

JOINT BAYESIAN ESTIMATION OF VOXEL ACTIVATION AND INTER-REGIONAL CONNECTIVITY IN FMRI EXPERIMENTS

DANIEL SPENCER[✉], RAJARSHI GUHANIYOGI AND RAQUEL PRADO

UNIVERSITY OF CALIFORNIA

Brain activation and connectivity analyses in task-based functional magnetic resonance imaging (fMRI) experiments with multiple subjects are currently at the forefront of data-driven neuroscience. In such experiments, interest often lies in understanding activation of brain voxels due to external stimuli and strong association or connectivity between the measurements on a set of pre-specified groups of brain voxels, also known as regions of interest (ROI). This article proposes a joint Bayesian additive mixed modeling framework that simultaneously assesses brain activation and connectivity patterns from multiple subjects. In particular, fMRI measurements from each individual obtained in the form of a multi-dimensional array/tensor at each time are regressed on functions of the stimuli. We impose a low-rank parallel factorization decomposition on the tensor regression coefficients corresponding to the stimuli to achieve parsimony. Multiway stick-breaking shrinkage priors are employed to infer activation patterns and associated uncertainties in each voxel. Further, the model introduces region-specific random effects which are jointly modeled with a Bayesian Gaussian graphical prior to account for the connectivity among pairs of ROIs. Empirical investigations under various simulation studies demonstrate the effectiveness of the method as a tool to simultaneously assess brain activation and connectivity. The method is then applied to a multi-subject fMRI dataset from a balloon-analog risk-taking experiment, showing the effectiveness of the model in providing interpretable joint inference on voxel-level activations and inter-regional connectivity associated with how the brain processes risk. The proposed method is also validated through simulation studies and comparisons to other methods used within the neuroscience community.

Key words: Bayesian inference, brain activation, brain connectivity, functional magnetic resonance imaging, graphical modeling, multiway stick-breaking prior, PARAFAC decomposition, tensor response.

1. Introduction

Recently, rapid advancements in different imaging modalities have generated massive neuroimaging data which are key in understanding how the human brain functions. For the present article, our motivation is mainly drawn from multi-subject functional MRI (fMRI) studies. Data from an fMRI scan are processed from the scanner as a four-dimensional tensor object, where the index of a datum within the tensor provides information about the location of the datum in space and time. In the context of an fMRI scan, the brain at a single point in time can be pictured as a three-dimensional tensor partitioned into small cubes, known as *voxels* (Lazar 2008). A relative measure of oxygen in the blood, referred to as the blood oxygen level-dependent (BOLD) response, is obtained from every voxel in each scan typically acquired around every two seconds. Such scans can be taken when a subject is in a resting state without any conditions imposed; when a subject is exposed to certain conditions/stimuli, such as noises or videos; or when a subject is actively involved in completing a task. As oxygen is required to perform functions in the brain, these readings are used to infer which parts of the brain are “active” in a given thought process. Expected activation patterns include local spatial dependence in the sense that voxels located next

The research of Rajarshi Guhaniyogi is partially supported by grants from the Office of Naval Research (ONR-BAA N000141812741) and the National Science Foundation (DMS-1854662). Raquel Prado was partially supported by NSF grant SES-1853210.

Correspondence should be made to Daniel Spencer, Department of Statistics, University of California, Santa Cruz, CA, USA. Email: daspence@ucsc.edu; URL: <https://users.soe.ucsc.edu/~daspence/>

to each other tend to be jointly activated, as well as non-local dependencies in which groups of voxels in distant regions of the brain are activated by a given thought process.

In addition to determining brain activation linked with a specific cognitive or sensorimotor function, neuroscientists are often interested in the way in which different spatially adjacent groups of voxels, referred to as regions of interest (ROIs) in the brain work together to process information. These types of relationships between ROIs are collectively referred to as *functional connectivity* (Hutchison et al. 2013). The major contribution of this article is the proposal of a Bayesian modeling framework that simultaneously detects voxel-level activation and connectivity between different ROIs with precise characterization of uncertainty for multi-subject fMRI data. These methods, applied to task-based fMRI experiments, lead to meaningful inferences about the functions of the human brain in such experimental settings.

Simultaneous analysis of multi-subject 3D fMRI scans is a challenging problem due to the sheer amount of data. Our modeling framework addresses this issue by using a mixed effects tensor response regression analysis in which low-rank tensor decompositions are combined with a multiway stick-breaking shrinkage prior to achieve parsimony in the estimation of voxel-level activation. Such framework provides a powerful and computationally feasible setting for inferring activation and connectivity in multi-subject task-related fMRI studies.

Before describing our proposed modeling approach, we present an overview of currently available methods that separately infer activation at the voxel level and connectivity at the region-specific level in Sects. 1.1 and 1.2, respectively. Section 1.3 then presents a review of models that jointly infer activation and connectivity. Due to the space constraint, we mainly focus on describing approaches that would be natural competitors to our proposed approach either because they share similarities in at least one of the modeling components (e.g., tensor-based approaches), or they use the same inferential paradigm (Bayesian approaches).

1.1. Activation Models

Several approaches have been proposed for the analysis of brain activation. Single-subject frameworks in particular have a rich background for modeling activation. The simplest of them, known as the general linear model (GLM), fits a regression model at each voxel with the observed voxel-specific BOLD response regressed on activation-related predictors and identifies if the response is significantly associated with the predictors in that voxel (Friston et al. 1995; Penny et al. 2011), after accounting for multiple testing corrections. Many implementations of the GLM include clusterwise corrections to account for spatial association using some variant of independent components analysis. However, work by Eklund et al. (2016) suggests that many of these methods inflate false-positive rates. The Benjamini–Hochberg correction (Benjamini and Hochberg 1995), which works by setting the false discovery rate, is a possible solution, but remains incomplete, as it does not take spatial information into account. Another idea, which addresses the sparse nature of fMRI activation (Olshausen and Field 2004), assigns a spike-and-slab prior on the regression coefficients (Brown et al. 1998; Yu et al. 2018) across all voxels. These priors take the form

$$\beta_v \sim \gamma_v \text{Normal}(0, v_1) + (1 - \gamma_v) \text{Normal}(0, v_0), \quad (1)$$

with β_v the activation parameter at voxel v . In this setting, v_0 is relatively small, v_1 is relatively large. γ_v is a zero-one random variable such that $\Pr(\gamma_v = 1) = \pi_v$, with π_v being the probability that β_v comes from the normal distribution with larger variance. While these methods tend to have robust sensitivity and specificity, they do not take spatial information into account. In fact, brain activation is a local process in the sense that spatially contiguous voxels generally tend

to be activated together while performing a task. Spike-and-slab prior distribution is unable to incorporate such structural information a priori.

Various approaches also account for spatial association in the neighboring voxels by inducing dependence among voxel-specific regression coefficients using Markov random fields (Zhang et al. 2014b; Kalus et al. 2014; Smith and Fahrmeir 2007; Lee et al. 2014). In these approaches, the activation coefficients are represented as having a multivariate normal prior with precision Σ^{-1} . Within Σ^{-1} the values on the diagonal are the number of neighbors that voxel v has. The off-diagonal elements σ_{vk} are set equal to 1 if voxel v and voxel k are considered neighbors, and 0 otherwise. Models with this type of Markov random fields components usually capture, at least partially, the underlying spatial structure in fMRI data but tend to be computationally expensive. More complex Markov random field structures may be proposed, but almost certainly at the cost of further decreased computational efficiency. In contrast, our proposed tensor mixed effect model improves inference by encouraging localized activation without explicitly employing spatial dependence.

Other sophisticated approaches include spatially varying coefficient (SVC) models which employ spatial basis functions to model activation-related coefficients (Flandin and Penny 2007; Zhu et al. 2014). Besides being computationally expensive, such models are sensitive to the selection of the basis functions, and require specific knowledge to appropriately calibrate them.

More recently, tensor-based approaches have been proposed and applied to analyze large-dimensional brain imaging data. Zhou et al. (2013) considers a framework that linearly models a given clinical outcome as a response and uses brain images as tensor covariates. Estimation in this setting is achieved via maximum likelihood and regularized tensor regression tools based on tensor decomposition are also used to determine which regions of the brain are associated with the clinical response. Tensor decomposition methods are covered in further detail in Sect. 2. Guhaniyogi et al. (2017) proposes a Bayesian tensor regression method in presence of a scalar response and a tensor predictor with shrinkage priors to identify cells in the tensor predictor significantly related to the scalar response. Tensor-based frequentist and Bayesian approaches for joint modeling of the BOLD response across all voxels in the form of a tensor have also been developed [e.g., Li and Zhang 2017; Guhaniyogi and Spencer 2018]. However, these approaches have not been developed for multi-subject studies and do not allow us to infer connectivity between brain regions.

1.2. Multi-subject and Connectivity Approaches

In order to overcome the computational challenge of having voxel-level data analyzed for multiple subjects, early approaches combined information across voxels, either using the general linear model (GLM) parameter estimates or residual variance. Two-stage methods, such as those by Bowman et al. (2008) and Sanyal and Ferreira (2012), fit subject-specific GLMs, and then use regularization on the parameter estimates to determine activation. Mixture models and nonparametric Bayesian models (Zhang et al. 2016; Xu et al. 2009) have also been proposed to analyze inter- and intra-subject variability, though they incur heavy computational cost for exact results via Markov Chain Monte Carlo (MCMC).

While there is considerable literature on activation-only models, literature on Bayesian functional connectivity offers a much smaller number of distinct approaches. Most of the available models consider connectivity across clusters of voxels that display similar activity patterns, or develop connectivity measures across predetermined ROIs. To this end, Patel et al. (2006a) and (2006b) discretized the fMRI time series between regions based on whether they had elevated activity according to a threshold, and then compared joint and marginal probabilities of elevated activity. Bowman et al. (2008) modeled similarity within and between ROIs based on estimates of elements of the covariance in a two-stage model. A Dirichlet process prior was used by Zhang

et al. (2014b) to cluster remote voxels together, asserting that the clustering inferred an inherent connectivity. Zhang et al. (2014a) went on to propose a dynamic functional connectivity model, estimating connective phases and temporal transitions between them.

1.3. Joint Estimation of Activation and Connectivity

As mentioned above, there are several Bayesian modeling frameworks for assessing activation or connectivity separately; however, models incorporating both of them jointly in multi-subject fMRI studies with voxel-level data are comparatively rare in the literature. In the recent past, Kook et al. (2019) proposed an approach in which a Dirichlet process (DP) mixture model is used to classify voxels as active or inactive via discrete wavelet transformations. The clustering of the voxels through time via the mixture components is then used to derive a measure of inter-voxel connectivity within- and between-subjects. While such model succinctly captures activation and connectivity, the use of a Dirichlet process hinders computational efficiency. Variational Bayes methods were used to speed-up computation, providing results in a fraction of the time that a full Markov chain Monte Carlo simulation would require; however, these variational approaches lead to approximate rather than exact posterior inference. In addition, the model of Kook et al. (2019) analyzes voxel-level connectivity, rather than region-of-interest connectivity. Connectivity-informed activation detection was proposed by Ng et al. (2011) using a two-step classical modeling process. First, a graphical LASSO or oracle-approximating shrinkage is used to cluster voxels into groups, and then cluster information is incorporated into the estimation for activation effects. This model differs from our proposal in that it does not provide robust uncertainty quantification, and regions cannot be defined *a priori*. In addition, activation is detected through an average over parcellated regions, and as such, voxel-level activation inference is not possible.

