



Beyond consensus: Embracing heterogeneity in curated neuroimaging meta-analysis

Gia H. Ngo^{a,b}, Simon B. Eickhoff^{c,d}, Minh Nguyen^a, Gunes Sevinc^e, Peter T. Fox^{f,g},
R. Nathan Spreng^{h,i}, B.T. Thomas Yeo^{a,j,k,l,*}

^a Department of Electrical and Computer Engineering, Clinical Imaging Research Centre, N.1 Institute for Health and Memory Networks Program, National University of Singapore, Singapore

^b School of Electrical and Computer Engineering, Cornell University, Ithaca, NY, USA

^c Institute for Systems Neuroscience, Medical Faculty, Heinrich Heine University Düsseldorf, Düsseldorf, Germany

^d Institute of Neuroscience and Medicine, Brain & Behaviour (INM-7), Research Centre Jülich, Jülich, Germany

^e Department of Psychiatry, Massachusetts General Hospital, Harvard Medical School, Boston, MA, USA

^f Research Imaging Institute, University of Texas Health Science Center at San Antonio, San Antonio, TX, USA

^g South Texas Veterans Health Care System, San Antonio, TX, USA

^h Laboratory of Brain and Cognition, Montreal Neurological Institute, McGill University, Montreal, QC, Canada

ⁱ Departments of Psychiatry and Psychology, McGill University, Montreal, QC, Canada

^j Martinos Center for Biomedical Imaging, Massachusetts General Hospital, Charlestown, MA, USA

^k Centre for Cognitive Neuroscience, Duke-NUS Medical School, Singapore

^l NUS Graduate School for Integrated Sciences and Engineering, National University of Singapore, Singapore

ARTICLE INFO

Keywords:

Theory of mind
Autobiographical memory
Executive function
Inhibition
Attentional control
Mental disorder subtypes

ABSTRACT

Coordinate-based meta-analysis can provide important insights into mind-brain relationships. A popular approach for curated small-scale meta-analysis is activation likelihood estimation (ALE), which identifies brain regions consistently activated across a selected set of experiments, such as within a functional domain or mental disorder. ALE can also be utilized in meta-analytic co-activation modeling (MACM) to identify brain regions consistently co-activated with a seed region. Therefore, ALE aims to find consensus across experiments, treating heterogeneity across experiments as noise. However, heterogeneity within an ALE analysis of a functional domain might indicate the presence of functional sub-domains. Similarly, heterogeneity within a MACM analysis might indicate the involvement of a seed region in multiple co-activation patterns that are dependent on task contexts. Here, we demonstrate the use of the author-topic model to automatically determine if heterogeneities within ALE-type meta-analyses can be robustly explained by a small number of latent patterns. In the first application, the author-topic modeling of experiments involving self-generated thought ($N = 179$) revealed cognitive components fractionating the default network. In the second application, the author-topic model revealed that the left inferior frontal junction (IFJ) participated in multiple task-dependent co-activation patterns ($N = 323$). Furthermore, the author-topic model estimates compared favorably with spatial independent component analysis in both simulation and real data. Overall, the results suggest that the author-topic model is a flexible tool for exploring heterogeneity in ALE-type meta-analyses that might arise from functional sub-domains, mental disorder subtypes or task-dependent co-activation patterns. Code for this study is publicly available (https://github.com/ThomasYeoLab/CBIG/tree/master/stable_projects/meta-analysis/Ngo2019_AuthorTopic).

1. Introduction

Brain imaging experiments are often underpowered (Carp, 2012; Poline et al., 2012; Button et al., 2013). Coordinate-based meta-analysis provides an important framework for analyzing underpowered studies

across different experimental conditions and analysis pipelines to reveal reliable trends (Wager et al., 2003; Fox et al., 2014; Poldrack and Yarkoni, 2016). Large-scale coordinate-based meta-analyses synthesize thousands of experiments across diverse experimental designs to discover broad and general principles of brain organization and disorder (Laird

* Corresponding author. ECE, CIRC, N.1 Institute & MNP, National University of Singapore, Singapore.

E-mail address: thomas.yeo@nus.edu.sg (B.T.T. Yeo).

<https://doi.org/10.1016/j.neuroimage.2019.06.037>

Received 9 June 2017; Received in revised form 17 May 2019; Accepted 17 June 2019

Available online 20 June 2019

1053-8119/© 2019 Elsevier Inc. All rights reserved.

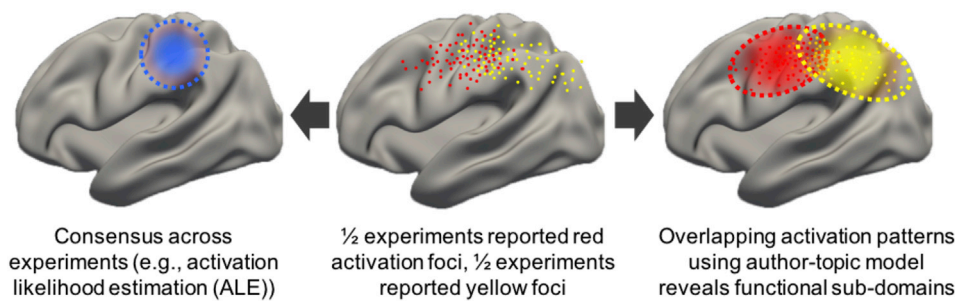


Fig. 1. Toy illustration of heterogeneity in neuroimaging meta-analysis. The middle panel shows activation peaks reported from neuroimaging experiments across multiple tasks within a functional domain. Half the experiments are red dots; half the experiments are yellow dots. The left panel illustrates a possible outcome of activation likelihood estimation (ALE), which converges on regions consistently activated across experiments (blue dotted circle). The right panel illustrates a possible estimate by the author-topic model (Yeo et al., 2015), which recovers overlapping patterns (red and yellow ovals) corresponding to two functional sub-domains. We note that the spatial spread of the activation foci was exaggerated to accentuate the overlaps and differences between the activation patterns of the two functional sub-domains.

et al., 2011; Poldrack et al., 2012; Crossley et al., 2014). By contrast, the vast majority of meta-analyses involve smaller number of experiments that are expertly chosen (curated) to generate consensus on specific functional domains (e.g., Binder et al., 2009), brain regions (e.g., Shackman et al., 2011) or disorders (e.g., Cortese et al., 2012).

A popular approach for smaller-scale meta-analyses is activation likelihood estimation or ALE (Turkeltaub et al., 2002; Laird et al., 2005; Eickhoff et al., 2009, 2012; Turkeltaub et al., 2012). ALE identifies brain regions consistently activated across neuroimaging experiments within a functional domain (Costafreda et al., 2008; Spaniol et al., 2009; Beissner et al., 2013) or within a disorder (e.g., Fitzgerald et al., 2008; Minzenberg et al., 2009; Di Martino et al., 2009). Thus, ALE treats heterogeneities across studies as noise. Consequently, ALE analysis might miss out on genuine biological heterogeneity indicative of functional sub-domains or disorder subtypes.

For example, Fig. 1 (middle panel) illustrates activation foci from experiments across multiple tasks associated with a hypothetical functional domain. These foci are generated by two latent sub-domains activating distinct, but overlapping, brain regions. Without prior knowledge of the two sub-domains from theory or previous empirical work, ALE will converge on regions commonly activated across both sub-domains (Fig. 1 left panel). To get around this issue, meta-analytic studies can sub-divide experiments into hypothetical functional sub-domains before applying ALE. For example, a recent meta-analysis divided working memory experiments into verbal versus non-verbal tasks, as well as tasks involving object identity versus object locations (Rottschy et al., 2012). However, manually subdividing experiments requires prior knowledge of the sub-domains and may reinforce biases towards existing concepts. By contrast, in this study, we explored whether a previously published data-driven approach (author-topic model; Yeo et al., 2015) can help uncover heterogeneities¹ within ALE-type meta-analyses in a bottom-up, data-driven fashion (Fig. 1 right panel).

1.1. Discovering sub-domains of self-generated thought

A good example in which ALE might miss out on functional sub-domains is the default network and self-generated thought (Smallwood, 2013; Andrews-Hanna et al., 2014). Self-generated thought involves associative and constructive processes that take place within an individual, and depends upon an internal representation to reconstruct or imagine a situation, understand a stimulus, or generate an answer to a

question. The term “self-generated thought” serves to contrast with thoughts where the primary referent is based on immediate perceptual input. By virtue of being largely stimulus independent or task unrelated, self-generated thought has been linked with the functions of the default network (Buckner et al., 2008; Andrews-Hanna et al., 2014). Previous ALE meta-analyses have implicated the default network in many tasks involving self-generated thought, including theory of mind, narrative fiction, autobiographical memory and moral cognition (Spreng et al., 2009; Binder et al., 2009; Mar, 2011; Sevinc and Spreng, 2014).

