights for the log-likelihood function. In the Maximization step the weighted log-likelihood function maximized with respect to each of the variables.

The joint log-likelihood is jointly non-concave, but individually concave w.r.t the variables $\{\Sigma_i^{-1}\}_1^N, \mathbf{B}$ and $\mathbf{C}$. Hence we will adopt a block optimization strategy that repeatedly solves for one variable (e.g. $\mathbf{B}$) while holding the others fixed (e.g. $\mathbf{C}$ and $\{\Sigma_i^{-1}\}_1^N$) until a local optimum is reached.

The optimal values for $\{\mu_i\}_1^N$ and $\{\Sigma_i\}_1^N$ have closed form expressions given by

$$\mu_i = \frac{\sum_{s,t} q_1(S_t^s = i) \mathbf{y}_t^s}{N_i}, \quad \Sigma_i = \frac{1}{v_1 + 1} \mathbf{S}_i + \frac{v_1}{v_1 + 1} \hat{\Sigma}_i \quad (9)$$

where

$$N_i = \sum_{s,t} q_1(S_t^s = i), \quad \mathbf{S}_i = \frac{1}{N_i} \sum_{s,t} q_1(S_t^s = i) (\mathbf{y}_t^s - \mu_i)(\mathbf{y}_t^s - \mu_i)^T, \quad i = 1, \dots, N \quad (10)$$

The optimization w.r.t the variables $\mathbf{B}$ and $\mathbf{C}$ is a constrained maximization problem without closed form solutions. We use the spectral projected gradient (SPG) solver with an efficient projection method, similar to the algorithm proposed in [10] to solve for $\mathbf{B}$ and $\mathbf{C}$ separately.

Matrices $\mathbf{B}$ and $\mathbf{C}$ are initialized randomly. The state transition matrix Π and occurrence probabilities δ are initialized as $\Pi_{ii} = 0.5, \Pi_{ij} = 0.5 * (N - 1), i \neq j, \delta_i = 1/N$ for $i, j = 1, 2, \dots, N$. Mean vectors μ_i and covariance matrices Σ_i are initialized by using random selection of data points. After the local optima for the unknowns are found using EM, the optimal state sequence for each subject

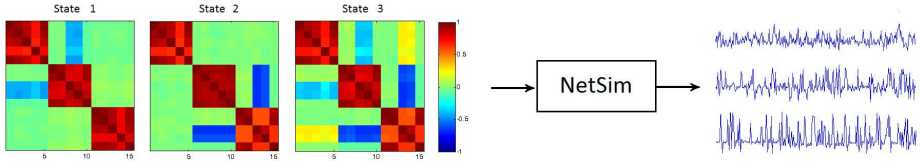


Fig. 2. Simulated data: Ground truth correlation matrices (left) and 50 randomly generated temporal state sequences was input to NetSim [13]. The resulting time-series form the input to our method.

can be found using the Viterbi algorithm [12]. The overall strategy is summarized in Algorithm 1.

2.5 Choice of Free Parameters

Parameters v_1 and v_2 control the effect of the prior variables $\mathbf{B}$ and $\mathbf{C}$ on the model. $v_1 = v_2 = 0$ reduces the model to a standard HMM. Observe from Equation 9 that the optimal value of Σ_i at every M-step is a weighted average of the sample covariance matrix $\mathbf{S}_i$ and the approximation $\hat{\Sigma}_i$. The weight is controlled by the free parameter v_1 . The v_2 parameter controls the orthogonality constraint on the coefficient vectors $\mathbf{c}_i$. For all our experiments in this paper, we set $v_1 = v_2 = 1$. Parameter λ controls the amount of sparsity in the basis vectors and can be set based on known clinical information, for e.g., the average size of a sub-network, like the default mode network(DMN) or the fronto-parietal network.