Our article proposes a multi-subject Bayesian tensor mixed effect model that provides exact posterior inference on voxel-level activation using the inherent spatial structure of fMRI and region-of-interest-level connectivity measures in a single model, while imposing shrinkage on both measures. Modeling activation and connectivity in a single step reduces model bias over a two-step modeling process by taking into account any interaction between the activation and connectivity components. In addition, our approach does not require a multiplicity correction method when detecting activation and/or connectivity. To elaborate further, the model envisions the BOLD response over all voxels together for a subject at any time as a tensor and regresses this tensor object on the activation-related predictors. A tensor decomposition representation is then used in conjunction with a novel prior structure to make the model more parsimonious while simultaneously capturing the underlying spatial structure of the data in an efficient manner. This is one of the key features of using a tensor representation, i.e., it enables the use of a tensor decomposition structure, which essentially treats the principle axes of the coordinate system as principal components, inherently preserving localized spatial dependence while reducing the number of model parameters. This is an advantage with respect to approaches such as those in Kook et al. (2019) which explicitly model local spatial correlation among voxels using spatially dependent prior distributions on voxel-specific activation coefficients, and hence becoming computationally prohibitive with a large number of voxels, particularly in large multi-subject studies. In addition, our proposed model incorporates subject-ROI-specific random effects with a Gaussian graphical prior, imposing regularization on the precision matrix of the effects between regions (Wang et al. 2014). Both the activation and connectivity parameters are then classified into zero- and nonzero-effect sizes using the sequential 2-means method proposed by Li and Pati (2017). As a result, the model produces accurate measures of voxel-wise activation and inter-regional connectivity with interpretable effect sizes without the need for fine-tuning hyperparameters, basis functions or multiplicity corrections. In addition, the model is computationally tractable to provide samples

from the exact posterior distribution for 2-D slices or 3-D volumes of brain images, as well as higher-order tensor images.

The model is tested with the Balloon Analog Risk Task data collected and analyzed by Schonberg et al. (2012), with the goal of replicating the findings of that study while providing additional inference about the connective networks between regions of interest in the Montreal Neurological Institute atlas (Collins et al. 1995; Mazziotta et al. 2001). More specifically, inference confirming activation in the frontal lobe and insula for the risk-processing task is shown. Comparisons are made to models with the Bayesian spike-and-slab prior and generalized double-Pareto prior on the tensor coefficient in order to show the effect of representing the tensor coefficient through a tensor decomposition.

The upcoming sections proceed as follows. The model, including the prior structure, is set forth in Sect. 2. This is followed by a section on posterior inference (Sect. 3) and a section on simulated studies to empirically validate the model (Sect. 4). Sensitivity to hyperparameter specification and comparisons with other models is also shown in Sect. 4. Section 5 then describes the multi-subject fMRI data from the balloon-analog risk-taking experiment in detail, as well as the results obtained from applying the proposed methodology to these data. Finally, Sect. 6 summarizes our findings and provides a description of some future extensions.

2. Methodology

This section begins by setting the tensor notation and the PARAFAC decomposition structure. The complete modeling framework, including prior structure and hyperparameter specification, is then detailed.

2.1. Notation and Preliminaries

A tensor $\mathbf{A}$ is a D -way array (also known as a D -th order tensor) of dimensions $p_1 \times \cdots \times p_D$ with $(i_1, \dots, i_D)$ -th cell entry denoted by $A[i_1, \dots, i_D] \in \mathbb{R}$, $i_1 = 1, \dots, p_1; \dots; i_D = 1, \dots, p_D$. For vectors $\mathbf{a}_1, \dots, \mathbf{a}_D$ of lengths $p_1, \dots, p_D$, respectively, define the outer product between the vectors, denoted by $\mathbf{A} = \mathbf{a}_1 \circ \cdots \circ \mathbf{a}_D$, as a D -way array with $(i_1, \dots, i_D)$ -th cell element $A[i_1, \dots, i_D] = \prod_{j=1}^D a_{j,i_j}$, where a_{j,i_j} is the i_j -th element of $\mathbf{a}_j$. $\mathbf{A}$ is referred to as a Rank 1 tensor with dimensions $p_1 \times \cdots \times p_D$. A rank- R tensor $\mathbf{A}$ is obtained by summing R Rank 1 tensors, $\mathbf{A} = \sum_{r=1}^R \mathbf{a}_{1,r} \circ \cdots \circ \mathbf{a}_{D,r}$. This is also referred to as the CP/PARAFAC decomposition of rank R (Kiers 2000) and is used due to its relative simplicity (Kolda and Bader 2009). In what follows, we will refer to $\mathbf{a}_{j,r}$'s as margins of the tensor $\mathbf{A}$.

2.2. Model Framework and Prior Structure

We assume that the whole tensor structure of the fMRI is partitioned into G distinct brain regions, and that we are interested in effectively measuring brain connectivity between these regions. Importantly, the proposed model does not assume that the voxel-level activation and the functional connectivity between predefined regions of interest are independent. As an aside, note that setting $G = 1$ reduces the model to a simpler tensor response regression model, which is explored in a single-subject context in Guhaniyogi and Spencer (2018). Our proposal extends that model to a multi-subject tensor response regression setting. In addition, the proposed model also takes inter-regional connectivity into account.

Let $\mathbf{Y}_{i,g,t}$ be the tensor of observed BOLD response data in brain region g for the i -th subject at the t -th time point. $\mathbf{Y}_{i,g,t}$ is observed in the form of a tensor with dimensions $p_{1,g} \times \cdots \times p_{D,g}$. In the context of fMRI data analysis, the tensor dimension for a fixed time, subject, and region, denoted as D , is two or three, depending on whether a single slice or regional volume is analyzed.

To simultaneously measure activation due to stimulus at voxels in the g -th brain region and connectivity among G brain regions, we employ an additive mixed effect model with tensor-valued BOLD response and activation-related predictor $x_{i,t} \in \mathbb{R}$,

$$\mathbf{Y}_{i,g,t} = \mathbf{B}_g x_{i,t} + d_{i,g} + \mathbf{E}_{i,g,t}, \quad (2)$$

for subject $i = 1, \dots, n$, in region of interest $g = 1, \dots, G$, and time $t = 1, \dots, T$. Elements in the error tensor $\mathbf{E}_{i,g,t}$ are assumed to be independent and identically distributed following a normal distribution with mean 0 and shared variance σ_y^2 , though our framework can be extended to incorporate temporally correlated errors. However, in testing with both simulated and task-based fMRI data, this does not appear to have a large effect on the model inference.

The model in (2) has 3 components, namely an activation component, a connectivity component, and an error component. The tensor coefficient $\mathbf{B}_g \in \mathbb{R}^{p_{1,g} \times \dots \times p_{D,g}}$ is used to infer the strength of the association between $x_{i,t}$ and each voxel in $\mathbf{Y}_{i,g,t}$. In particular, $B_g[i_1, \dots, i_D] = 0$ implies that the $(i_1, \dots, i_D)$ -th voxel in the g -th ROI is *not activated* by the stimulus. In fact, the activation pattern is typically sparse and localized with only a few nonzero elements in $\mathbf{B}_g$ (Olshausen and Field 2004). $d_{i,g} \in \mathbb{R}$ are region- and subject-specific random effects which are jointly modeled to borrow information across ROIs. This model views connectedness through the elements of the precision matrix corresponding to different, pre-specified regions of interest, rather than between individual voxels. In this way, the detected relationship between regions is not directly determined by the detected activation of a voxel. In the present context, the conditional distributions $(d_{i,g}, d_{i,g'}) | \{d_{i,g''} : g'' \neq g, g'\}$ are investigated to assess the strength of connectivity between a pair of regions. As part of the model development, we impose prior distributions that favor conditional independence between most pairs $d_{i,g}$ and $d_{i,g'}$, effectively favoring connectivity only among a few pairs of regions. Choosing a different number of regions will obviously change the network of regions that inference can be applied to and also the inference in the model parameters. However, since the partial correlation is used to measure connectivity, the connectivity measured between two given regions should not be greatly affected by adding or removing other regions from the model.

As mentioned above, the coefficient tensor $\mathbf{B}_g \in \mathbb{R}^{p_{1,g} \times \dots \times p_{D,g}}$ in Eq. (2) characterizes a sparse relationship between the tensor response and the time-varying covariate $x_{i,t}$ in region g . In order to achieve parsimony in the number of estimated parameters, $\mathbf{B}_g$ is assumed to have a rank R parallel factorization (PARAFAC) decomposition, which takes the form:

$$\mathbf{B}_g = \sum_{r=1}^R \boldsymbol{\beta}_{g,1,r} \circ \dots \circ \boldsymbol{\beta}_{g,D,r}, \quad (3)$$

with tensor margin effects $\boldsymbol{\beta}_{g,1,r}, \dots, \boldsymbol{\beta}_{g,D,r}$. The PARAFAC tensor decomposition dramatically reduces the number of parameters in $\mathbf{B}_g$ from $\prod_{j=1}^D p_{j,g}$ to $R \sum_{j=1}^D p_{j,g}$, where the level of parameter reduction depends on R . Note that a smaller value of R leads to more parsimony and computational gain, perhaps at the cost of inferential accuracy. By contrast, a choice of even moderately large R entails higher computation cost. Using R as a model parameter often increases computation cost and is deemed unnecessary (Guhaniyogi et al. 2017). In view of the earlier literature, this article proposes fitting the model with various choices of R and chooses the one that yields the lowest Deviance Information Criterion (DIC) (Gelman et al. 2014). More discussion on the choice of R is provided in the Simulation Studies section.

A critical question remains how to devise a prior distribution on the low-rank decomposition (3) to facilitate identifying geometric subregions in the tensor response which share an association

with the predictor. Additionally, the model intends to build joint priors on region-specific random effects $d_{i,g}$ s to assess connectivity patterns. The next two subsections propose careful elicitation of the prior distributions on $\mathbf{B}_g$ and $d_{i,g}$ to achieve our stated goals.

2.3. Multiway Stick-Breaking Shrinkage Prior on $\mathbf{B}_g$ to Assess Activation

Although the spike-and-slab priors for selective predictor inclusion (George and McCulloch 1993; Ishwaran et al. 2005) possess attractive theoretical properties and an easy interpretation, they often lose their appeal due to their inability to explore a large parameter space. As a computationally convenient alternative, an impressive variety of shrinkage priors (Carvalho et al. 2010; Armagan et al. 2013) in the context of ordinary Bayesian high-dimensional regression have been developed.

Shrinkage architecture relies on shrinking coefficients corresponding to unimportant predictors, while maintaining accurate estimation with uncertainty for important predictor coefficients. The existing shrinkage prior literature serves as a basis to the development of shrinkage priors on the tensor coefficients. However, constructing such a prior on $\mathbf{B}_g$ presents additional challenges. To elaborate on it, notice that proposing a prior on a low-rank PARAFAC decomposition of $\mathbf{B}_g$ is equivalent to specifying priors over tensor margins $\boldsymbol{\beta}_{g,j,r}$. Since every cell coefficient in $\mathbf{B}_g$ is a nonlinear function of the tensor margins, careful construction of shrinkage priors on $\boldsymbol{\beta}_{g,j,r}$ s is important to impose desirable tail behavior of $\mathbf{B}_g[i_1, \dots, i_D]$ parameters. To this end, this article employs a *multiway stick-breaking* shrinkage prior on $\mathbf{B}_g$ to ensure desirable tail behavior. More specifically, the following shrinkage prior is proposed on the tensor margins

$$\boldsymbol{\beta}_{g,j,r} \sim N(\mathbf{0}, \phi_{g,r} \tau_g \mathbf{W}_{g,j,r}), \quad \mathbf{W}_{g,j,r} = \text{diag}(\omega_{g,j,r,1}, \dots, \omega_{g,j,r,p_j}),$$

where

$$\omega_{g,j,r,\ell} \sim \text{Exp}\left(\frac{\lambda_{g,j,r}^2}{2}\right), \quad \lambda_{g,j,r} \stackrel{\text{iid}}{\sim} \text{Gamma}(a_\lambda, b_\lambda),$$

for $j = 1, \dots, D$ and $g = 1, \dots, G$. Flexibility in modeling tensor margins are accommodated by introducing $\mathbf{W}_{g,j,r}$ s. In fact, integrating out $\mathbf{W}_{g,j,r}$ and $\lambda_{g,j,r}$ yields a generalized double Pareto shrinkage prior for the elements of $\boldsymbol{\beta}_{g,j,r}$ conditional on $\phi_{g,r}$. The proposed prior defines a set of rank-specific scale parameters $\phi_{g,r}$ using a stick-breaking construction of the form $\phi_{g,r} = \xi_{g,r} \prod_{l=1}^{r-1} (1 - \xi_{g,l})$, $r = 1, \dots, R-1$, and $\phi_{g,R} = 1 - \sum_{r=1}^{R-1} \phi_{g,r} = \prod_{l=1}^{R-1} (1 - \xi_{g,l})$ that achieves efficient shrinkage across ranks, where $\xi_{g,r} \stackrel{\text{iid}}{\sim} \text{Beta}(1, \alpha_g)$. Finally, the global scale parameters are modeled as $\tau_1, \dots, \tau_G \stackrel{\text{iid}}{\sim} \text{Gamma}(a_\tau, b_\tau)$.