However, many studies have suggested that the default network might be fractionated into sub-systems. For example, Andrews-Hanna and colleagues have proposed a dorsomedial prefrontal subsystem preferentially specialized for social cognition and narrative processing (Andrews-Hanna et al., 2014; Spreng and Andrews-Hanna, 2015) and a medial temporal lobe sub-system preferentially specialized for mnemonic constructive processes (Andrews-Hanna et al., 2014; Christoff et al., 2016). Both sub-systems might spatially overlap or inter-digitate across multiple brain regions (Andrews-Hanna et al., 2014; Braga and Buckner, 2017), which would be challenging to ALE without assuming prior knowledge of the sub-systems (Fig. 1). Furthermore, specific default network fractionation details differed across studies (Laird et al., 2009a; Andrews-Hanna et al., 2010b; Mayer et al., 2010; Humphreys et al., 2015; Kernbach et al., 2018), so application of the author-topic model might potentially clarify sub-systems subserving self-generated thought.

1.2. Discovering multiple co-activation patterns of the left inferior frontal junction (IFJ)

Another common application of ALE is meta-analytic connectivity modeling (MACM), which identifies brain regions that consistently co-activate with a particular seed region (Toro et al., 2008; Koski and Paus, 2000; Robinson et al., 2010; Eickhoff et al., 2010). The assumption is that the seed region exhibits a *single* co-activation pattern regardless of the actual task activating the seed region (Robinson et al., 2010). However, studies have shown the existence of multiple hub regions in the brain (e.g., dorsal anterior insula, dorsal anterior cingulate cortex) that are activated across many different tasks and might adapt their connectivity pattern depending on task context (Cole et al., 2013; Uddin, 2015; Bertolero et al., 2017). Thus, a seed region might be involved in *multiple* task-dependent co-activation patterns (McIntosh, 2000).

A good example in which MACM might miss out on multiple co-activation patterns is the left inferior frontal junction (IFJ; Muhle-Karbe et al., 2015). The IFJ has been implicated in many cognitive processes (Brass et al., 2005; Chikazoe et al., 2009a, b; Asplund et al., 2010) and is a key node of the multiple-demand system (Duncan, 2010; Fedorenko et al., 2013). IFJ might also coordinate information among modules by adapting its connectivity patterns across different resting and task states (Cole et al., 2013; Bertolero et al., 2018). Therefore, one might

¹ We note that when estimating functional sub-domains, we are not interested in capturing idiosyncrasies of individual experiments or even individual tasks. Instead, we are hoping to estimate a small number of overlapping, but distinct activation patterns (cognitive processes) that are recruited to different extents across tasks.

expect the IFJ region to exhibit multiple co-activation patterns that are dependent on task contexts. Since ALE cannot capture heterogeneity across experiments, MACM might be insensitive to such task-dependent co-activation patterns. On the other hand, application of the author-topic model to the IFJ region might yield multiple meaningful co-activation patterns.

1.3. Author-topic model

In this work, we propose the use of the author-topic model to automatically make sense of heterogeneity within ALE-type meta-analyses. We have previously utilized the author-topic model (Fig. 2; Yeo et al., 2015; Bertolero et al., 2015) to encode the intuitive notion that a behavioral task recruits multiple cognitive components, which are in turn supported by overlapping brain regions (Poldrack, 2006; Leech et al., 2012; Barrett and Satpute, 2013). While our previous work focused on large-scale meta-analysis across many functional domains (Yeo et al., 2015; Bertolero et al., 2015), the current study focuses on heterogeneity within a functional domain (self-generated thought) or co-activation heterogeneity of a seed region (left IFJ). These applications of the author-topic model are made possible by the development of a novel inference algorithm for the author-topic model (Ngo et al., 2016) that is sufficiently robust for smaller-scale meta-analyses.

Our choice of self-generated thought is motivated by previous work suggesting the possibility of fractionating self-generated thought into functional sub-domains (Section 1.1). Similarly, our choice of left IFJ is motivated by previous work suggesting that IFJ might adaptively modify its connectivity patterns across task contexts (Section 1.2). There are of course other functional domains (e.g., executive function) that might be fractionated and other hub regions (e.g., dorsal anterior insula) that might exhibit task-dependent co-activation patterns. Therefore, we have made our code publicly available for researchers to explore the heterogeneity of their preferred functional domain, hub region or mental disorder.

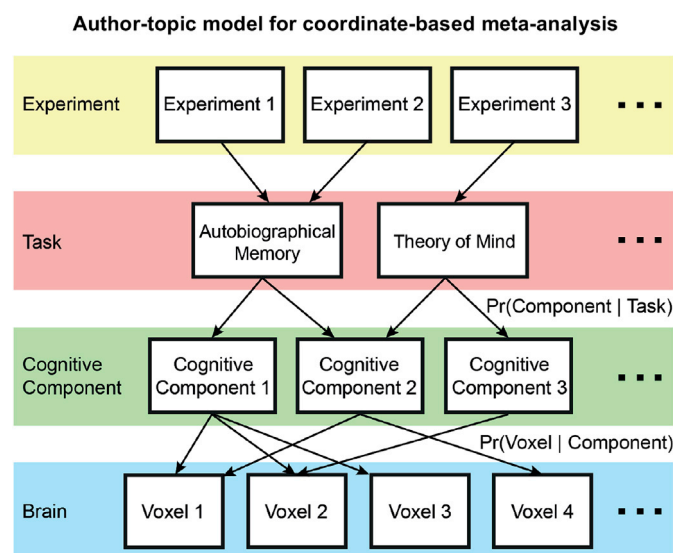


Fig. 2. Author-topic model for coordinate-based meta-analysis (Yeo et al., 2015). The underlying premise of the model is that behavioral tasks recruit multiple cognitive components, which are in turn supported by overlapping brain regions. The model parameters are the probability that a task would recruit a cognitive component ($\Pr(\text{component} | \text{task})$) and the probability that a component would activate a brain voxel ($\Pr(\text{voxel} | \text{component})$). The author-topic model can be directly applied to estimate cognitive components (sub-systems) of self-generated thought.

2. Methods

2.1. Overview

In Section 2.2, we reviewed the author-topic model and how it could be applied to coordinate-based meta-analysis (Yeo et al., 2015). Section 2.3 discussed simulations and comparisons with spatial independent component analysis. Finally, the model was utilized in two different applications. In the first application (Section 2.4), we applied the author-topic model to discover cognitive components subserving self-generated thought. In the second application (Section 2.5), we estimated the co-activation patterns of the left IFJ.

2.2. Author-topic model

2.2.1. Intuition behind the model

The author-topic model was originally developed to discover topics from a corpus of text documents (Rosen-Zvi et al., 2010). The model represents each text document as an unordered collection of words written by a group of authors. Each author is associated with a probability distribution over topics, and each topic is associated with a probability distribution over a dictionary of words. Given a corpus of text documents, there are algorithms to estimate the distribution of topics associated with each author and the distribution of words associated with each topic. A topic is in some sense abstract, but is made concrete by its association with certain words and its association with certain authors. For example, if the author-topic model was applied to neuroimaging research articles, the algorithm might yield a topic associated with the author “Stephen Smith” and words like “fMRI”, “resting-state” and “ICA”. One might then interpret the topic posthoc as a “resting-state fMRI” research topic.

In a previous study (Yeo et al., 2015), the author-topic model was applied to neuroimaging meta-analysis (Fig. 2) by treating task contrasts in the BrainMap database (Fox and Lancaster, 2002) as text documents, 83 BrainMap task categories (e.g., n-back) as authors, cognitive components as topics, and activation foci as words in the documents. Thus, the model encodes the premise that different behavioral tasks recruit multiple cognitive components, supported by overlapping brain regions.

Suppose a study utilizes one or more task categories, resulting in an experimental contrast yielding a collection of activation foci. Under the author-topic model, each activation focus is assumed to be generated by first randomly selecting a task from the set of tasks utilized in the experiment. Given the task, a component is randomly chosen based on the probability of a task recruiting a component ($\Pr(\text{component} | \text{task})$). Given the component, the location of the activation focus is then randomly chosen based on the probability that the component would activate a voxel ($\Pr(\text{voxel} | \text{component})$). The entire collections of $\Pr(\text{component} | \text{task})$ and $\Pr(\text{voxel} | \text{component})$ are denoted as matrices θ and β , respectively. For example, the 2nd row and 3rd column of θ corresponds to $\Pr(\text{3rd component} | \text{2nd task})$ and the 4th row and 28th column of β corresponds to $\Pr(\text{28th voxel} | \text{4th component})$. Therefore, each row of θ and β sums to 1. The formal mathematical definition of the model is provided in Supplemental Method S1.