The number of basis vectors K and the number of HMM states N can be chosen based on how the estimated values for Σ_i , $\mathbf{B}$ and $\mathbf{C}$ generalize. To assess generalizability, we will resort to Monte-Carlo split-sample cross-validation. All the other parameters being held fixed, for every value of K and N , the dataset is split into two halves. The model is trained on one half, and using the parameters Σ_i , $\mathbf{B}$, $\mathbf{C}$ computed from this half, the weights $q_t^s(i)$ are computed for the second half. This procedure is repeated multiple times and the average cross-validated log-likelihood is computed. The optimal choice of the parameters is considered to be the value at which the average log-likelihood does not significantly change. In this paper we will only examine the case when $K = N$.

3 Validation Using Simulated Data

3.1 Data

We used NetSim [13] to generate time-series data in order to evaluate our method. This software takes as input the underlying network configuration(s) and temporal state sequences. It returns realistic BOLD time series while incorporating neural lag (50 ms), variability in Hemodynamic Response Function (0.5 s) and thermal noise(1% of signal power).

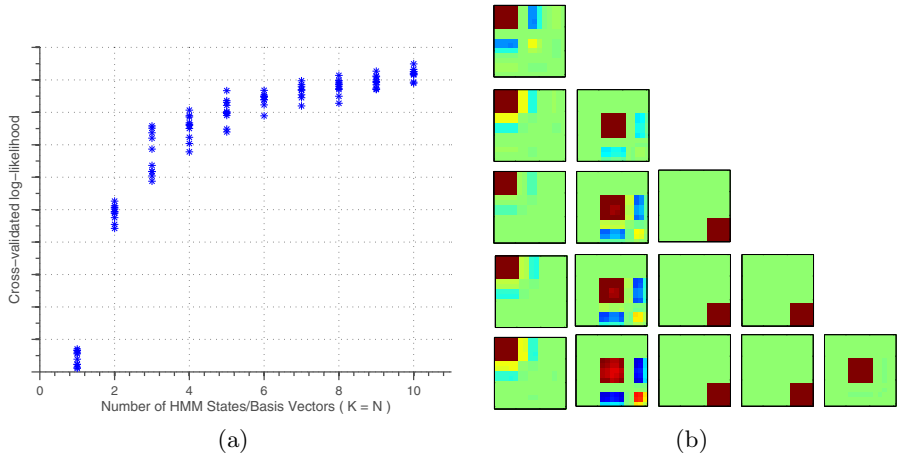


Fig. 3. (a) Monte-Carlo cross-validation log-likelihood for simulated data (b) Basis vectors for $K \in \{1, 2, 3, 4, 5\}$. Each row corresponds to a fixed value for K .

At any point in time and in any subject, the data is generated from one of the three network configurations, characterized by the covariance matrices shown in Figure 2. Our simulation consists of 15 nodes arranged in three subnetworks, which are positively correlated within each other (color scale red shown in Figure 2). The between-network connections vary with time - they are either zero (green), or negative (blue). Fifty temporal state sequences are used as input - the mean duration for each state was 40s (~ 13 TRs). The data was generated by applying Gaussian noise with mean zero as the stimulus at nodes 1,6 and 11. This ensures that the resulting data has mean value close to zero for all the nodes. The basis networks and the temporal state sequences was input to NetSim. This resulted in BOLD time series data for 50 “subjects”, with TR=3 s and 120 time-points each.

3.2 Results on Simulated Data

Figure 3a shows the results of the Monte-Carlo cross-validation procedure on the simulated data as N is varied. The average cross-validated log-likelihood with increases with increasing $K = N$, showing that the HMM clusters generalize well. The other parameters are fixed at $v_1 = v_2 = 1$ and $\lambda = 0.2$. The gain in generalizability is reduced after $K \sim 3$ or 4.

The rank-one basis matrices $\mathbf{b}_k \mathbf{b}_k^T$, $k = 1, 2, \dots, K$ computed for $K \in \{1, 2, 3, 4, 5\}$ are shown in Figure 3b. Each row corresponds to a fixed value for K . The values of the other parameters were fixed at $v_1 = v_2 = 1$ and $\lambda = 0.2$. It is evident that our algorithm effectively recovers the network basis. Our basis clearly identifies the clustering of the nodes into sub-networks and the anti-correlated relationship between them. Also, observe that each time K is increased by one, the method incrementally adds to the previous basis.