Without constraints on the values for $\phi_{g,r} \in \Phi_g$, where $r = 1, \dots, R$, identifiability issues arise in the posterior sampling for the variance terms for $\boldsymbol{\beta}_{g,j,r} \in \mathbf{B}_g$. In order to address this issue, a stick-breaking structure is imposed on $\phi_{g,r}$'s, as described above. In effect, this prevents $\phi_{g,r}$'s from switching labels across ranks in which the variance of $\boldsymbol{\beta}_{g,j,r}$ may be close together. The result of this constraint is a more stable MCMC for the posterior draws of $\boldsymbol{\beta}_{g,j,r}$. The tuning parameter α_g in the stick-breaking construction determines which tensor rank R is favored by data. In particular, $\alpha_g \rightarrow 0$ favors small values of most $\phi_{g,r}$ a-priori. Therefore, a data-dependent learning of α_g is essential in order to tune to the desired sparsity in $\mathbf{B}_g$. The subsection on hyperparameter specification discusses a model-based choice of α_g , along with the specific choices for $a_\lambda, b_\lambda, a_\tau$, and b_τ .

2.4. Bayesian Graphical Lasso Prior for Modeling Connectivity

Following Wang et al. (2012), to capture connectivity between different regions for individuals, $d_{i,g}$ s are jointly modeled with a Gaussian graphical lasso prior. To be more precise,

$$\begin{aligned} \mathbf{d}_i &= (d_{i,1}, \dots, d_{i,G})' \sim N(\mathbf{0}, \mathbf{\Sigma}^{-1}), \quad i = 1, \dots, n, \\ p(\sigma|\zeta) &= C^{-1} \prod_{k < l} [DE(\sigma_{kl}|\zeta)] \prod_{k=1}^G \left[\text{Exp}\left(\sigma_{kk} \middle| \frac{\zeta}{2}\right) \right] \mathbf{1}_{\mathbf{\Sigma} \in \mathcal{P}^+}, \end{aligned} \quad (4)$$

where $\mathcal{P}^+$ is the class of all symmetric positive definite matrices and C is a normalization constant. $\sigma = (\sigma_{kl} : k \leq l)$ is a vector of upper triangular and diagonal entries of the precision matrix $\mathbf{\Sigma}$. Using properties of the multivariate Gaussian distribution, a small value of σ_{kl} stands for weak connectivity between ROIs k and l , given the other ROIs. In fact, $\sigma_{kl} = 0$ ($k < l$) implies that there is no connectivity between ROIs k and l , given the other ROIs. Thus, to favor shrinkage among off-diagonal entries of $\mathbf{\Sigma}$ for drawing inference on connectivity between pairs of ROIs, the Bayesian graphical lasso prior employs double exponential prior distributions on the off-diagonal entries of this precision matrix. The diagonals of $\mathbf{\Sigma}$ are assigned exponential distributions. Let $\eta = (\eta_{kl} : k < l)$ be a set of latent scale parameters. Using the popular scale mixture representation of double exponential distributions (Wang et al. 2012), we can write

$$p(\sigma|\zeta) = \int p(\sigma|\zeta, \eta) p(\eta|\zeta) d\eta,$$

with $p(\sigma|\zeta, \eta)$ given by

$$p(\sigma|\zeta, \eta) = C_\eta^{-1} \prod_{k < l} \left[\frac{1}{\sqrt{2\pi\eta_{kl}}} \exp\left(-\frac{\sigma_{kl}^2}{2\eta_{kl}}\right) \right] \prod_{k=1}^G \left[\frac{\zeta}{2} \text{Exp}\left(-\frac{\zeta}{2}\sigma_{kk}\right) \right] \mathbf{1}_{\mathbf{\Sigma} \in \mathcal{P}^+}, \quad (5)$$

where C_η is the normalizing constant, which is an analytically intractable function of η . The mixing density of η in the representation above is given by

$$p(\eta|\zeta) \propto C_\eta \prod_{k < l} \frac{\zeta^2}{2} \exp\left(-\frac{\zeta^2}{2}\eta_{kl}\right). \quad (6)$$

The hierarchy is completed by adding a Gamma prior on ζ , $\zeta \sim \text{Gamma}(a_\zeta, b_\zeta)$.

Finally, an inverse gamma prior $\sigma_y^2 \sim \text{Inverse Gamma}(a_\sigma, b_\sigma)$ is used on the variance parameter σ_y^2 . A plate diagram of the model structure can be seen in Fig. 1.

2.5. Hyperparameter Specification

The hyperparameters α_g in the stick-breaking construction are assigned a discrete uniform prior over 10 equally spaced values in the interval $[R^{-D}, R^{-0.10}]$, which will allow the data to dictate the level of sparsity appropriate for the prior (Guhaniyogi et al. 2017). The posterior updating of α_g under the griddy-Gibbs sampling algorithm can be found in the Online Resource. Although our extensive exploration with simulation data shows that a careful choice of fixed α_g value is able to produce inference as good as is offered by allowing α_g to vary, in practice, it is not clear how to choose such values. Consequently, a naive choice of α_g may offer incorrect

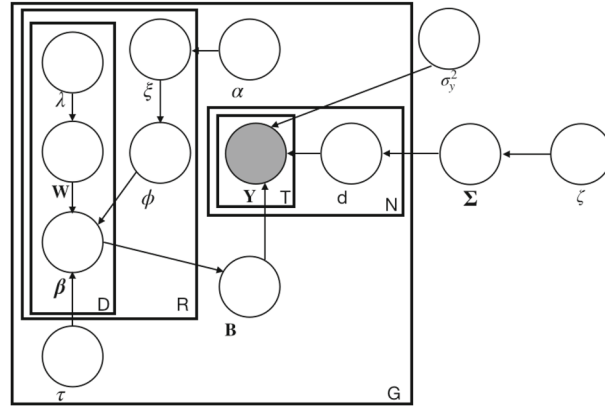


FIGURE 1.
Plate diagram representation of the proposed model

conclusion from the model. On the other hand, the posterior updating of α_g in each iteration adds burden to computation. Therefore, this article allows posterior updating of α_g in the MCMC iteration, and fixes α_g values after the burn-in, at the value drawn in the n_{burn} -th iteration, where n_{burn} is the burn-in for the MCMC chain. As we discuss later, $n_{\text{burn}} = 100$ for simulation studies and equals 200 for the real data analysis. This both allows α_g to be learned while reducing any research bias stemming from a user dependent fixed value for α_g . The strategy works for various simulation studies and moderate perturbation of the prior range seems to produce robust inference. The values chosen for a_λ and b_λ have a strong effect on the shrinkage properties of the generalized double-Pareto prior, and setting $a_\lambda = 3$ and $b_\lambda = \sqrt[2D]{a_\lambda}$ prevents the prior for $\lambda_{g,j,r}$ from allowing for insufficient variance for $B_g[i_1, \dots, i_D]$ to detect nonzero coefficients. Similar to Guhaniyogi et al. (2017), the hyperparameters a_τ and b_τ are set to $D - 1$ and $R^{1/D-1}$, respectively, in order to prevent overshrinkage with higher tensor response dimensions. Following Wang et al. (2014), a_ζ and b_ζ are set to 1 and 0.01, respectively, in order to preserve relative noninformativity of the Gaussian graphical prior. Finally, for both simulation studies and the real data analysis, a_σ and b_σ were set to be 1 and $-\log 0.95$, respectively. While these hyperparameters are specified to provide readers a specific set of choices and they produce desirable results, we establish in the Simulation Studies section that the inference is fairly robust with moderate perturbation of these hyperparameters.

3. Posterior Computation

The model framework and prior structure allow sampling from the posterior distribution using the Markov Chain Monte Carlo (MCMC) algorithm outlined in the Online Resource. In order to speed convergence of the MCMC chain, values for $\beta_{g,j,r}$ are initialized using the singular value decomposition of the mode- j matricization of each $\hat{\mathbf{B}}_{g,\text{MLE}}$, the maximum likelihood estimate of the tensor-valued coefficient for the activation, assuming no connectivity component. This particular initialization method limits the rank R of the model to an upper bound of $\min_{g,j} p_{g,j}$. The posterior distributions of unknown quantities of interest are approximated by their empirical distributions from post burn-in MCMC samples.

Of particular interest in neuroscience is the assessment of whether a brain voxel is active or not, which, in our modeling framework translates to verifying whether $B_g[i_1, \dots, i_D]$ is nonzero for voxel $(i_1, \dots, i_D)$. It is well-acknowledged that the problem of selecting important cell coefficients

is a challenging task when $\mathbf{B}_g$ is assigned a continuous shrinkage prior, since none of the cell coefficients is exactly zero in any MCMC iteration. The sequential 2-means method recently developed by Li and Pati (2017) and outlined in Algorithm 1 is thus used to classify whether a coefficient is zero or nonzero from their post burn-in MCMC iterates. This is done separately for each reconstructed regional coefficient tensor $\mathbf{B}_g$. Following the suggestion in Li and Pati (2017), the value of b in Algorithm 1 is set to be the median of the standard deviations of the elements within $\mathbf{B}_g$. Within the estimates obtained through the use of the sequential 2-means method, nonzero-valued voxels are considered to be active.

Algorithm 1: Sequential 2-means for posterior draws $s = 1, \dots, S$ for parameter θ

Result: Final estimate of θ with small elements set to be equal to 0

for $s \leftarrow 1$ **to** S **do**

 Cluster the absolute value of elements in $\theta^{(s)}$ into two clusters, $\mathcal{A}$ and $\mathcal{B}$, where $\bar{\mathcal{A}} \leq \bar{\mathcal{B}}$, where $\bar{\mathcal{A}}$ and $\bar{\mathcal{B}}$ denote the mean of elements in the clusters $\mathcal{A}$ and $\mathcal{B}$, respectively;

 Cluster the elements of $\mathcal{A}$ into two clusters, $\mathcal{A}$ and $\mathcal{A}'$ such that $\bar{\mathcal{A}} < \bar{\mathcal{A}}'$;

while $|\bar{\mathcal{A}} - \bar{\mathcal{A}}'| > b$ **do**

 Cluster the elements of $\mathcal{A}$ into two clusters, $\mathcal{A}$ and $\mathcal{A}'$ such that $\bar{\mathcal{A}} < \bar{\mathcal{A}}'$;

end

 The number of elements remaining in $\mathcal{A}$ is the estimated number of true zero-valued elements, $n_z^{(s)}$, in $\theta^{(s)}$

end

Find $\hat{n}_z = \text{median value of } n_z$;

Find $\hat{\theta} = \text{median values of the elements in } \theta^{(1:S)}$;

Set elements in $\hat{\theta}$ with the $\hat{n}_z$ smallest absolute values to 0

In order to obtain an interpretable measure of the connectivity between regions, the partial correlation between regions is examined. Since the partial correlation accounts for the correlation between two regions after removing the influence of all other regions (Das et al. 2017), it is expected to be an appropriate measure of pairwise connections. First, the partial correlation matrix was calculated from each posterior draw of Σ using the `prec2part` function in the `DensParcorr` package in *R* (Wang et al. 2018). Next, the sequential 2-means method (Li and Pati 2017) was used on the upper-triangular elements of the partial correlation matrix to classify them as zero or nonzero. Regions with nonzero partial correlations are said to be connected (Warnick et al. 2018).