A key property of the author-topic model is that the ordering of words within a document is exchangeable. When applied to meta-analysis, the corresponding assumption is that the ordering of activation foci is arbitrary. Although the ordering of words within a document is obviously important, the ordering of activation foci is not. For example, in the context of text documents, “dog has a bone” has a different meaning from “bone has a dog”. On the other hand, in the context of a fMRI experiment, reporting parietal activation coordinates followed by prefrontal activation coordinates is equivalent to reporting prefrontal activation coordinates followed by parietal activation coordinates. Therefore, the author-topic model is arguably more suitable for meta-analysis than topic discovery from documents.

2.2.2. Estimating the model parameters

Given a collection of experiments with their associated activation coordinates and task categories, as well as the number of cognitive components K , the probabilities θ and β can be estimated using various algorithms (Rosen-Zvi et al., 2010; Yeo et al., 2015; Ngo et al., 2016). Here, we chose to utilize the CVB algorithm because the algorithm was more robust to the choice of hyperparameters in smaller datasets. Although the CVB algorithm for the author-topic model was first introduced in a conference article (Ngo et al., 2016), detailed derivations have not been published. For completeness, detailed derivations of the author-topic CVB algorithm are provided in Supplemental Method S2. Explanations of why the CVB algorithm is theoretically better than the EM algorithm and standard variational Bayes inference are found in Supplemental Method S3. In this work, Bayesian information criterion (BIC) was used to estimate the optimal number of cognitive components (Supplemental Method S4). Further implementation details are found in Supplemental Method S5.

2.2.3. Input to the author-topic model

Each task activation contrast was associated with a set of activation foci. The spatial locations (i.e., coordinates) of the activation foci were reported in or transformed to the MNI152 coordinate system (Lancaster et al., 2007). Using standard meta-analysis procedure (Wager et al., 2009; Yarkoni et al., 2011; Yeo et al., 2015), a 2 mm-resolution binary activation image was created for each experimental contrast, in which a voxel was given a value of 1 if it was within a 10 mm-radius of any activation focus, and 0 otherwise. Thus, the set activated voxels of each experiment in the author-topic model corresponds to the set of voxels with a value of 1 in the corresponding 2 mm-resolution binary activation image. We note that the exact choice of smoothing radius did not significantly affect the results (see Section 3.4).

2.3. Simulations

2.3.1. Independent component analysis (ICA)

ICA is a data-driven technique that has been widely applied to fMRI (Calhoun et al., 2001; Beckmann and Smith, 2004). ICA has also been successfully applied to coordinate-based meta-analysis (Smith et al., 2009). However, the author-topic model has a few significant advantages over ICA in the case of coordinate-based meta-analysis. First, activation foci are binary data in the sense that a voxel is either reported to be activated or not in an experiment. However, ICA requires positive and negative values in the input data, which involves demeaning the binary values at each voxel (across experiments). In contrast, the author-topic model makes direct use of the binary activation data. Second, the author-topic model is able to exploit task categorical information (red task layer in Fig. 2), which is non-trivial to introduce in ICA.

Most importantly, ICA estimates can be negative, which do not make

sense in the case of coordinate-based meta-analysis. For example, a task should not be allowed to be negatively associated with a component, since task activation and de-activation in a coordinate-based meta-analysis are typically handled separately. Similarly, it does not make sense for the activation maps associated with each component to be negative. The situation is of course reversed in image-based meta-analysis (Salimi-Khorshidi et al., 2009), where there might be both activation and de-activation. For image-based meta-analysis, it does make sense to talk about components being negatively recruited by a task and ICA makes more theoretical sense than the author-topic model.

2.3.2. Simulation details

Here, we considered simulations to compare the effectiveness of the author-topic model and ICA. More specifically, we considered a hypothetical situation in which five tasks from a functional domain recruited two cognitive components with different probabilities (Fig. 3A). The two components have distinct activation patterns on a 2D “brain” of 256 by 256 pixels. More specifically, each component is associated with activations within two Gaussian distributions centered at two opposite quadrants of the 2D brain (Fig. 3B). Given the activation foci of multiple experiments (task contrasts), the goal was to automatically recover the two cognitive components using either the author-topic model or ICA.

A single simulation run comprised 150 experiments (task contrasts), which is comparable to a typical meta-analysis (c.f. self-generated thought in Section 2.4). Each experiment (task contrast) was randomly assigned to one of the five tasks, with the contrast distributions skewed towards two of the five tasks to simulate the fact that some tasks are more popular than others in the literature. Furthermore, each task contrast is randomly chosen to have between 1 and 10 activation foci. For each activation focus, a component was randomly sampled based on the probability of components given the task assigned to the experiment. For the given component, one of the 2-D Gaussian distributions of each component was randomly chosen with equal probabilities (Fig. 3B). The spatial location of the activation focus was then randomly sampled from the Gaussian distribution. The activation focus was smoothed with a binary smoothing kernel, such that all pixels within 10 voxels from an activation focus were given a value of 1, and 0 otherwise.

For a given simulation run, the latent components were estimated using either ICA or the author-topic model. We considered three possible ICA setups. The first two setups (ICA1 and ICA2) utilized CanICA (Varoquaux et al., 2010), an ICA decomposition implementation provided with Nilearn (Abraham et al., 2014). CanICA extracts representative patterns of multi-subject fMRI data by performing ICA on a data subspace common to the group (Varoquaux et al., 2010). In the two setups ICA1 and ICA2, each task was treated as a subject. In ICA1, the activation maps of all experiments assigned to the same task were summed together, i.e., each task was treated as a single subject with a single time point. In ICA2, each task was treated as a single subject, but the experiments assigned to

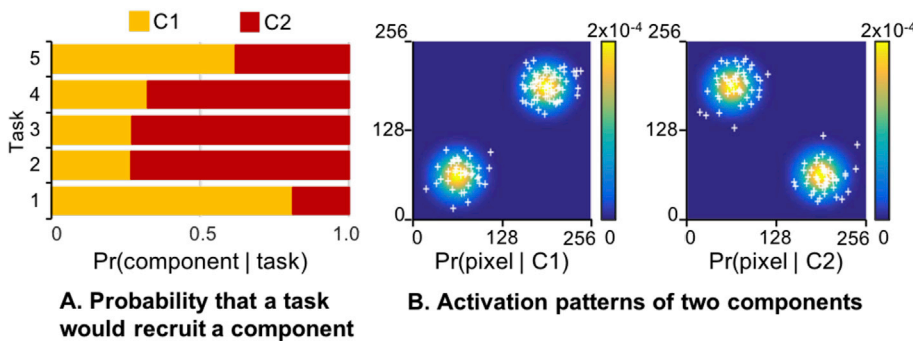


Fig. 3. Simulation of heterogeneity in coordinate-based meta-analysis. (A) Bar chart shows five tasks from a functional domain recruiting two cognitive components with different probabilities. (B) Activation patterns of two components on a 2D “brain” of 256 by 256 pixels. Each component is associated with activations (white crosses) within two Gaussian distributions centered at two opposite quadrants of the 2D “brain”. For each simulation run, the probability of a task recruiting a component and the covariances of each component’s 2D Gaussian distributions were randomly generated. The author-topic model and ICA were then applied to recover the two components. We note that ICA mixture weights can be negative, which does not make sense in the context of coordinate-based meta-analysis. As such, we discarded simulation runs if any of the ICA estimates yielded negative weights.

the given task were treated as separate time points of the subject. The third setup (ICA3) utilized the MELODIC implementation of ICA from the FSL package (Beckmann and Smith, 2004; Smith et al., 2004).

To evaluate the estimation quality, Pearson's correlation coefficient was computed between the groundtruth probability distribution of a component activating a vertex ($\text{Pr}(\text{vertex} | \text{component})$) against the estimates from the author-topic model or ICA. Pearson's correlation coefficient was also computed between the groundtruth distribution of components given a task ($\text{Pr}(\text{component} | \text{task})$) and estimates from the author-topic model or ICA.

The simulation was repeated multiple times. For a given simulation run, the covariances of each component's 2D Gaussian distributions were randomly generated (Fig. 3B). The probability of a task recruiting a component was also randomly generated (Fig. 3A). As explained previously (Section 2.3.1), ICA's mixture weights can be negative, which implies negative associations between tasks and components. This does not make sense in the case of coordinate-based meta-analysis, so we discarded simulation runs if any of the ICA estimates yielded negative weights. Overall, we ran roughly 300 simulation runs in order to yield exactly 100 simulation runs, in which ICA estimates were valid.