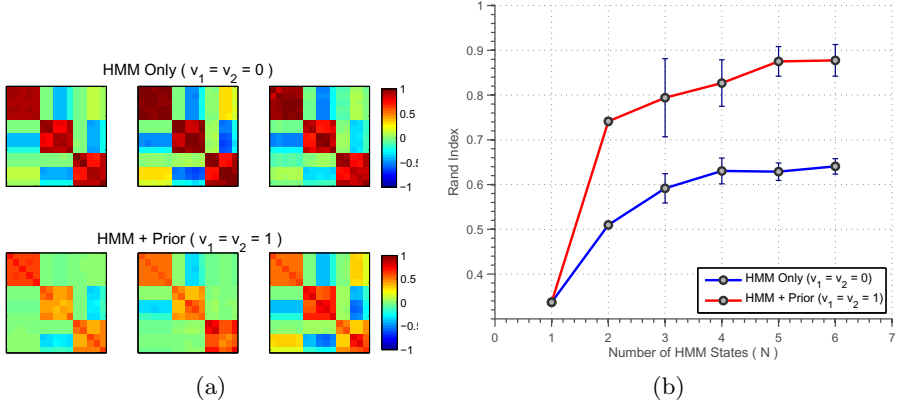


Fig. 4. (a) Covariance matrices Σ_i for HMM only (top) and our method (bottom) (b) Cross-validated Rand Index vs. N for both cases

The comparison between the performance of our method with the HMM alone is shown in Figure 4. Figure 4a shows the estimated covariance matrices $\{\Sigma_i\}_{i=1}^N$ for $N = K = 3$, when only the HMM is used (top) and with the priors (bottom). Clearly, our method accurately estimates the underlying covariance matrices for all three states. Due to the lack of a significant difference in the mean vectors, the HMM alone performs poorly, with little or no difference between the three states. The optimal state-paths output by the Viterbi algorithm are also compared with the ground truth for both cases using the Rand Index (Figure 4b). Our method is able to achieve close to 90% clustering accuracy, while the HMM alone fails.

4 Application to Resting State fMRI Data

4.1 Data

BOLD fMRI was acquired with a Siemens 3 Tesla system using a whole-brain, gradient-echo echo planar sequence with TR/TE = 3000/32 ms, voxel resolution = 3x3x3 mm and number of time-points = 120. We used data from 420 normal participants, with age range 15.9 ± 3 years.

Pre-processing. Functional images were motion corrected, spatially smoothed (6mm FWHM) and temporally altered to retain frequencies 0.01-0.1 Hz. Several sources of confounding variance, including six motion parameters and mean whole-brain, WM, CSF time-courses were regressed out. The residual time course for each subject was transformed to standard MNI anatomical space. We used 160 regions of interest (ROIs) described by Dosenbach et al. [14], which were derived from a meta-analysis of a large sample of task-based fMRI studies. Each ROI was a non-overlapping 10mm diameter sphere, and was categorized by Dosenbach et al. as belonging to one of six networks, including: default-mode, cingulo-opercular, fronto-parietal, sensorimotor, occipital or cerebellar. The mean time-series of each ROI is

extracted from the registered fMRI image. The time-series is demeaned and scaled to have an average variance value of unity.

4.2 Results on fMRI Data

Fig. 5. Monte-Carlo cross-validation log-likelihood for fMRI data. The log-likelihood is greatest for $K = N \in \{2, 3, 4, 5\}$. The generalizability is variable for $N \in \{6, 7\}$ and begins to fall after $N = 7$, showing that our model begins to over-fit the data after this value. Thus, from the given data, we are able to obtain $N = 6$ distinct HMM states.