4. Simulation Studies

To validate the proposed methods, we simulate synthetic data with similar structure to that found in data collected from human fMRI studies. The tensor responses are simulated considering a block experimental design from the likelihood in (2). In each simulation study, we construct $G = 10$ different coefficient tensors corresponding to disjoint spaces, hereafter referred to as regions. For ease of visualization, the coefficient tensors are created to be three-dimensional, but can be generalized to any arbitrary dimension D so that the method may be applied to other scenarios. Throughout the simulation study, a sample size of $n = 20$ subjects is used, with the number of time points per subject being fixed at $T = 100$.

The covariate, $x_{i,t} = x_t$, was set to be the same for all of the subjects, without any loss of generality. A block experimental design is employed to generate the covariate, which consists of several discrete epochs of activity-rest periods, with the “activity” representing a period of stimulus presentations, and the “rest” referring to a state of rest or baseline. These activity-rest

periods are alternated throughout the experiment to ensure that signal variation, scanner sensitivity, and subject movement have the similar effect throughout the experiment. To simulate activity-rest periods, we use the stimulus indicator function z_t as:

$$z_t = \begin{cases} 1, & \text{for } kP < t < kP + P/2, \quad k = 0, 1, \dots \\ 0, & \text{otherwise} \end{cases}$$

for all t , given a defined period P for the block design. In our simulations, P is set to be 30. Next, the `canonicalHRF` function in the `neuRosim` package in R (Welvaert et al. 2011) is used to convolve the stimulus indicator z_t with the double-gamma canonical hemodynamic response function (HRF), which corrects for the expected delay between a stimulus and the resultant physiological response in the brain (Friston et al. 1998). This HRF is set using the default function values in `neuRosim` to have a delay of response relative to onset equal to 6 time steps, a delay of undershoot relative to onset of 12, a dispersion of response equal to 0.9, a dispersion of undershoot equal to 0.9, and a scale of undershoot equal to 0.35. The resulting covariate x_t is plotted in Fig. 1 of the Online Resource. The dimensions of response tensor margins $p_{1,g}$, $p_{2,g}$, and $p_{3,g}$ all set to 10 for each region g , resulting in 10 regions with 1,000 voxels in each.

In order to demonstrate the effectiveness of the shrinkage component of the model, the true tensor coefficient values were randomly assigned using the `specifyregion` function from `neuRosim`. This function allows for the definition of tensors such that nonzero elements are spatially contiguous spheres. In this simulation, the coefficient tensors are designed such that all elements took the value of either zero or 0.05. In real fMRI data, activation is typically observed in a small number of voxels/regions. Therefore, we set the sizes of the true activated cells in our simulated data to be no greater than 5% of the total tensor size. The true values for a slice of one of the coefficient tensors can be seen in Fig. 2.

The contrast-to-noise ratio, defined as $B_{g,v}/\sigma_y$ for $B_{g,v} \neq 0$, was set to be equal to 0.05, which is proposed as a realistic value for neuroimaging data by Welvaert and Rosseel (2013). The connectivity between tensor regions was simulated by setting two pairs of the ten regions to have a region-wide correlation of 0.9, while all other regions were assigned correlations of zero. A covariance matrix (Σ^{-1}) was created from this correlation matrix, and the region effects for subject i were simulated from a multivariate normal distribution with mean zero and covariance Σ^{-1} . The signal-to-noise ratio, defined as $\Sigma_{g,g'}^{-1}/\sigma_y^2$ for $\Sigma_{g,g'}^{-1} \neq 0$ was set to 1, a realistic value based on Welvaert and Rosseel (2013). This quantity can be thought of as the relative effect of the connectivity on the observed response tensors. Finally, the observation-level variance (σ_y^2) was set to be 1.

4.1. Competitors

We fitted our proposed Bayesian model to the simulated data using different choices of rank R . In most of the real life applications, small values of R are sufficient to attain the desired inference, so models up to Rank 7 were fit to the simulated data.

The performance of the proposed model is compared to that obtained from the following models: a vectorized model with a Generalized Double Pareto (GDP) shrinkage prior on the activation coefficients and a Gaussian graphical prior on the connectivity parameters, referred to as the *vectorized-GDP* approach; another vectorized model with a spike-and-slab prior on the activation coefficients referred to as the *spike-and-slab* approach; a general linear model (GLM). Details about these models are provided below.

The vectorized GDP model vectorizes the tensor response and builds a voxel-specific linear mixed effect model by regressing the response on predictors, followed by jointly estimating

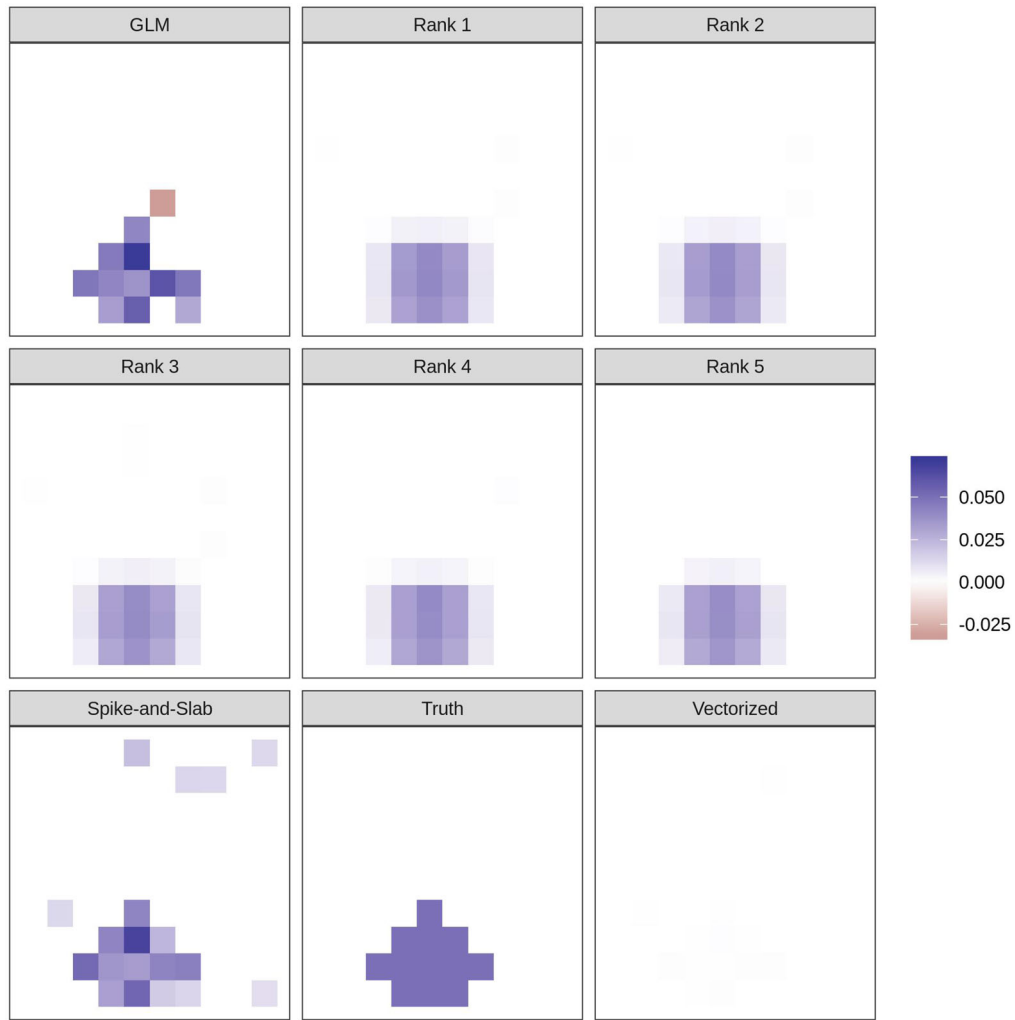


FIGURE 2.

Rank model estimates and true value for a single slice of a three-dimensional coefficient tensor. Estimates are found using the sequential 2-means variable selection method (Li and Pati 2017). The spike-and-slab and vectorized model estimates are also included for comparison

the voxel-specific regression coefficients using a shrinkage prior distribution. More precisely, if Y_{g,i,t,v_1,v_2,v_3} is the response at voxel (v_1, v_2, v_3) in region g at time t for individual i , this mixed effect model proposes

$$Y_{g,i,t,v_1,v_2,v_3} \stackrel{\text{ind.}}{\sim} N(b_{g,v_1,v_2,v_3}^* x_{i,t} + d_{g,i}^*, \sigma^{*2}), \quad \beta_g^* \sim N(\mathbf{0}, \tau_g^* \mathbf{W}_g^*), \quad (7)$$

where $\beta_g^* = (b_{g,v_1,v_2,v_3}^* : v_1 = 1 : p_{g,1}, v_2 = 1 : p_{g,2}, v_3 = 1 : p_{g,3})' \in \mathbb{R}^{p_{g,1} \times p_{g,2} \times p_{g,3}}$ is the vector of fixed effects and $\mathbf{W}_g^* = (\omega_{g,v_1,v_2,v_3}^* : v_1 = 1 : p_{g,1}, v_2 = 1 : p_{g,2}, v_3 = 1 : p_{g,3})$. The random effects $d_{g,i}^*$'s are jointly assigned a Gaussian graphical prior similar to (4). The hierarchical specification is completed by assigning $\tau_g^* \sim \text{Gamma}(a_\tau, b_\tau)$, $\omega_{g,v_1,v_2,v_3}^* \sim \text{Exp}\left(\frac{\lambda_g^{*2}}{2}\right)$, $\lambda_g^* \sim$

$\text{Gamma}(a_\lambda, b_\lambda)$. The vectorized GDP prior is a shrinkage prior on activation coefficients, which are envisioned as approximations to the Bayesian variable selection priors. Hence, as it is also the case with the spike-and-slab priors below, they take care of the multiplicity issues. Comparison with this vectorized-GDP reveals the advantage of retaining the tensor structure of the response to capture underlying local spatial structure while simultaneously inferring connectivity, as well as the advantage due to the parsimony offered by the PARAFAC decomposition. We also attempted to implement a spatially varying coefficient (SVC) model (Zhang et al. 2015) and found it to be extremely computationally demanding due to large matrix inversions in each MCMC iteration. Hence the comparison with SVC is not reported.

The spike-and-slab model utilizes the spike-and-slab prior, introduced in Eq. (1) with $v_0 = 0.1$ for the “spike,” or the part of the prior distribution with a high density around zero, and $v_1 = 10$ for the “slab,” or the part of the prior distribution centered around zero, but with a much lower density around zero itself. The probability of inclusion in the spike prior component for voxel v , γ_v , was assigned a Bernoulli prior distribution with probability π_v . Finally, each π_v was assigned a Beta(2, 2) prior distribution. As shown by Scott and Berger (2010), keeping π_v random as opposed to fixed addresses the multiple correction issue in the spike-and-slab prior. As in the vectorized-GDP competitor, the random effects $d_{g,i}^*$ for the subjects and regions are assigned a Gaussian graphical prior as in (4).

Rather than using a cutoff to determine whether there is activation for a particular voxel within the coefficient tensor and for consistency in the comparison with the proposed approach, the sequential 2-means method outlined in Algorithm 1 is used on the posterior values of the coefficient tensor for the vectorized-GDP and spike-and-slab models. Similarly, the sequential 2-means method is also used to infer connectivity in both these competitors.