2.4. Self-generated thought

2.4.1. Activation foci of experiments involving self-generated thought

To explore cognitive components subserving self-generated thought, we considered 1812 activation foci from 179 experimental contrasts across 167 imaging studies, each employing one of seven task categories subjected to prior meta-analysis with GingerALE (Fox and Lancaster, 2002; Laird et al., 2009b, 2011; Fox et al., 2014; <http://brainmap.org/ale>). Of the 167 studies, 48 studies employed “Autobiographical Memory” ($N = 19$), “Navigation” ($N = 13$) or “Task Deactivation” ($N = 16$) tasks. The 48 studies were employed in a previous meta-analysis examining the default network (Spreng et al., 2009). There were 79 studies involving “Story-based Theory of Mind” ($N = 18$), “Nonstory-based Theory of Mind” ($N = 42$) and “Narrative Comprehension” ($N = 19$) tasks. The 79 studies were utilized in a previous meta-analysis examining social cognition and story comprehension (Mar, 2011). Finally, there were 40 studies involving the “Moral Cognition” task that was again utilized in a previous meta-analysis (Sevinc and Spreng, 2014). The list of all experiments included in the dataset are provided in Supplemental Method S7. The criteria for selecting the experiments can be found in the original meta-analyses (Spreng et al., 2009; Mar, 2011; Sevinc and Spreng, 2014). All foci coordinates were in or transformed to the MNI152 coordinate system (Lancaster et al., 2007).

2.4.2. Discovering cognitive components of self-generated thought

The application of the author-topic model to discover cognitive components subserving self-generated thought (Fig. 2) is conceptually similar to the original application to the BrainMap (Yeo et al., 2015). The key difference is that the current application is restricted to seven related tasks in order to discover heterogeneity within a single functional domain, while the original application sought to find common and distinct cognitive components across domains.

The model parameters are the probability of a task recruiting a component ($\text{Pr}(\text{component} | \text{task})$) and the probability of a component activating a brain voxel ($\text{Pr}(\text{voxel} | \text{component})$). The parameters were estimated from the 1812 activation foci from the previous section using the CVB algorithm (Supplemental Methods S2 and S5). BIC was used to estimate the optimal number of cognitive components (Supplemental Method S4).

2.4.3. Interpreting cognitive components of self-generated thought

We note that $\text{Pr}(\text{component} | \text{task})$ can be interpreted as follows. Suppose $\text{Pr}(\text{component C1} | \text{autobiographical memory})$ is equal to 0.63 and an autobiographical memory experiment reports 100 activation foci. Then, on average, 63 of the 100 foci will fall inside component C1.

The matrix $\text{Pr}(\text{voxel} | \text{component})$, β , can be interpreted as K brain images in MNI152 coordinate system (Lancaster et al., 2007), where K is the number of cognitive components. Volumetric slices highlighting specific subcortical structures were displayed using FreeSurfer (Fischl, 2012). The cerebral cortex was visualized by transforming the volumetric images from MNI152 space to fs_LR surface space using Connetome Workbench (Van Essen et al., 2013) via the FreeSurfer surface space (Buckner et al., 2011; Fischl, 2012). For visualization purpose, isolated surface clusters with less than 20 vertices were removed. $\text{Pr}(\text{component} | \text{task})$ was thresholded at $1/K$, and $\text{Pr}(\text{voxel} | \text{component})$ was thresholded at $1e-5$, consistent with previous work (Yeo et al., 2015). Unthresholded maps of the components are available on NeuroVault (Gorgolewski et al., 2015) at <https://neurovault.org/collections/4684/>.

2.4.4. Goodness of fit

For each task, we computed the weighted average of the cognitive components ($\text{Pr}(\text{voxel} | \text{component})$), where the weights corresponded to the probabilities of the task recruiting the components ($\text{Pr}(\text{component} | \text{task})$). This weighted average spatial map could be interpreted as the model estimate of the “ideal” (reconstructed) activation map for a particular task. The model fit was good if a task's reconstructed activation map was similar to the empirical activation map of the task (obtained by averaging the activation maps of all experiments employing the task). Therefore, we computed Pearson's correlation coefficient between all pairs of reconstructed and empirical activation maps, yielding a 7×7 correlation matrix (since there were 7 tasks).

2.4.5. Correspondence between cognitive components and resting-state networks

Motivated by similarities between task and resting-state networks (Smith et al., 2009; Laird et al., 2011; Yeo et al., 2015), we compared the cognitive components of self-generated thought with a previously published set of 17 resting-state networks (Yeo et al., 2011). For each resting-state network and each cognitive component, the probability of the cognitive component activating a voxel ($\text{Pr}(\text{voxel} | \text{component})$) was averaged across all voxels within the network, resulting in an average probability of a component activating the given network.

2.5. Left inferior frontal junction (IFJ)

2.5.1. Activation foci of experiments activating the left IFJ

To explore task-dependent co-activation patterns expressed by the left IFJ, we considered activation foci from experiments reporting activation within a left IFJ seed region (Fig. S1) delineated by a previous study (Muhle-Karbe et al., 2015). Muhle-Karbe and colleagues performed a co-activation-based parcellation of a left lateral prefrontal region into six parcels, including an IFJ region (Muhle-Karbe et al., 2015). The parcellation procedure assumed that voxels within a parcel exhibited a single co-activation pattern. Thus, the advantage of using this particular IFJ seed region (instead of an IFJ region from a different study) is that this region is thought to exhibit a single co-activation pattern (according to MACM).

This seed region is publicly available on ANIMA (Reid et al., 2016a; http://anima.fz-juelich.de/studies/MuhleKarbe_2015_IFJ). We selected experiments from the BrainMap database with at least one activation focus falling within the IFJ seed region. We further restricted our analyses to experimental contrasts involving normal subjects. Overall, there were 323 experiment contrasts from 238 studies with a total of 5201 activation foci. The list of all experiments included in the dataset are provided in Supplemental Method S8.

2.5.2. Discovering co-activation patterns of the IFJ

To apply the author-topic model to discover co-activation patterns, we consider each of the 323 experimental contrasts to employ its own unique task category (Fig. 4). In the parlance of the author-topic model, we assumed each document (experiment) has its own unique author

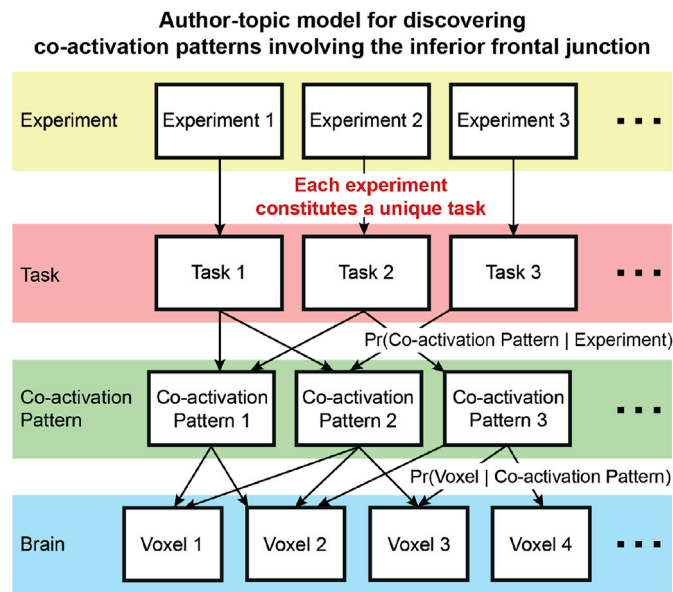


Fig. 4. Author-topic model for discovering co-activation patterns of the inferior frontal junction (IFJ). In contrast to Fig. 2, this instantiation of the model assumes that each experiment constitutes a unique task. The premise of the model is that the IFJ expresses multiple overlapping task-dependent co-activation patterns. The model parameters are the probability of an experiment recruiting a co-activation pattern ($\Pr(\text{co-activation pattern} | \text{experiment})$), and the probability of a voxel being associated with a co-activation pattern ($\Pr(\text{voxel} | \text{co-activation pattern})$).

(task). The premise of the model is that the IFJ expresses one or more overlapping co-activation patterns depending on task contexts. A single experiment activating the IFJ might recruit one or more co-activation patterns. The model parameters are the probability that an experiment would recruit a co-activation pattern ($\Pr(\text{co-activation pattern} | \text{experiment})$), and the probability that a voxel would be involved in a co-activation pattern ($\Pr(\text{voxel} | \text{co-activation pattern})$). The parameters were estimated from the 5201 activation foci from the previous section using the CVB algorithm (Supplemental Method S2 and S5). BIC was used to estimate the optimal number of co-activation patterns (Supplemental Method S4).

2.5.3. Interpreting co-activation patterns of the IFJ

Similar to the previous application on self-generated thought, the matrix $\Pr(\text{voxel} | \text{co-activation pattern})$, β , was visualized as K brain images in both fs_LR surface space and MNI152 volumetric space. Like before, isolated surface clusters with less than 20 vertices were removed for the purpose of visualization. Unthresholded spatial maps of the co-activation patterns are available on NeuroVault (Gorgolewski et al., 2015) at <https://neurovault.org/collections/4718/>.