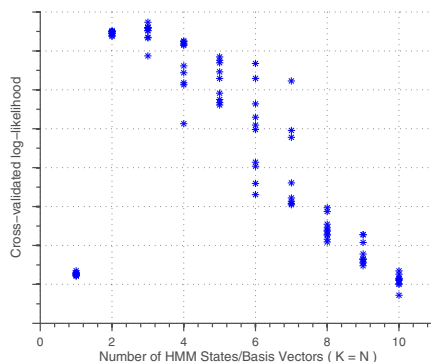


Figure 6 shows the six rank-one basis matrices obtained from our method. The ROIs are sorted according to their Dosenbach[14] network labels. It is easy to observe that our method is fairly accurate in identifying the general clustering

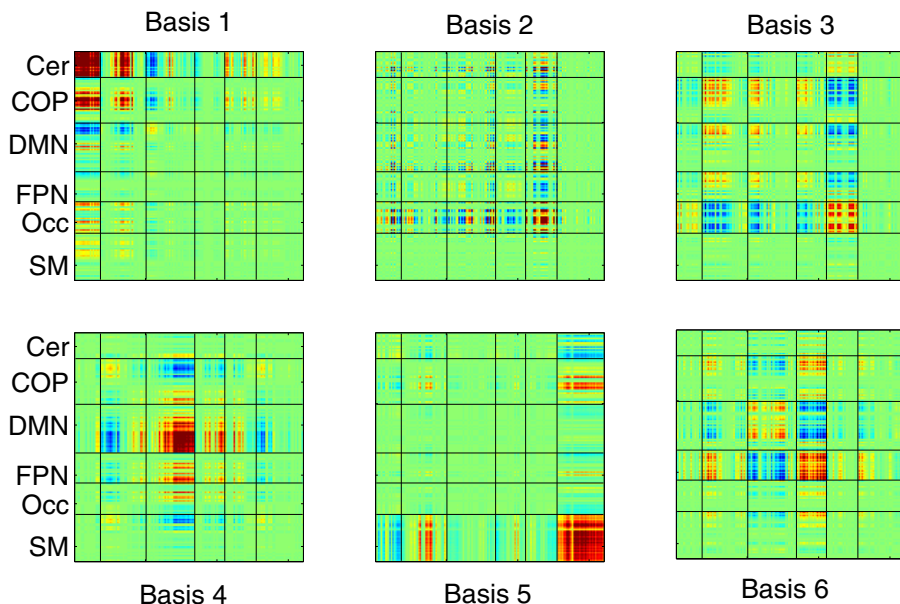


Fig. 6. Rank-one basis matrices for fMRI data. The Dosenbach [14] labels are given on the left. Abbreviations Cer: Cerebellum, COP: Cingulo-opercular network, DMN: Default-mode network, FPN: Fronto-Parietal network, Occ: Occipital network, SM: Sensori-motor network.

of the ROIs, since most ROIs belonging to a basis network are assigned the same Dosenbach labels. For example, all the ROIs belonging to the cerebellum are clustered in basis 1. Same is the case with occipital cortex (basis 3), default mode (basis 4), sensori-motor (basis 5) and fronto-parietal (basis 6).

We looked at two sub-networks in particular - the fronto-parietal network (“dorsal attentional network”) and the cingulo-opercular network (“ventral attentional network”). Both are task-positive networks that activate when the subject is in an “extrospective” state, and it is well-established that activation of either of these networks causes the default mode network (DMN) to deactivate [1]. This behavior is captured in basis 6, which shows the anti-correlated nature of the fronto-parietal network and default mode. The anti-correlation between the cingulo-opercular network (COP) and the DMN is captured across multiple basis networks (1, 4 and 6).