Finally, in order to improve the interpretability of the proposed model for those familiar with neuroimaging, the general linear model (GLM) is also used for comparison. It is of importance to note that the GLM is only placed on the activation components, as it fits each voxel to the covariate separately, which assumes that each voxel is independent of the others. In addition, sparsity in the activation cannot be built into the model assumptions, which has an effect on the uncertainty quantification of the tensor coefficient elements. The connectivity estimates do not have any uncertainty quantification at all. In order to determine activation, the Benjamini–Hochberg method for multiple testing correction Benjamini and Hochberg (1995) is used, as it allows for the control of the false discovery rate, which is desirable in high-dimensional regression settings. The residuals from the GLM are then summed within each region of interest and used to estimate the covariance between the regions. The covariance estimate is then used to estimate the precision and the partial correlation.

4.2. Comparison Metrics

MCMC is run for 1,100 iterations for all competitors, with a 100-iteration burn-in and the remaining used for inference. The assessment of convergence is made by the Raftery–Lewis diagnostic test implemented in the R package “coda.” It shows a median effective sample size of 1,000 for the elements of all the $\mathbf{B}_g$ s in the Rank 1 model, and around 840 for the Rank 2 through 7 models. Comparisons among competitors are based on a model fitting statistic and point estimation of $\mathbf{B}_g$ ’s. The accuracy of detecting active and inactive voxels for each competitor are also reported. Finally, we also compare competitors in terms of identifying connectivity between regions as the main goal of our proposed approach is to jointly detect activation and connectivity.

Model fitting is compared using the deviance information criterion (DIC), defined in Gelman et al. (2014) as $\text{DIC} = -2 \log p(\mathbf{Y}|\hat{\mathbf{B}}, \hat{\mathbf{d}}, \mathbf{X}, \hat{\sigma}_y^2) + 2p_{\text{DIC}}$, where $p_{\text{DIC}} = 2 \left(\log p(\mathbf{Y}|\hat{\mathbf{B}}, \hat{\mathbf{d}}, \mathbf{X}, \hat{\sigma}_y^2) - \frac{1}{S} \sum_{s=1}^S \log p(\mathbf{Y}|\mathbf{B}^s, \mathbf{d}^s, \mathbf{X}, \sigma_y^{2(s)}) \right)$, $\hat{\theta}$ is the posterior mean of any parameter θ , S is the total

number of post burn-in posterior samples. The superscript s denotes s -th post burn-in posterior sample for a parameter, $\mathbf{Y}$, $\mathbf{X}$ are the collection of all responses and predictors, respectively.

For comparison between the models in terms of point estimation of $\mathbf{B}_g$'s, we compute the square root of the mean squared error (RMSE) between the estimated tensor coefficient and true tensor coefficient, $\sqrt{\sum_{g=1}^G \sum_{v \in \mathcal{R}_g} (\bar{B}_{g,v} - B_{g,v}^0)^2}$, where $\mathcal{R}_g$ represents region g , $B_{g,v}^0$ and $\bar{B}_{g,v}$ are the true and the posterior mean of the v the cell coefficient in the g -th region, respectively. In addition, as mentioned above, given the posterior mean estimates of $B_{g,v}$, sequential 2-means approach (Li and Pati 2017) is employed to identify active and inactive voxels. The true positive rate (TPR) and false positive rate (FPR) are computed for the different approaches. Finally, a summary of the performance of the connectivity for each model is given as the Frobenius norm of the difference between the point estimate of the partial correlation ($\hat{\Omega}$) and the true partial correlation matrix (Ω), with $\Omega = \Sigma^{-1}$. The Frobenius norm is defined for matrix $\mathbf{A}$ as $\|\mathbf{A}\|_F = \sqrt{\text{trace}(\mathbf{A}'\mathbf{A})}$.

4.3. Results

In order to clearly define applications under which the proposed tensor model is expected to perform well, the following research hypotheses will be tested. Under the structure of the proposed tensor model, it is expected that sparse, hypercubic activation regions will be recovered well by the model. Sparse connectivity is also expected to be recovered effectively if the true partial correlation between regions is far enough away from zero. The model will also be tested in cases in which the observation error in the simulated data is strongly autocorrelated. The tensor models demonstrate benefit in terms of activation estimation, as seen in Fig. 2. The slice of activation data shown was selected by finding the slice of one of the tensor coefficients that had the most true activation to clearly show the differences in the patterns of activation detection. The proposed model excels in its ability to identify activation only in voxels in and around where there is true activation, adequately capturing the localized spatial structure underlying the activation mechanism. This is an advantage over all other competitors. In particular, note that the spike-and-slab approach leads to a relatively large number of isolated false positives, while the vectorized-GDP identifies no active sites.

Performance measures for all the competitors are summarized in Table 1. Using the DIC as a model selection criterion, the Rank 3 model is chosen among the proposed models. Overall, tensor models with rank greater than 1 outperform the vectorized-GDP in terms of sensitivity, with a sensitivity that is on par with the spike-and-slab. Due to the shrinkage imposed by the Bayesian models, the GLM has a better sensitivity and specificity for activation; however, it does not offer reliable uncertainty quantification, spatially localized activation, or reliable estimates of connectivity. Note that the spike-and-slab and vectorized-GDP competitors use far more parameters in their hierarchical models than the proposed tensor regression models to estimate the tensor coefficient. Therefore, the spike-and-slab and vectorized-GDP competitors are not compared with respect to the deviance information criterion, as we have found this criterion to be unreliable to compare models that have very large differences in the number of parameters due to underestimation of the penalty term.

When viewing Table 1, it is important to keep in mind that we are interested in joint detection of activation and connectivity. Similar to the tensor coefficient, the sequential 2-means method (Li and Pati 2017) is used on the off-diagonal elements of the partial correlation matrix calculated from the precision matrix Σ to recover the connectivity structure among regions in the simulated data. In terms of the connectivity we see that the best performance is obtained by the tensor model of Rank 3 and the worst is that obtained by the GLM approach. All unconnected regions are classified as having a partial correlation of zero, and the connected regions have nonzero partial correlations. The estimates from the model with the optimal rank as determined by the DIC and

TABLE 1.

Performance diagnostics based on 1,100 draws from the posterior distribution with multiple different models using the same simulated data

	# Parameters in $\mathbf{B}$	Time (h)	DIC	RMSE for B	Sensitivity	Specificity	$\ \hat{\boldsymbol{\Omega}} - \boldsymbol{\Omega}\ _F$
Rank 1	300	3.36	577863793	0.0247	0.4537	0.8124	1.1928
Rank 2	600	4.08	577825593	0.0251	0.6146	0.7896	1.0930
Rank 3	900	4.93	577724587	0.0252	0.6098	0.7891	0.9623
Rank 4	1,200	5.55	577877043	0.0252	0.6098	0.7911	0.9822
Rank 5	1,500	6.33	577823303	0.0252	0.6146	0.7934	1.1084
Rank 6	1,800	7.07	577849650	0.0252	0.6098	0.7945	1.1701
Rank 7	2,100	7.23	577806310	0.0252	0.6146	0.7954	1.1366
Vectorized	10,000	1.42		0.0232	0.5756	0.8871	0.9989
Spike-and-Slab	10,000	2.22		0.0241	0.6049	0.8696	1.2795
GLM	10,000	0.02		0.0037	0.8634	0.9955	2.7857

For the performance measures of the Bayesian models, the first 100 draws from the posterior distribution are discarded as a burn-in

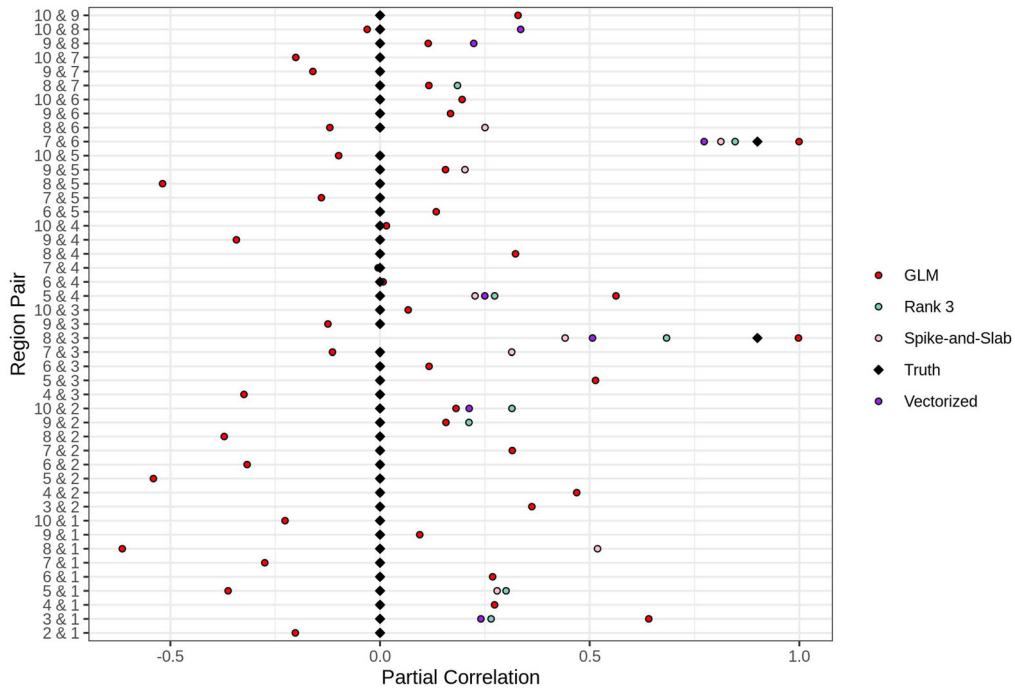


FIGURE 3.

Estimates of the partial correlation for all possible region pairs after using the sequential 2-means method from Li and Pati (2017). The true partial correlation values for all region pairs are shown for comparison

the competitor models are shown in Fig. 3, which shows that the effect sizes are smaller than the true generative values. However, in settings with low signal-to-noise ratio and sample size, the proposed tensor model with Rank 3 generally gives nonzero values for region pairs with true nonzero values, while leading to significant shrinkage in the true zero values. The GLM on the other hand does not adequately estimate connectivity.

TABLE 2.
Comparison of performance for correlated and uncorrelated error

	Autocorrelated error		Uncorrelated error	
	Sensitivity	Specificity	Sensitivity	Specificity
Rank 1	0.5660	0.4836	0.5660	0.5290
Rank 2	0.5660	0.4182	0.5660	0.5090
Rank 3	0.5660	0.4134	0.5660	0.4979
Rank 4	0.5660	0.4113	0.5660	0.4836
Rank 5	0.5660	0.4097	0.5660	0.4667
Vectorized	0.5660	0.4266	0.5660	0.5026
Spike-and-Slab	0.5660	0.4261	0.5660	0.5048
GLM	0.0283	0.9974	0.0000	1.0000

4.4. Sensitivity Analyses

4.4.1. Temporally Correlated Errors Given that the model makes assumptions about independent errors, simulation tests were also conducted under data generation scenarios that violate this assumption. Therefore, in addition to having a scenario with independent errors $e_{g,i,t,\ell} \in \mathbf{E}_{g,i,t}$ as proposed in the original setting, we considered a new smaller scale simulation scenario in which the errors have an autoregressive structure of order 1. That is, $e_{g,i,t,\ell} = 0.9e_{g,i,t-1,\ell} + u_{g,i,t,\ell}$, with $u_{g,i,t,\ell}$ assumed to be independent and identically distributed with mean zero and variance $\sigma_y^2 = 1$. This structure results in simulated data that are highly autocorrelated. In this new scenario, we simulated $G = 5$ regions, each with response tensors of size $\mathbf{Y}_g \in \mathbb{R}^{20 \times 20}$ for $n = 20$ subjects and $T = 100$ time steps. Furthermore, in order to directly show the effect of the autocorrelated errors on the model performance, an otherwise identical dataset of the same size was created in which the error in the data generation model was not autocorrelated.