Because each of the 323 experiments was treated as employing a unique task category, $\Pr(\text{co-activation pattern} | \text{experiment})$, θ , is a matrix of size $K \times 323$. θ was further mapped onto BrainMap task categories to assist in the interpretation. More specifically, since the experiments were extracted from the BrainMap database, each experiment was tagged with one or more BrainMap task categories (Table S1). The $\Pr(\text{co-activation pattern} | \text{experiment})$ was averaged across experiments employing the same task category to estimate the probability that a task category would recruit a co-activation pattern ($\Pr(\text{co-activation pattern } c | \text{task } t)$). Further details of this procedure are found in Supplemental Method S6. The $\Pr(\text{co-activation pattern} | \text{task})$ can be interpreted as follows. Suppose $\Pr(\text{co-activation pattern } C1 | \text{semantic monitoring/discrimination})$ is equal to 0.51 and we have a semantic monitoring/discrimination experiment that reports activation in the left IFJ and 100 activation foci. Then, on average, 51 foci will fall inside co-activation

pattern C1.

We note that directly using the BrainMap task categories to interpret the co-activation patterns is tricky. This is because a BrainMap task category might only have a very small percentage of experiments activating the IFJ, so these experiments might not be representative of the task category. For example, of the 230 experiments in the BrainMap database labeled as the “Encoding” task category, only 13 experiments reported activations in the left IFJ. Thus, the 13 experiments were not simply encoding tasks, but encoding tasks that happened to activate the IFJ. This is the reason why the BrainMap task categories were not directly utilized in the author-topic model for the IFJ analysis and that each experiment was treated as employing a unique task category (c.f. self-generated thought in Section 2.4).

To ensure an appropriate interpretation, we inspected the original publications associated with the top three experiments with the highest $\Pr(\text{co-activation pattern} | \text{experiment})$ for each of the top three tasks associated with each co-activation pattern, i.e., nine publications for each co-activation pattern. The literature analysis allowed us to determine if there were common neural processes underlying the subset of experiments within each task category that strongly activated the IFJ.

2.6. Goodness of fit

For each co-activation pattern, activation maps of the top three experiments with the highest probability of recruiting a co-activation pattern (i.e., $\Pr(\text{co-activation pattern} | \text{experiment})$) for each of the top three tasks associated with the co-activation pattern (i.e., nine activation maps in total) were averaged, resulting in an empirical activation map associated with each co-activation pattern. The model fit was good if the empirical activation map was similar to the estimated co-activated pattern. Therefore, we computed Pearson's correlation coefficient between all pairs of empirical activation maps and co-activation maps, yielding a $K \times K$ correlation matrix, where K is the number of co-activation patterns estimated by BIC.

2.7. Data and code availability

Activation foci from the meta-analysis of self-generated thought and the source code of the author-topic model, including the visualization and analysis tools, are publicly available at https://github.com/ThomasYeoLab/CBIG/tree/master/stable_projects/meta-analysis/Ngo2019_AuthorTopic. The activation foci from the meta-analysis of IFJ can be obtained via a collaborative-use license agreement with BrainMap (<http://www.brainmap.org/collaborations.html>).

3. Results

3.1. Overview

In Section 3.2, we show simulation results suggesting that the author-topic model compares favorably with ICA in the goal of discovering latent patterns in coordinate-based meta-analysis. We then explored the cognitive components of self-generated thought (Section 3.3) and the co-activation patterns of the IFJ (Section 3.4). Finally, Section 3.5 discusses a few control analyses.

3.2. Simulations

Fig. 5 shows the results of one representative simulation (see Section 2.3 for details). Fig. 5A shows the groundtruth 2D “brain” maps for this representative simulation run. The two leftmost columns show simulated activation foci as white crosses overlaid on top of the 2D Gaussian distributions used to generate the foci. The rightmost bar chart shows the probability of each of the 5 tasks recruiting a component.

The rightmost column of Fig. 5B shows the author-topic model estimates of the probability of each of the 5 tasks recruiting a component.

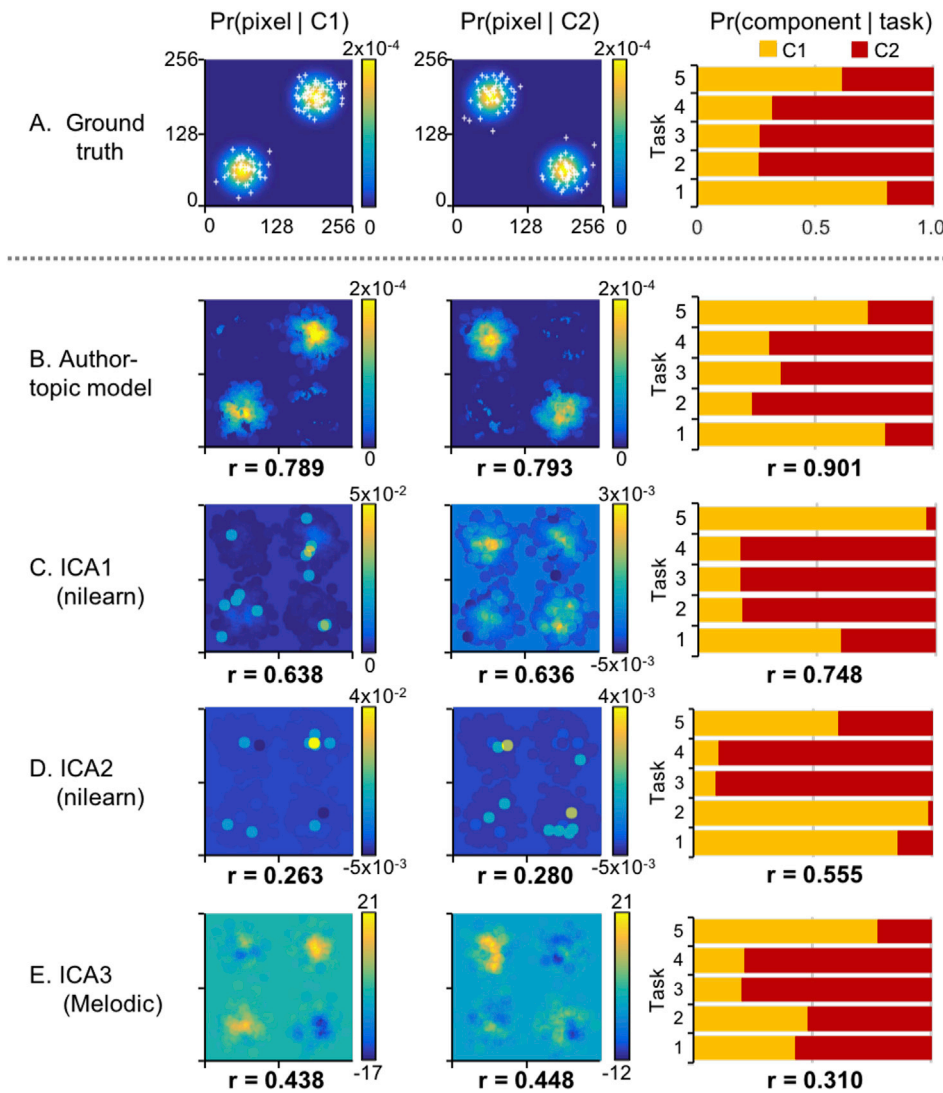


Fig. 5. Simulation comparing the author-topic model and ICA. (A) Single representative simulation run. Two leftmost columns show activation foci (white crosses) on top of Gaussian distributions used to generate the foci. Rightmost bar chart shows the probability of each of the 5 tasks recruiting a component. (B) Author-topic model estimates. (C-E) ICA estimates. Number below each panel is the correlation between model estimates and groundtruth averaged across 100 simulation runs. Observe that ICA can yield negative weights, which do not make sense in the context of a coordinate-based meta-analysis (see discussion in Section 2.3.1). We note that about 300 simulation runs were performed in order to generate 100 simulation runs in which ICA estimates of mixture weights were non-negative.

The rightmost column of Fig. 5C to E shows the ICA mixture weights, normalized so they sum to one.² The mixture weights represent the association between the tasks and the components. The numbers at the bottom of each panel are the correlations between the estimates and groundtruth averaged across 100 simulation runs. In general, the author-topic model yielded better estimates of the associations between tasks and components.