We also note that the most amount of overlap between the basis networks occurs at the COP, with different aspects of it positively correlating with the cerebellum, sensori-motor and fronto-parietal networks(in basis 1, 5 and 6 respectively) and negatively correlating with the default mode and occipital

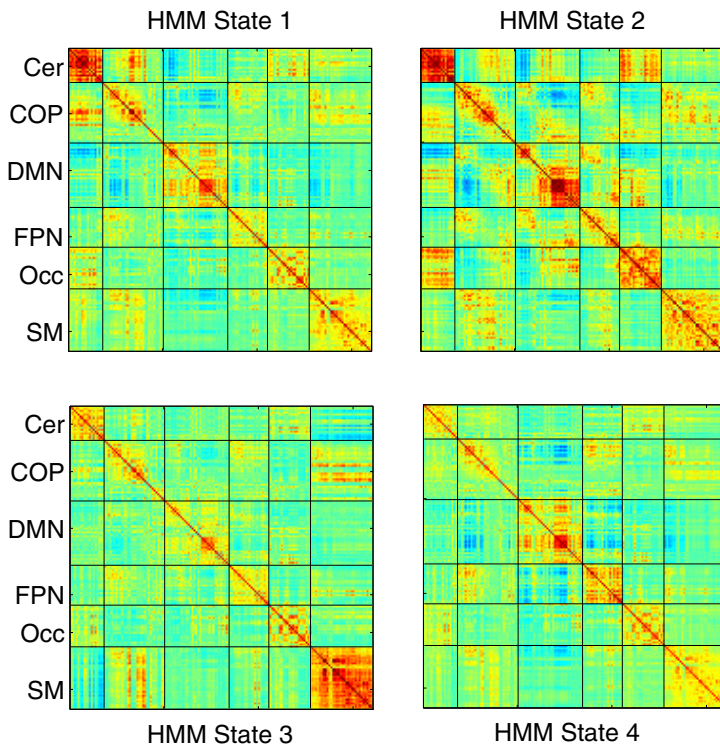


Fig. 7. Four HMM States obtained from resting state fMRI data. Abbreviations Cer: Cerebellum, COP: Cingulo-opercular network, DMN: Default-mode network, FPN: Fronto-Parietal network, Occ: Occipital network, SM: Sensori-motor network.

networks (in basis 1 and 3). The fact that these correlation/anti-correlation relationships are split amongst the basis networks suggests that they have different temporal behavior.

The covariance matrices Σ_i are shown in Figure 7. Due to the page limit only the four of the six covariance matrices are shown. It is clear that the states are separated based on the dominant basis networks. The cingulo-opercular network is most active in states 1 and 2, positively correlating with the cerebellum and negatively correlating with the DMN. State 2 shows greater activity in the occipital network. The sensori-motor network is active in state 3. State 4 is dominated by the anti-correlated pair of the fronto-parietal network and the DMN. The average duration the subjects existed in each of these states was between 12 and 20 time-points.

5 Conclusion

To our knowledge, this is the first attempt at resolving functional connectivity in resting state fMRI data into discrete temporal states, each associated with a distinct connectivity pattern. Each subject is assigned a sequence of states, which can then be used for group comparisons.

Our basis learning formulation provides sparse and possibly overlapping components without having to use strong constraints like orthogonality or independence of the basis. Further more, our basis decomposition only allows non-negative combinations of basis vectors, making the resulting basis more interpretable. These properties make our method better suited than spatial or temporal ICA [15] for decomposing brain activity into interpretable components.

Hidden Markov Models have been used in the context of fMRI, but primarily for task-based experiments [16][17]. In our method the emphasis is on finding hidden brain-states when an external stimulus is not provided to the subject. The HMM is strongly driven by the differences in the covariance matrices of these hidden brain states.

As a part of our future work, we hope to analyze the effect of model selection on our method in greater detail. A thorough examination of the functional interpretability of the basis vectors and HMM states is needed. As an additional validation step, this method can be applied to task-fMRI to recover the stimulus sequence.