For each of the two settings in this new simulation scenario, $\hat{\mathbf{B}}_g$ for the spike-and-slab competitor, the vectorized-GDP competitor, the GLM and the tensor model are created using the sequential 2-means method. The sensitivity and specificity were then found for each model, and can be seen in Table 2. It is important to note that the GLM does not do well in this scenario due to the low contrast-to-noise ratio combined with a smaller sample size in which only around 5% of the voxels in the coefficient tensor are actually nonzero. In spite of using the same multiple corrections method used for the GLM in the calculation of Table 1, in this case the GLM leads to a very poor sensitivity of 0.0283. Note that the specificity and sensitivity also go down for the other models also due to the smaller sample size, but the key part here is that the sensitivity measures are not affected by the induced temporal autocorrelation in the tensor models.

4.4.2. Contrast-to-Noise Comparisons Next, we ran a set of scenarios in which the only change in the simulated data was the contrast-to-noise ratio, something that should have a significant impact on the activation inference. In each simulated dataset, again, the number of subjects was set to $n = 20$, the number of time steps was set to $T = 100$, the number of regions was set to $G = 5$, the signal to noise ratio was set to 5, and the observation noise was set to $\sigma_y^2 = 1$. Figure 4 shows the sensitivity and specificity for different values of the contrast-to-noise ratio. This shows that the proposed model is more sensitive than Bayesian competitors at low contrast-to-noise levels. This also shows that the higher-rank models have a higher specificity than the Bayesian competitors at higher contrast-to-noise ratios.

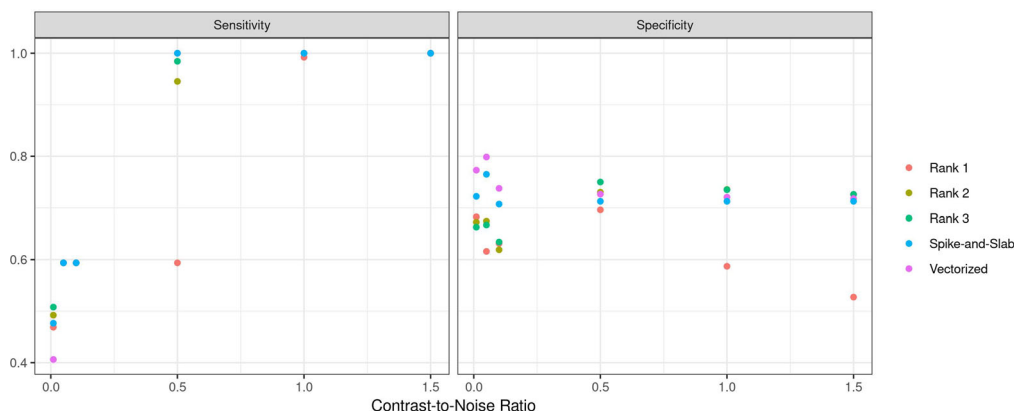


FIGURE 4.
Sensitivity and specificity under varying contrast-to-noise ratios

4.4.3. Hyperparameter Sensitivity Finally, in order to test the robustness of the model to choices of the hyperparameters, a grid of hyperparameter values was made by scaling each of the “standard” values for a_λ , b_λ , a_τ , b_τ , a_ζ , b_ζ , a_σ , and b_σ , defined in Sect. 2.5, by 0.01, 1, and 100, resulting in 6,561 different combinations. Of these, 100 settings were randomly sampled from the list and then tested with tensor model corresponding to Rank 3. We graphed boxplots of the RMSE as well as length and coverage of 95% CI for all these hyperparameter combinations (these are available in Fig. 2 of the Online Resource). The results are fairly robust with all three metrics varying within a small range under all different hyperparameter combinations. Overall, the simulation study reveals excellent recovery of activation and connectivity among regions by the proposed model. Although the computation time for the proposed model may be a bit on the higher side, the burden is somewhat lessened by the rapid MCMC convergence for the model parameters allowing accurate inference even with a small burn-in.

5. Real Data Analysis

We analyze data collected in a study examining the fMRI scans of individuals undergoing a test which introduces risk-taking scenarios. This study is known as the balloon-analog risk-taking task experiment, which requires participants to make active decisions, and whose design has been found to correlate with a number of naturalistic risk-taking behaviors (Schonberg et al. 2012). The data are available from the OpenfMRI project and the OpenNeuro platform at <https://openneuro.org/datasets/ds000001/versions/00006?app=MRIQC&version=33&job=5978f5dca1f52600019e85c4>. It consists of 16 individuals who were scanned using a 3T Siemens AG Allegra MRI machine in the Ahmanson-Lovelace Brain Mapping Center at UCLA. While in the scanner, the subjects inflated simulated balloons. A trial is defined as a balloon that can be pumped a certain number of times. Each trial could end in one of two ways. First, the subject could “cash-out” at any point during the trial and add the cumulative winnings for that balloon to their collective “bank.” Second, upon pumping, a balloon may explode and the participants would lose the cumulative winnings for that balloon and nothing would be added to their collective “bank.” The subjects interacted with the simulation by pushing one of two buttons with their right pointer finger or right middle finger. Each trial began with winnings of \$0.25, displayed below the balloon, and each successive pump added \$0.25 to the cumulative winnings for that balloon. The balloons were red, green, or blue in color, and the maximum number of pumps for

a balloon was drawn from a discrete uniform distribution between 1 and 8, 12, or 16, depending on the color of the balloon. Intermittently, subjects would be shown a gray control balloon with a maximum of 12 pumps that did not explode and did not have any associated monetary value. Unlike with the colored balloons, subjects did not have the option to “cash-out” when inflating the control balloon. Each run for each subject was ten minutes in length, during which each color balloon could be presented no more than 12 times.

The preprocessing was done using FSL following what was done by Schonberg et al. (2012) as closely as possible. The fMRI has a repetition time (TR) of 2 seconds. In order to allow for T1 equilibrium effects, the first two scans were dropped. The EPI images were motion corrected, then high-pass filtered using a Gaussian least-squares linear fit with $\sigma = 50.0$ seconds. Brain extraction was done using the BET function in FSL. The anatomic (T1-weighted) scans were registered using an affine transformation to standard Montreal Neurological Institute (MNI) space, and the EPI scans were then registered to each subject’s corresponding anatomic scan. Finally, the data were spatially smoothed using a Gaussian kernel with a 5 mm full-width half-maximum (FWHM). As these methods are implemented on the whole brain volumes all at once, the EPI scans were downsampled to have voxels with volume 8 mm^3 to find areas of increased activity within the entire brain in order to choose slices within the brain that can be analyzed at a higher resolution with voxels of volume 2 mm^3 . For both cases, data were separated into 9 regions of interest based on the MNI structural atlas provided within FSL (Collins et al. 1995; Mazziotta et al. 2001). The MNI structural atlas is a hand-segmented atlas developed by Mazziotta et al. (2001) and Collins et al. (1995) and distributed within the FSL library of neuroimaging tools (Jenkinson et al. 2012). Choice of atlas is dependent on particular hypotheses, and can have a large effect on the connectivity inference of the proposed model. Splitting large regions of interest into two subregions has been found to be practically unnecessary, as the partial correlation within these subregions has been observed to be very high, even when the regions that were split apart are not physically contiguous within the brain. This does not present meaningful additional inference in the results of the model, so the regions of interest are kept as defined by the structural atlas in order to ease interpretation of the results. The voxel-level activation results were observed to be practically unchanged by splitting the regions of interest. One subject (subject 15) was removed from the dataset after exploratory data analysis showed unusually high variance and temporal patterns that were not present in the scans of other subjects. In the whole-brain analysis, the regions of interest varied in size between 99 voxels and 1,667 voxels after the BOLD response tensors were multiplied by binary masks, with a median region volume of 649 voxels.

To measure the level of risk being processed by a subject at a given time, we slightly modified the procedure used in Schonberg et al. (2012), described as follows. Begin with the centered number of pumps that an individual gave a “treatment” balloon before they “cashed-out” or the balloon exploded. It is assumed that the higher the number of pumps becomes, the more risk is present to the individual. This value was then convolved with the double-gamma hemodynamic response function (HRF), which takes into account the physiological lag between stimulus and response, and smooths the stepwise function for the centered number of pumps. In this analysis, all subjects are similar in age, and so they are assumed to have the same HRF. Future work may be done to expand the model to account for variance in the HRF for different subjects. The HRF used the default values in the `canonicalHRF` function in the `neuRosim` package in *R*, which are described in the Simulated Data Analysis section. Finally, we deviated from Schonberg et al. (2012) by subtracting the centered, convolved number of pumps on the control balloon from the treatment series to provide a basis of comparison between the two balloon types. To summarize, the final covariate is calculated as

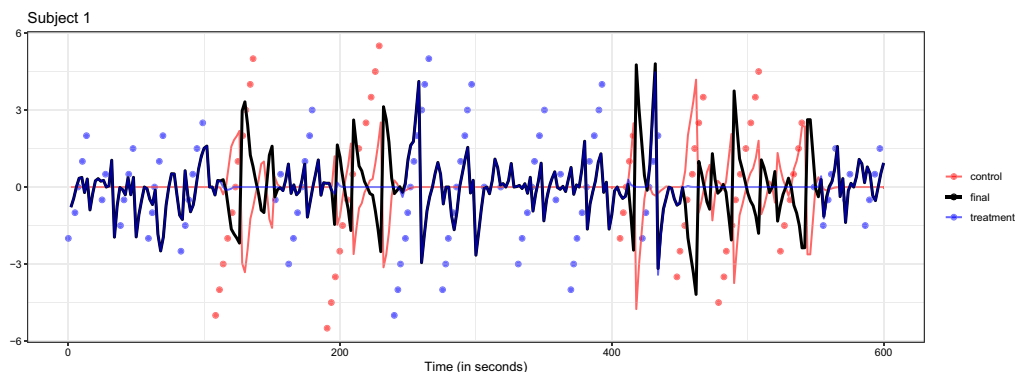


FIGURE 5.

The numeric values for the centered number of pumps for the control and treatment balloons, their convolutions with the double-gamma hemodynamic response function, and the final covariate used for these analyses

$$\begin{aligned} (\text{covariate}) &= (\text{centered, convolved number of treatment pumps}) \\ &- (\text{centered, convolved number of control pumps}). \end{aligned}$$

Figure 5 shows the raw values for the centered number of control and treatment pumps, as well as the convolved pump functions and the final values for the covariate that were used in these analyses.

The independent variable was then created as the difference between parametric modulation of the number of pumps on the treatment balloons and on the control balloons. The first two covariate values for each subject were removed to match the two dropped volumes in the EPI images. This covariate can be viewed as a contrast that accounts for the effect of the treatment balloons that mitigates any activation by subtracting the effect of increased pumps on the control balloon, when subjects are aware that there is no risk (see Fig. 5). This is similar to, but not the same as the analyses done in Schonberg et al. (2012), which attempts to measure activity by subtracting effect estimates associated with the control balloons from the effect estimates associated with the treatment balloons. In this scenario, positive coefficients imply more activity associated with treatment balloons, negative coefficients suggest higher activity levels for the control balloons, and values close to zero imply no activity or similar activity for the treatment and control balloons.

The previous work done by Schonberg et al. (2012) concludes that areas within the frontal lobe, insula, and occipital lobe show BOLD response associations with risk-associated tasks. In this analysis, the data will be analyzed in order to verify these conclusions and explore the functional connectivity between the defined regions of interest. We hypothesize that our model will recover activations in the frontal lobe and insula, and that there will be some positive partial correlations between the two groups.

The proposed Bayesian tensor mixed effect models were fitted first on whole-brain data with low ranks in order to identify regions of the brain that should be examined further in a full-resolution analysis of a slice within the scans. This is done in order to perform an analysis over the whole brain, which produces a dataset that does not fit into computer memory all-at-once at full resolution, without sacrificing precision on the estimate of voxel-level activation. Next, low-rank models were fitted to data from an axial slice in which $z = 18$ in the MNI standard space. This slice was identified as being within a region showing activity in the whole-brain analysis.

Due to the large size of the data, 2,200 samples were drawn from the joint posterior distribution of all of the parameters, and the first 200 samples were discarded as a “burn-in” measure. We

TABLE 3.