Fig. 5B shows the author-topic model estimates, while Fig. 5C to E shows the ICA estimates. The two leftmost columns show the spatial maps of the two components estimated by the author-topic model or ICA. The numbers at the bottom of each panel are the correlations between the estimated and groundtruth “brain” maps averaged across 100 simulation runs. In general, the author-topic model yielded better estimates of the groundtruth “brain” maps. It is also worth noting that the ICA spatial maps showed negative values, even though the simulation runs had been constrained to those where ICA mixture weights were positive.³ As previously explained (Section 2.3.1), negative values are not meaningful in the context of coordinate-based meta-analysis.

3.3. Self-generated thought

3.3.1. ALE meta-analysis of self-generated thought

Fig. 6 shows the activation likelihood estimate (ALE) of experiments involving self-generated thought. Statistical significance was established with 1000 permutations. The map was thresholded at a voxel-wise uncorrected threshold of $p < 0.001$ and cluster-level family-wise error rate threshold of $p < 0.01$. Consistent with previous studies, ALE reveals a constellation of regions typically referred to as the default network (Raichle et al., 2001; Buckner et al., 2008; Spreng et al., 2009). However, as previously discussed, ALE cannot reveal functional sub-domains within self-generated thought without prior assumptions about the sub-domains. Therefore, in the next section, we explored the use of the author-topic model.

3.3.2. Cognitive components of self-generated thought

Fig. 7 shows the cognitive components of self-generated thought estimated by the author-topic model. Fig. 7A shows the BIC score as a function of the number of estimated cognitive components. A higher BIC score indicates a better model. Because the 2-component estimate achieved the highest BIC score, subsequent results will focus on the 2-component estimate.

The 2-component estimate is shown in Fig. 7B. The seven tasks recruited the two cognitive components to different degrees. The top

² Recall that simulation runs were discarded if ICA yielded negative weights.

³ Note that this is after adding back the mean signal removed during the ICA de-meaning step.

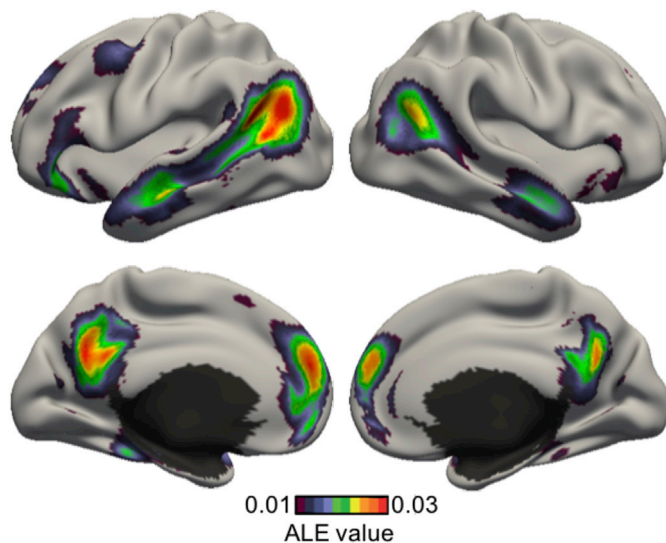


Fig. 6. Activation likelihood estimate (ALE) of experiments involving self-generated thought. Consistent with previous studies, ALE reveals a set of regions corresponding to the default network. However, ALE cannot provide insights into functional sub-domains without prior assumptions about the sub-domains.

tasks recruiting component C1 were “Navigation” and “Autobiographical Memory”. In contrast, the top tasks recruiting component C2 were “Narrative Comprehension”, “Theory of Mind (story-based)”, “Task deactivation”, “Theory of Mind (nonstory-based)”, and “Moral Cognition”.

Compared with Fig. 6, the two cognitive components appeared to decompose the activation pattern revealed by ALE. The two cognitive components appeared to activate different portions of the default network (Fig. 7B). Focusing our attention to the medial cortex, both components had high probability of activating the medial parietal cortex. However, while component C2's activation was largely limited to the precuneus, component C1's activation also included the posterior cingulate and retrosplenial cortices in addition to the precuneus. Both components also had high probability of activating the medial prefrontal cortex (MPFC). However, component C1's activations were restricted to the middle portion of the MPFC, while component C2's activations were restricted to the dorsal and ventral portions of the MPFC. Finally, component C1, but not component C2, had high probability of activating the hippocampal complex.

Switching our attention to the lateral cortex, component C1 had high probability of activating the posterior inferior parietal cortex, while component C2 had high probability of activating the entire stretch of cortex from the temporo-parietal junction to the temporal pole. Component C2 was significantly more likely than component C1 to activate the inferior frontal gyrus.

3.3.3. Goodness of fit

Fig. 8 shows the correlation matrix between the empirical activation maps of seven tasks involving self-generated thought (rows) and seven task activation maps reconstructed from the author-topic model parameter estimates (columns). The diagonal entries of the correlation matrix were significantly higher than the off-diagonal entries: average diagonal entry was 0.69, while the average off-diagonal entry was 0.50. Overall, this suggests a good model fit. However, the diagonal entries were not always the highest and there was a clear block-diagonal structure. Not surprisingly, the top left block corresponded to the top two tasks recruiting component C1 (Fig. 7), while the bottom right block corresponded to the top five tasks recruiting component C2 (Fig. 7).

3.3.4. Correspondence with resting-state networks

The average probability of each self-generated thought cognitive component activating each resting-state network (Yeo et al., 2011) is shown in Fig. S2. Four resting-state networks with the highest probability of being activated by either component are shown in Fig. 9. Three of these resting-state networks were previously considered to be fractionation of the default network (Yeo et al., 2014).

The Default C resting-state network was most strongly activated by component C1, while the temporal parietal resting-state network was most strongly activated by component C2. On the other hand, Default A and B resting-state networks were preferentially activated by components C2.

3.4. Left inferior frontal junction (IFJ)

3.4.1. ALE meta-analysis of the left IFJ's co-activation pattern

Fig. 10 shows the co-activation pattern of the left IFJ estimated by the application of ALE to meta-analytic co-activation modeling (Muhle-Karbe et al., 2015). Statistical significance was established with 1000 permutations. The map was thresholded at a voxel-wise uncorrected threshold of $p < 0.001$ and cluster-level family-wise error rate threshold of $p < 0.01$. The co-activation pattern was mostly bilateral and involved dorsolateral prefrontal cortex, anterior insula, superior parietal lobules, posterior medial frontal cortex and the fusiform gyri. As previously discussed, ALE delineates regions consistently activated across studies, but cannot reveal potential task-dependent co-activation patterns. Therefore, in the next section, we explored the use of the author-topic model.

3.4.2. Task-dependent co-activation patterns of the left IFJ

Fig. 11 shows the co-activation patterns of the left IFJ estimated by the author-topic model. Fig. 11A shows the BIC score as a function of the number of estimated co-activation patterns. There were two peaks corresponding to the 3-pattern and 5-pattern estimates. Fig. S3 shows the 5-pattern estimate. Although the 5-pattern estimate had a higher BIC score than the 3-pattern estimate, the co-activation patterns appeared to fractionate the IFJ into smaller territories. While this fractionation was intriguing, our goal was to examine if the IFJ exhibited task-dependent co-activation patterns and not whether it can be further fractionated. Thus, the 5-pattern estimate represented a degenerate solution from this perspective.⁴

Fig. 11B shows the co-activation patterns from the 3-pattern estimate. Unlike the 5-pattern estimate, the 3 co-activation patterns appeared to overlap completely within the IFJ. Therefore, subsequent results will focus on the 3-pattern estimate. Overall the 3 co-activation patterns appeared to decompose the consensus co-activation pattern revealed by ALE (Fig. 10).

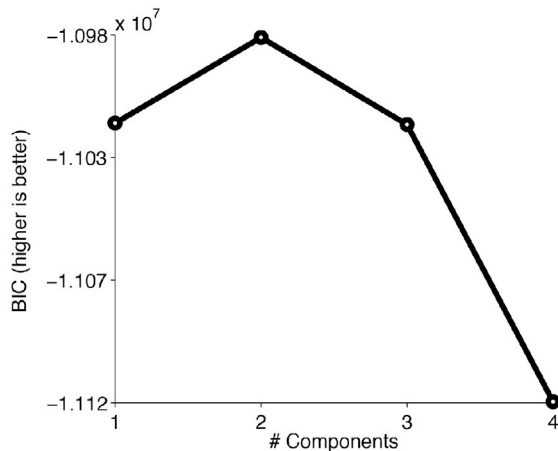
Co-activation pattern C1 was left lateralized and might be recruited in tasks involving language processing. Co-activation pattern C2 involved bilateral superior parietal and posterior medial frontal cortices, and might be recruited in tasks involving attentional control. Co-activation pattern C3 involved bilateral frontal cortex, anterior insula and posterior medial frontal cortex, and might be recruited in tasks involving inhibition or response conflicts.