References

1. Raichle, M.E., MacLeod, A.M., Snyder, A.Z., Powers, W.J., Gusnard, D.A., Shulman, G.L.: A default mode of brain function. *Proceedings of the National Academy of Sciences* 98(2), 676–682 (2001)
2. Broyd, S.J., Demanuele, C., Debener, S., Helps, S.K., James, C.J., Sonuga-Barke, E.J.S., et al.: Default-mode brain dysfunction in mental disorders: a systematic review. *Neuroscience and Biobehavioral Reviews* 33(3), 279 (2009)
3. Chang, C., Glover, G.H.: Time-frequency dynamics of resting-state brain connectivity measured with fMRI. *Neuroimage* 50(1), 81–98 (2010)

4. Majeed, W., Magnuson, M., Hasenkamp, W., Schwarb, H., Schumacher, E.H., Barsalou, L., Keilholz, S.D.: Spatiotemporal dynamics of low frequency bold fluctuations in rats and humans. *Neuroimage* 54(2), 1140–1150 (2011)
5. Hutchison, R.M., Womelsdorf, T., Gati, J.S., Everling, S., Menon, R.S.: Resting-state networks show dynamic functional connectivity in awake humans and anesthetized macaques. *Human Brain Mapping* (2012)
6. Buckner, R.L., Vincent, J.L., et al.: Unrest at rest: default activity and spontaneous network correlations. *Neuroimage* 37(4), 1091–1096 (2007)
7. Sra, S., Cherian, A.: Generalized dictionary learning for symmetric positive definite matrices with application to nearest neighbor retrieval. In: Gunopulos, D., Hofmann, T., Malerba, D., Vazirgiannis, M. (eds.) *ECML PKDD 2011, Part III. LNCS (LNAI)*, vol. 6913, pp. 318–332. Springer, Heidelberg (2011)
8. Sivalingam, R., Boley, D., Morellas, V., Papanikolopoulos, N.: Positive definite dictionary learning for region covariances. In: 2011 IEEE International Conference on Computer Vision (ICCV), pp. 1013–1019. IEEE (2011)
9. Baum, L.E., Petrie, T., Soules, G., Weiss, N.: A maximization technique occurring in the statistical analysis of probabilistic functions of markov chains. *The Annals of Mathematical Statistics*, 164–171 (1970)
10. Batmanghelich, N.K., Taskar, B., Davatzikos, C.: Generative-discriminative basis learning for medical imaging. *IEEE Transactions on Medical Imaging* 31(1), 51–69 (2012)
11. Bishop, C.M., et al.: *Pattern recognition and machine learning*, vol. 4. Springer, New York (2006)
12. Viterbi, A.: Error bounds for convolutional codes and an asymptotically optimum decoding algorithm. *IEEE Transactions on Information Theory* 13(2), 260–269 (1967)
13. Smith, S.M., Miller, K.L., Salimi-Khorshidi, G., Webster, M., Beckmann, C.F., Nichols, T.E., Ramsey, J.D., Woolrich, M.W.: Network modelling methods for fmri. *Neuroimage* 54(2), 875–891 (2011)
14. Dosenbach, N.U.F., Fair, D.A., Miezin, F.M., Cohen, A.L., et al.: Distinct brain networks for adaptive and stable task control in humans. *Proceedings of the National Academy of Sciences* 104(26), 11073 (2007)
15. Smith, S.M., Miller, K.L., Moeller, S., Xu, J., Auerbach, E.J., Woolrich, M.W., Beckmann, C.F., Jenkinson, M., Andersson, J., Glasser, M.F., et al.: Temporally-independent functional modes of spontaneous brain activity. *Proceedings of the National Academy of Sciences* 109(8), 3131–3136 (2012)
16. Faisan, S., Thoraval, L., Armspach, J.P., Heitz, F.: Hidden markov multiple event sequence models: A paradigm for the spatio-temporal analysis of fmri data. *Medical Image Analysis* 11(1), 1 (2007)
17. Janoos, F., Machiraju, R., Singh, S., Morocz, I.Á.: Spatio-temporal models of mental processes from fmri. *Neuroimage* 57(2), 362–377 (2011)