The median effective sample size and log deviance information criterion for the five tensor decomposition models and a vectorized model for comparison

	Slice ESS	DIC	Whole brain ESS	DIC
Rank 1	1,826	1107186130	1,743	1777033580
Rank 2	1,742	1107078983	1,583	1777227494
Rank 3	1,742	1107160202	1,211	1777227494
Rank 4	1,709	1107130129		
Rank 5	1,705	1107130129		
Spike-and-Slab	2,000		2,000	
Vectorized	1,774		2,000	

believe that this is a sufficient posterior sample size based on examinations of the log-likelihood and autocorrelation functions, which can be seen in the online resource Figs. 3 and 4. In addition, the `effectiveSize` function within the `coda` package in R is used to calculate median values for the effective sample size for the 2,000 posterior draws of the elements in all $\mathbf{B}_g$ for the different rank models, see Table 3. This table indicates uncorrelated post burn-in posterior samples to draw reliable posterior inference.

The final estimates of the activation tensors were found following the sequential 2-means variable selection method as described in Li and Pati (2017), using the median posterior standard deviations of the elements within each tensor $\mathbf{B}_g$ as the tuning parameter. These estimates within the whole brain were then reorganized to their original positions, and can be seen for a single axial slice in Fig. 6. Results for the high-resolution analysis of a single axial slice based on the whole-brain analysis can be seen in Fig. 7. Higher values of the coefficient suggest that there is an increase in the BOLD response associated with higher levels of perceived risk. Larger positive values suggest that blood flow increases in these regions as risk increases. Larger negative values would suggest regions that exhibit a decrease in blood flow as risk increases, though no such regions were observed in this analysis. The figures do show activations in the left posterior region of the frontal lobe and the anterior portion of the left insula, as concluded in Schonberg et al. (2012).

Similar to the simulation studies, the vectorized GDP, spike-and-slab, and GLM competitors are also fitted to the data to assess the advantages of preserving the tensor structure of the brain image in our proposed model. According to the deviance information criterion (DIC) (Gelman et al. 2014) given in Table 3, Rank 2 is the best performing tensor mixed effect model for the higher-resolution 2D slice data, while Rank 1 is the best performing model in the whole volume data. Figures 6 and 7 show that all the models that use tensor decompositions generally agree in terms of the posterior activation results. The vectorized model provides much lower estimates of activation strength than those obtained from the tensor decomposition models. These results suggest that the tensor mixed effects models using the tensor decomposition is more sensitive than the GLM or vectorized GDP models, while also being more specific in detecting non-activation in regions further from active regions than the spike-and-slab model.

The estimates for the significantly nonzero partial correlations between regions for the whole brain shown in Fig. 8 indicate a functional connectivity network that also supports our hypothesis that the frontal lobe and the insula are connected. Other rank models lead to results with similar connective networks. The whole-brain connectivity network is shown because the results are more interpretable. Our finding agrees with earlier experiments suggesting that the frontal lobe plays a

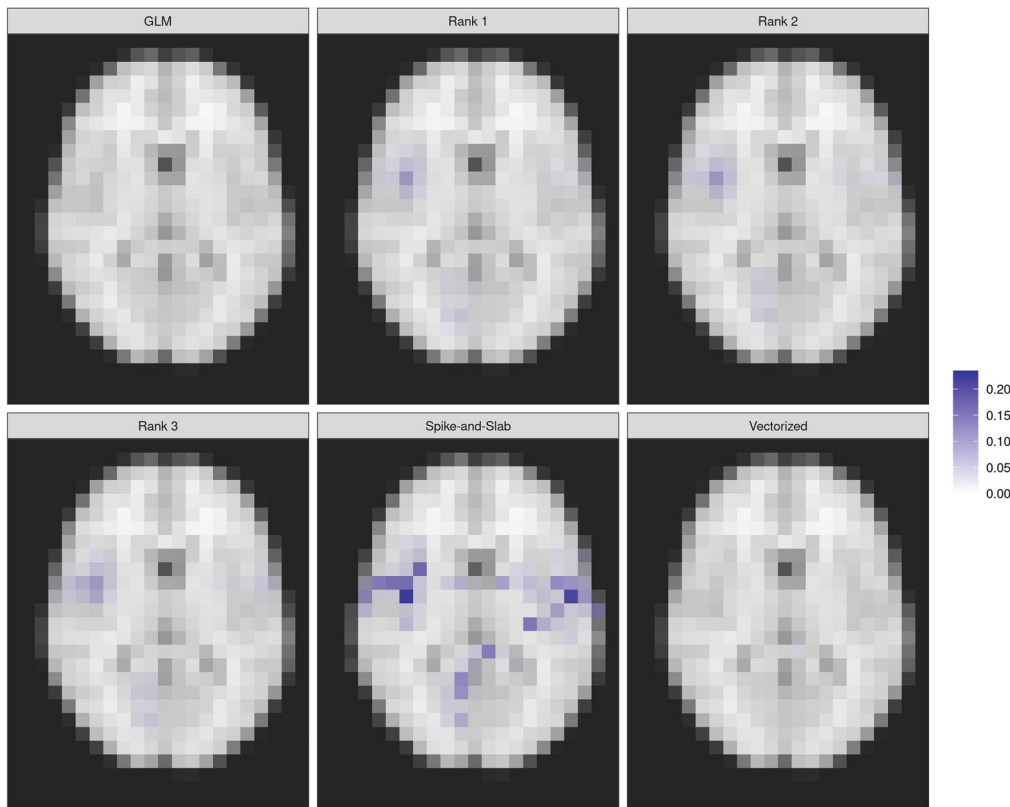


FIGURE 6.
One slice of activity estimates—whole brain

role in the assessment of risk (Miller and Milner 1985). The numeric estimates for these partial correlations are small, which is to be expected in high-noise smoothed data, especially given the strength of the regularization in the Gaussian graphical prior (note that this prior induces strong shrinkage, see, for example the strength of the regularization in Fig. 3 that is obtained in a much simpler simulation scenario with reasonable signal-to-noise data). However, these detected connectivity network in the regions of interest are significant and may be of investigative interest to neuroscientists.

6. Conclusion and Future Work

We present a new Bayesian tensor model for joint detection of voxel-level activation and region-specific connectivity in multi-subject studies. The proposed model produces markedly improved inference over a vectorized GDP model both in terms of identifying point estimation and quantifying uncertainty in a statistically principled manner. In addition, the model performs especially well in scenarios with low contrast-to-noise ratios, properly identifying hypercubic nonzero-valued regions within tensor coefficients while also finding functional connectivity between predefined regions. We also found that the proposed model exhibits an advantage over other Bayesian models in that nonzero estimates of activation tend to remain tightly clustered around true activation regions, preserving the underlying localized spatial structure. This is in

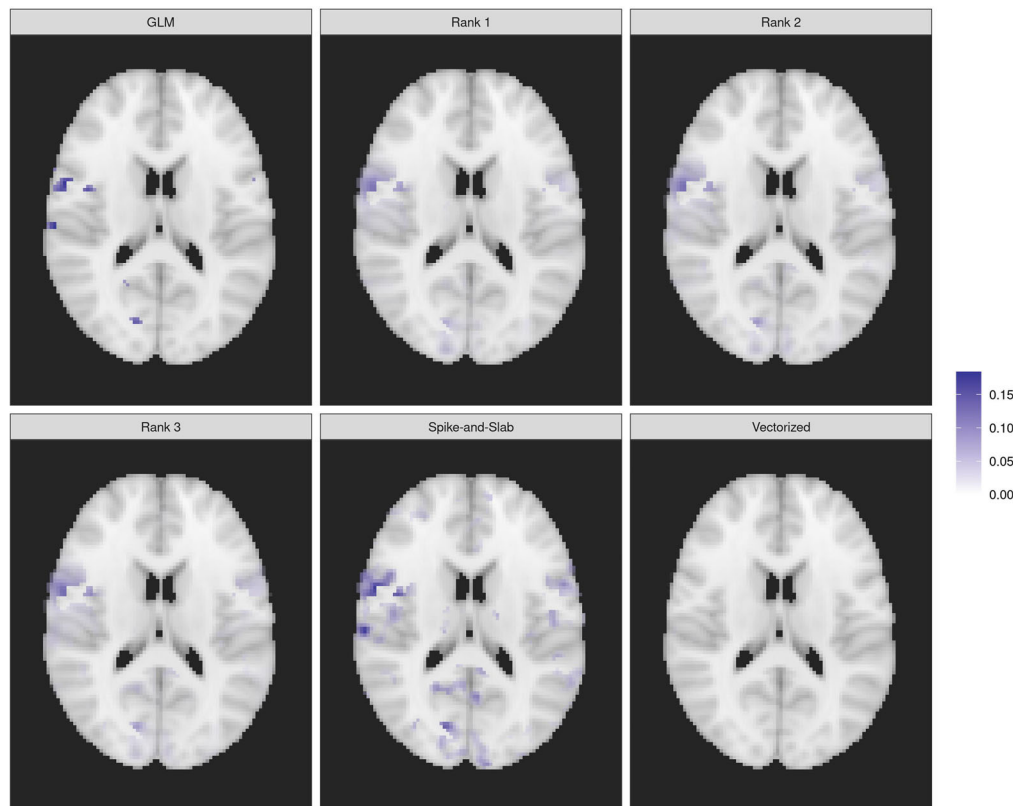


FIGURE 7.
One slice of activity estimates—single slice

contrast to both the spike-and-slab and vectorized GDP models considered as competitors. Our sensitivity analysis exhibited the robustness of the model to choices of hyperparameters. The proposed modeling structure is also flexible as it does not require extensive parameter tuning, adjustments for multiple testing, or selecting specific basis representations, making it accessible to a wide range of neuroscientists and statisticians alike. Analysis of the model's performance under misspecification showed that it does not appear to be strongly impacted by the presence of temporal correlated errors. Due to the shrinkage priors imposed on the activation and connectivity components of the model, effect sizes are mildly underestimated, but activation and connectivity are still detected in low contrast-to-noise and signal-to-noise settings. Furthermore, our simulation studies show that the proposed tensor models perform better than competing models, particularly in comparison to the GLM, in terms of inferring the connectivity structure across multiple regions.

Our analysis of a subset of the brain data examined by Schonberg et al. (2012) confirmed that increased risk was associated with activity in the insula and frontal lobe. The inference was further improved by examining functional connectivity between regions of interest to detect a functional connectivity network, which we hope can be further explored in future research.

Our proposed approach assumes several important extensions. Notably, the parsimony in activation coefficients achieved by a PARAFAC decomposition may appear to be restrictive in certain applications, and can be replaced by a more flexible Tucker decomposition. Additional shrinkage prior models may also be explored under different tensor decomposition prior structures to explore the effects on posterior inference. Extensions of (2) that incorporate nonlinear

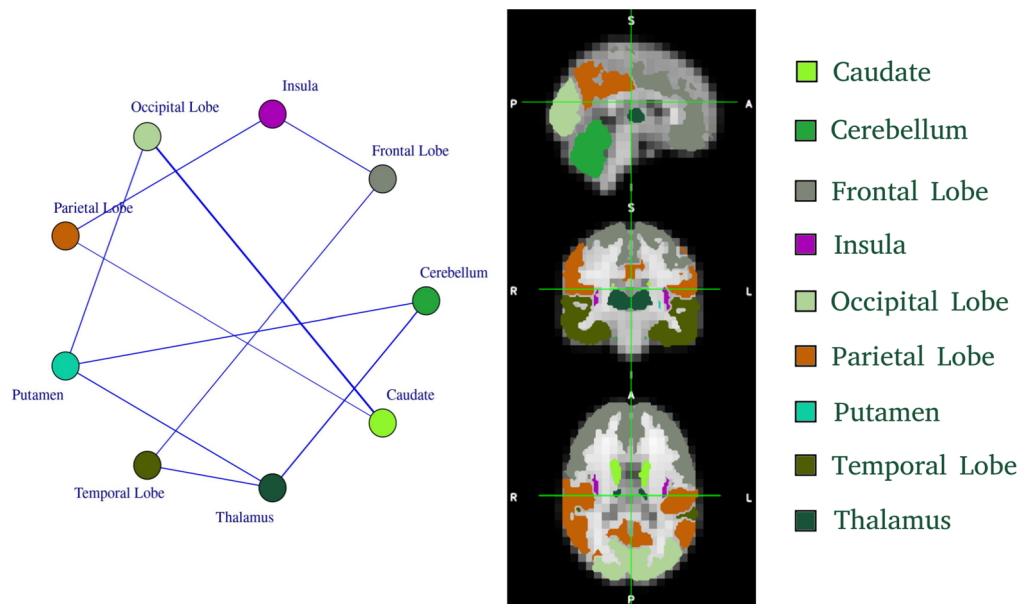


FIGURE 8.