We now discuss in detail spatial differences among the co-activation patterns. Co-activation pattern C3 strongly engaged bilateral anterior insula, while co-activation pattern C1 only engaged left anterior insula. The activation of the anterior insula was much weaker in co-activation pattern C2.

In the frontal cortex, co-activation pattern C1 had high probability of activating the left inferior frontal gyrus, while co-activation pattern C3

⁴ Given that the “undesirable” 5-pattern estimate had the highest BIC, these results emphasized the fact that the BIC should only be treated as a guide to the number of cognitive components or co-activation patterns, rather than providing a definitive answer.

A. Bayesian Information Criterion (BIC)



B. Cognitive components of self-generated thought

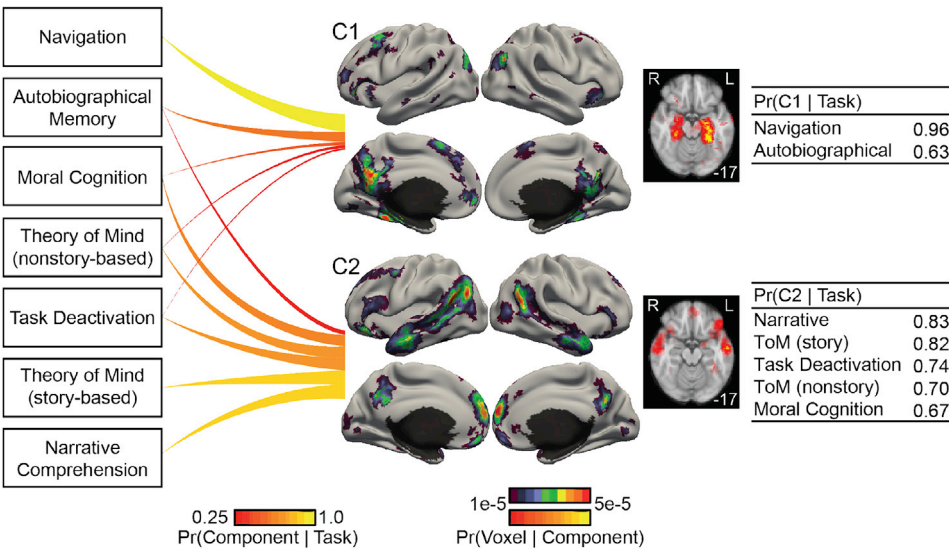


Fig. 7. Cognitive components of self-generated thoughts. (A) Bayesian Information Criterion (BIC) plotted as a function of the number of estimated cognitive components. A higher BIC indicates a better model. BIC peaks at 2 components. (B) 2-component model estimates. Each line connects 1 task with 1 component. The thickness and brightness of the lines are proportional to the magnitude of $\Pr(\text{component} | \text{task})$. For each component, the four leftmost figures show the surface-based visualization for the probability of components activating different brain voxels (i.e., $\Pr(\text{voxel} | \text{component})$), whereas the rightmost figure show a volumetric slice highlighting subcortical structures being activated differently across components. The top color bar is utilized for the surface-based visualization, whereas the bottom color bar is utilized for the volumetric slices. The tables on the right show the top tasks most likely to recruit the two components. The numbers in the right column correspond to $\Pr(\text{component} | \text{task})$. Navigation and Autobiographical Memory preferentially recruited component C1, whereas Narrative Comprehension, Theory of Mind (ToM), Task Deactivation and Moral Cognition preferentially recruited component C2.

had high probability of activating bilateral dorsal lateral prefrontal cortex. Although all three co-activation patterns also had high probability of activating the posterior medial frontal cortex (PMFC), the activation shifted anteriorly from co-activation patterns C1 to C2 to C3.

In the parietal cortex, co-activation pattern C2 included the superior parietal lobule and the intraparietal sulcus in both hemispheres. C1 and C3 did not activate the superior parietal cortex. Finally, co-activation pattern C1 engaged bilateral superior temporal cortex, which might overlap with early auditory regions. Both co-activation patterns C1 and C2 also had high probability of activating ventral visual regions, especially in the fusiform gyrus.

The top three tasks recruiting each co-activation pattern is shown in Fig. 11B. For completeness, the top five tasks recruiting each co-activation pattern are shown in Table S2. The top tasks with the highest probability of recruiting co-activation pattern C1 were “Semantic Monitoring/Discrimination”, “Covert Reading”, and “Phonological Discrimination”. The top tasks recruiting co-activation pattern C2 were “Counting/Calculation”, “Task Switching”, and “Wisconsin Card Sorting Test”. The top tasks recruiting co-activation pattern C3 were “Go/No-Go”, “Encoding”, and “Overt Word Generation”.

At first glance, the top three tasks for co-activation pattern C3 (“Go/No-Go”, “Encoding”, and “Overt Word Generation”) might not seem to be similar tasks. The reason for this incongruence was previously explained in Section 2.5.3 and was due to the fact that the experiments activating

IFJ might not be representative of their task categories. Indeed, of the 123 BrainMap experiments labeled as the “Overt Word Generation” task, only 6 experiments reported activation in the IFJ. Thus, the 6 experiments were not simply “Overt Word Generation” task, but “Overt Word Generation” experiments that happened to activate the IFJ. This motivated further examination of the original publications associated with the top experiments activating IFJ in order to interpret the co-activation patterns (see Section 4.2 for discussion).

Table S3-A to S3-C list the top three experiments with the highest $\Pr(\text{co-activation pattern} | \text{experiment})$ for each of the top three tasks associated with each co-activation pattern. For example, Tables S3-A lists the top three experiments employing “Semantic Monitoring/Discrimination”, “Covert Reading” or “Phonological Discrimination” with the highest $\Pr(\text{co-activation pattern C1} | \text{experiment})$.

To further ensure that the 3 co-activation patterns were not fractionating IFJ (like the 5-pattern estimate), Fig. S4 illustrates the activation foci of the top three experiments with the highest $\Pr(\text{co-activation pattern} | \text{experiment})$ for each of the top three tasks associated with each co-activation pattern falling inside the IFJ. Table 1 shows the mean and standard deviation of the coordinates of these activation foci (within IFJ) for each co-activation pattern. The mean locations of the IFJ activations across co-activation patterns did not differ by more than 2.5 mm along any dimension, suggesting that the co-activation patterns were probably not simply sub-dividing the IFJ.

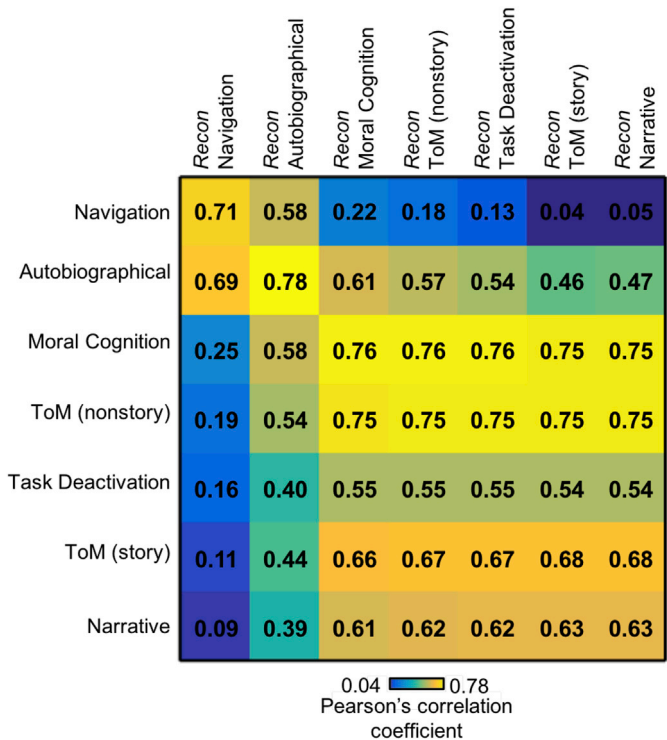


Fig. 8. Goodness of fit of the author-topic model for self-generated thought. The matrix represents the correlations between the empirical activation maps (rows) and reconstructed activation maps (columns) of seven tasks. The tasks follow the same ordering as in Fig. 7. The diagonal values (average $r = 0.69$) were larger than off-diagonal values (average $r = 0.50$), suggesting a good model fit.