The connected regions of the whole brain in the Rank 1 model, based on the partial correlation. The partial correlation here was found after using the sequential 2-means method (Li and Pati 2017) on the partial correlation matrix elements across all MCMC samples. Thicker lines correspond to larger partial correlations, as all of the estimates of the nonzero partial correlations are positive

regional effects through time will also be explored. Very importantly, extending the model to include temporal and spatiotemporal dependent error structures may further improve the model in its use with fMRI data. Finally, investigation into model-driven choices for subject-specific hemodynamic response functions may improve upon the accuracy of the proposed approach in real data applications.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

References

- Armagan, A., Dunson, D. B., & Lee, J. (2013). Generalized double Pareto shrinkage. *Statistica Sinica*, 23(1), 119.
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: A practical and powerful approach to multiple testing. *Journal of the Royal Statistical Society: Series B (Methodological)*, 57(1), 289–300.
- Bowman, F. D., Caffo, B., Bassett, S. S., & Kilts, C. (2008). A Bayesian hierarchical framework for spatial modeling of fMRI data. *Neuroimage*, 39(1), 146–156.
- Brown, P. J., Vannucci, M., & Fearn, T. (1998). Multivariate Bayesian variable selection and prediction. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 60(3), 627–641.
- Carvalho, C. M., Polson, N. G., & Scott, J. G. (2010). The horseshoe estimator for sparse signals. *Biometrika*, 97(2), 465–480.
- Collins, D. L., Holmes, C. J., Peters, T. M., & Evans, A. C. (1995). Automatic 3-D model-based neuroanatomical segmentation. *Human Brain Mapping*, 3(3), 190–208.
- Das, A., Sampson, A. L., Lainscsek, C., Muller, L., Lin, W., Doyle, J. C., et al. (2017). Interpretation of the precision matrix and its application in estimating sparse brain connectivity during sleep spindles from human electrocorticography recordings. *Neural Computation*, 29(3), 603–642.
- Eklund, A., Nichols, T. E., & Knutsson, H. (2016). Cluster failure: Why fMRI inferences for spatial extent have inflated false-positive rates. *Proceedings of the National Academy of Sciences*, 113(28), 7900–7905.
- Flandin, G., & Penny, W. D. (2007). Bayesian fMRI data analysis with sparse spatial basis function priors. *Neuroimage*, 34(3), 1108–1125.

- Friston, K. J., Ashburner, J., Frith, C. D., Poline, J.-B., Heather, J. D., & Frackowiak, R. S. J. (1995). Spatial registration and normalization of images. *Human Brain Mapping*, 3(3), 165–189.
- Friston, K. J., Fletcher, P., Josephs, O., Holmes, A., Rugg, M. D., & Turner, R. (1998). Event-related fMRI: Characterizing differential responses. *Neuroimage*, 7(1), 30–40.
- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A., & Rubin, D. B. (2014). *Bayesian data analysis* (Vol. 2). Boca Raton: CRC Press.
- George, E. I., & McCulloch, R. E. (1993). Variable selection via Gibbs sampling. *Journal of the American Statistical Association*, 88(423), 881–889.
- Guhaniyogi, R., & Spencer, D. (2018). *Bayesian tensor response regression with an application to brain activation studies*. Technical report, UCSC.
- Guhaniyogi, R., Qamar, S., & Dunson, D. B. (2017). Bayesian tensor regression. *Journal of Machine Learning Research*, 18(79), 1–31.
- Hutchison, R. M., Womelsdorf, T., Allen, E. A., Bandettini, P. A., Calhoun, V. D., Corbetta, M., et al. (2013). Dynamic functional connectivity: Promise, issues, and interpretations. *Neuroimage*, 80, 360–378.
- Ishwaran, H., Rao, J. S., et al. (2005). Spike and slab variable selection: Frequentist and Bayesian strategies. *The Annals of Statistics*, 33(2), 730–773.
- Jenkinson, M., Beckmann, C. F., Behrens, T. E. J., Woolrich, M. W., & Smith, S. M. (2012). Fsl. *Neuroimage*, 62(2), 782–790.
- Kalus, S., Sämann, P. G., & Fahrmeir, L. (2014). Classification of brain activation via spatial Bayesian variable selection in fMRI regression. *Advances in Data Analysis and Classification*, 8(1), 63–83.
- Kiers, H. A. L. (2000). Towards a standardized notation and terminology in multiway analysis. *Journal of Chemometrics: A Journal of the Chemometrics Society*, 14(3), 105–122.
- Kolda, T. G., & Bader, B. W. (2009). Tensor decompositions and applications. *SIAM Review*, 51(3), 455–500.
- Kook, J. H., Guindani, M., Zhang, L., & Vannucci, M. (2019). NPBates-fMRI: Non-parametric Bayesian general linear models for single- and multi-subject fMRI data. *Statistics in Biosciences*, 11(1), 3–21.
- Lazar, N. (2008). *The statistical analysis of functional MRI data*. New York: Springer.
- Lee, K.-J., Jones, G. L., Caffo, B. S., & Bassett, S. S. (2014). Spatial Bayesian variable selection models on functional magnetic resonance imaging time-series data. *Bayesian Analysis (Online)*, 9(3), 699.
- Li, H., & Pati, D. (2017). Variable selection using shrinkage priors. *Computational Statistics & Data Analysis*, 107, 107–119.
- Li, L., & Zhang, X. (2017). Parsimonious tensor response regression. *Journal of the American Statistical Association*, 112(519), 1131–1146.
- Mazziotta, J., Toga, A., Evans, A., Fox, P., Lancaster, J., Zilles, K., et al. (2001). A probabilistic atlas and reference system for the human brain: International consortium for brain mapping (ICBM). *Philosophical Transactions of the Royal Society of London. Series B: Biological Sciences*, 356(1412), 1293–1322.
- Miller, L., & Milner, B. (1985). Cognitive risk-taking after frontal or temporal lobectomy—ii. The synthesis of phonemic and semantic information. *Neuropsychologia*, 23(3), 371–379.
- Ng, B., Abugharbieh, R., Varoquaux, G., Poline, J. B., & Thirion, B. (2011). Connectivity-informed fMRI activation detection. In *International conference on medical image computing and computer-assisted intervention* (pp. 285–292). Springer.
- Olshausen, B. A., & Field, D. J. (2004). Sparse coding of sensory inputs. *Current Opinion in Neurobiology*, 14(4), 481–487.
- Patel, R. S., Bowman, F. D. B., & Rilling, J. K. (2006a). A Bayesian approach to determining connectivity of the human brain. *Human Brain Mapping*, 27(3), 267–276.
- Patel, R. S., Bowman, F. D. B., & Rilling, J. K. (2006b). Determining hierarchical functional networks from auditory stimuli fMRI. *Human Brain Mapping*, 27(5), 462–470.
- Penny, W. D., Friston, K. J., Ashburner, J. T., Kiebel, S. J., & Nichols, T. E. (2011). *Statistical parametric mapping: The analysis of functional brain images*. Boston: Elsevier.
- Sanyal, N., & Ferreira, M. A. R. (2012). Bayesian hierarchical multi-subject multiscale analysis of functional MRI data. *Neuroimage*, 63(3), 1519–1531.
- Schonberg, T., Fox, C. R., Mumford, J. A., Congdon, E., Trepel, C., & Poldrack, R. A. (2012). Decreasing ventromedial prefrontal cortex activity during sequential risk-taking: An fMRI investigation of the balloon analog risk task. *Frontiers in Neuroscience*, 6, 80.
- Scott, J. G., & Berger, J. O. (2010). Bayes and empirical-Bayes multiplicity adjustment in the variable-selection problem. *The Annals of Statistics*, 38(5), 2587–2619.
- Smith, M., & Fahrmeir, L. (2007). Spatial Bayesian variable selection with application to functional magnetic resonance imaging. *Journal of the American Statistical Association*, 102(478), 417–431.
- Wang, H., et al. (2012). Bayesian graphical LASSO models and efficient posterior computation. *Bayesian Analysis*, 7(4), 867–886.
- Wang, X., Nan, B., Zhu, J., & Koeppe, R. (2014). Regularized 3D functional regression for brain image data via Haar wavelets. *The Annals of Applied Statistics*, 8(2), 1045.
- Wang, Y., Kang, J., Kemmer, P. B., & Guo, Y. (2018). *DensParcorr: Dens-based method for partial correlation estimation in large scale brain networks*. Retrieved January 18, 2019 from <https://CRAN.R-project.org/package=DensParcorr>. R package version 1.1.
- Warnick, R., Guindani, M., Erhardt, E., Allen, E., Calhoun, V., & Vannucci, M. (2018). A Bayesian approach for estimating dynamic functional network connectivity in fMRI data. *Journal of the American Statistical Association*, 113(521), 134–151.

- Welvaert, M., Durnez, J., Moerkerke, B., Verdoolaege, G., & Rosseel, Y. (2011). neuRosim: An R package for generating fMRI data. *Journal of Statistical Software*, 44(10), 1–18.
- Welvaert, M., & Rosseel, Y. (2013). On the definition of signal-to-noise ratio and contrast-to-noise ratio for fMRI data. *PLoS one*, 8(11), e77089.
- Xu, L., Johnson, T. D., Nichols, T. E., & Nee, D. E. (2009). Modeling inter-subject variability in fMRI activation location: A Bayesian hierarchical spatial model. *Biometrics*, 65(4), 1041–1051.
- Yu, C.-H., Prado, R., Ombao, H., & Rowe, D. (2018). A Bayesian variable selection approach yields improved detection of brain activation from complex-valued fMRI. *Journal of the American Statistical Association*, 113(524), 1395–1410.
- Zhang, J., Li, X., Li, C., Lian, Z., Huang, X., Zhong, G., et al. (2014a). Inferring functional interaction and transition patterns via dynamic Bayesian variable partition models. *Human Brain Mapping*, 35(7), 3314–3331.
- Zhang, L., Guindani, M., & Vannucci, M. (2015). Bayesian models for functional magnetic resonance imaging data analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 7(1), 21–41.
- Zhang, L., Guindani, M., Versace, F., Engelmann, J. M., & Vannucci, M. (2016). A spatiotemporal nonparametric Bayesian model of multi-subject fMRI data. *The Annals of Applied Statistics*, 10(2), 638–666.
- Zhang, L., Guindani, M., Versace, F., & Vannucci, M. (2014b). A spatio-temporal nonparametric Bayesian variable selection model of fMRI data for clustering correlated time courses. *Neuroimage*, 95, 162–175.
- Zhou, H., Li, L., & Zhu, H. (2013). Tensor regression with applications in neuroimaging data analysis. *Journal of the American Statistical Association*, 108(502), 540–552.
- Zhu, H., Fan, J., & Kong, L. (2014). Spatially varying coefficient model for neuroimaging data with jump discontinuities. *Journal of the American Statistical Association*, 109(507), 1084–1098.

Manuscript Received: 19 AUG 2019

Final Version Received: 2 SEP 2020

Published Online Date: 19 SEP 2020