3.4.3. Goodness of fit

Fig. 12 shows the correlation matrix between IFJ's co-activation patterns (columns) and the average activation maps of the top three tasks associated with each co-activation pattern (rows). The diagonal entries of the correlation matrix were significantly higher than the off-diagonal entries: average diagonal entry was 0.75, while the average

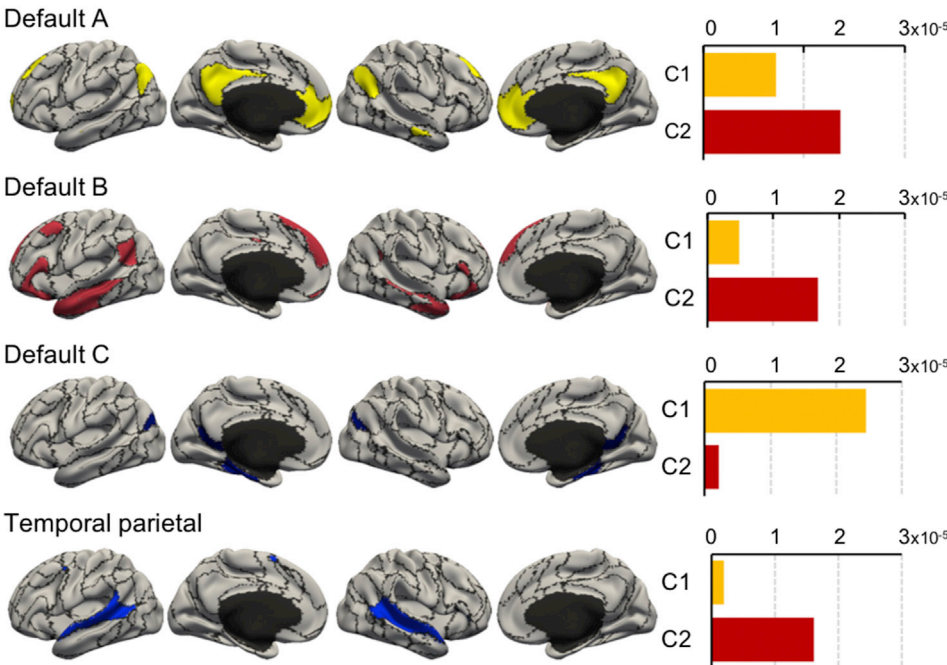


Fig. 9. Average probability of self-generated thought cognitive components activating voxels within 4 resting-state networks (Yeo et al. 2011). The naming of the four resting-state networks followed the convention of previous literature (Kong et al., 2018; Li et al., 2019). Default C resting-state network was primarily activated by component C1, while the temporal parietal resting-state network was primarily activated by component C2. On the other hand, Default A and B resting-state networks were preferentially activated by component C2.

off-diagonal entry was 0.31. Overall, this suggests a good model fit.

3.5. Control analyses

3.5.1. Smoothing

To create the input data for the author-topic model, the activation foci were smoothed with a 10 mm binary smoothing kernel (see Section 2.2.3), consistent with previous work (Wager et al., 2003; Yarkoni et al., 2011; Yeo et al., 2015). Using different smoothing radii yielded similar cognitive components of self-generated thought (Fig. S5) and co-activation patterns of the IFJ (Fig. S6).

3.5.2. Independent component analysis

For comparison, Fig. S7 shows the ICA (ICA1-nilearn) estimate of 2 components of self-generated thought. The estimates were quite similar

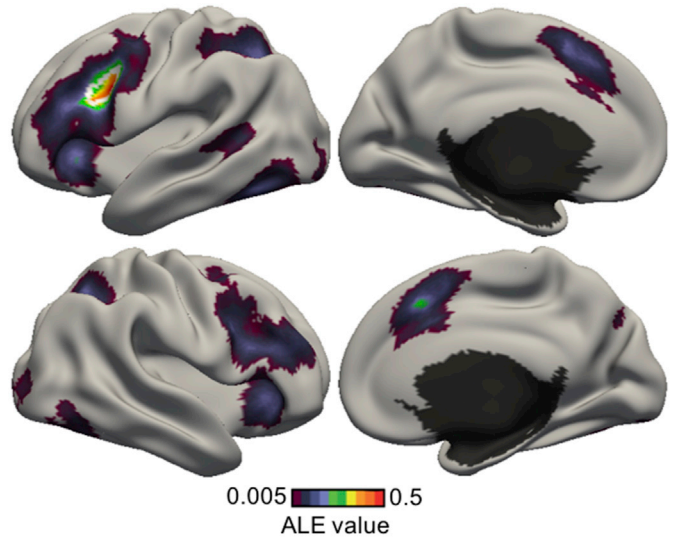
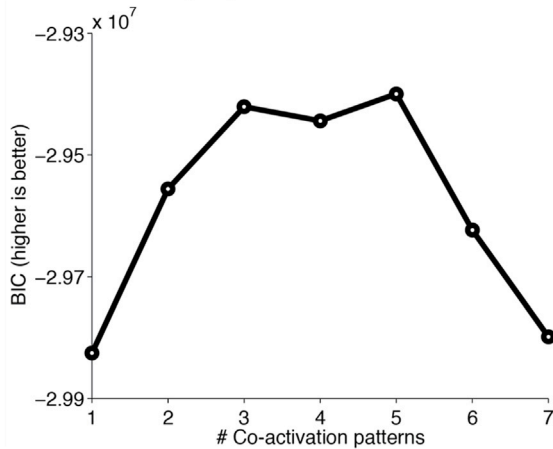


Fig. 10. Co-activation pattern of the left inferior frontal junction (IFJ) estimated by the application of ALE to perform meta-analytic co-activation mapping. The IFJ seed region is delineated by a white boundary.

A. Bayesian Information Criterion (BIC)



B. Co-activation patterns involving the inferior frontal junction (IFJ)

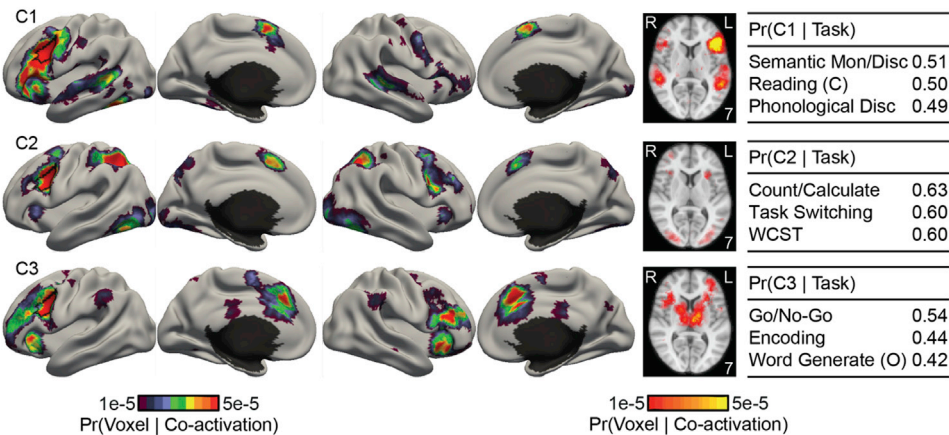


Fig. 11. Co-activation patterns involving the inferior frontal junction (IFJ). (A) Bayesian Information Criterion (BIC) plotted as a function of the number of estimated co-activation patterns. BIC peaks at 3 co-activation patterns. (B) 3-coactivation-pattern model estimates for the IFJ. Format follows Fig. 7. “(C)” and “(O)” indicate “covert” and “overt” respectively. “Mon”, and “Disc” are short for “monitor” and “discrimination” respectively. “Count/Calculate” is short for “Counting/Calculation”. “WCST” is short for “Wisconsin Card Sorting Test”. The left IFJ is delineated by the black boundary in the left hemisphere.

Table 1

Spatial locations of activation foci within the IFJ. Each row of the table shows the mean (standard deviation) of the coordinates of the activation foci (within IFJ) reported by the top 3 experiments with the highest Pr (co-activation pattern | experiment) for each of the top three tasks associated with each co-activation pattern falling inside the IFJ. See Fig. S4 for volumetric slices illustrating the locations of the activation foci within the IFJ. Across the 3 co-activation patterns, the mean coordinates of the top experiments do not differ by more than 2.5 mm in any dimension, suggesting that the co-activation patterns were not fractionating the IFJ.

	x/mm	y/mm	z/mm
Co-activation pattern C1	−40.33 (1.80)	3.89 (3.55)	30.67 (4.21)
Co-activation pattern C2	−39.33 (3.16)	5.33 (4.92)	31.78 (5.45)
Co-activation pattern C3	−39.40 (4.60)	6.40 (5.68)	29.70 (4.64)

to the author-topic estimate. However, the spatial maps contained negative values, which was inappropriate in the context of coordinate-based meta-analysis.

Fig. S8 shows the ICA (ICA1-nilearn) estimate of 3 co-activation patterns of left IFJ. However, the 3 independent components appeared to fractionate the left IFJ into smaller territories (Fig. S8), suggesting a degenerate solution to our problem, similar to the situation with the 5-pattern author-topic estimate (Fig. S3). Furthermore, both the mixture weights and spatial maps contained negative values, which were not interpretable in the context of coordinate-based meta-analysis (Section 2.3.1).

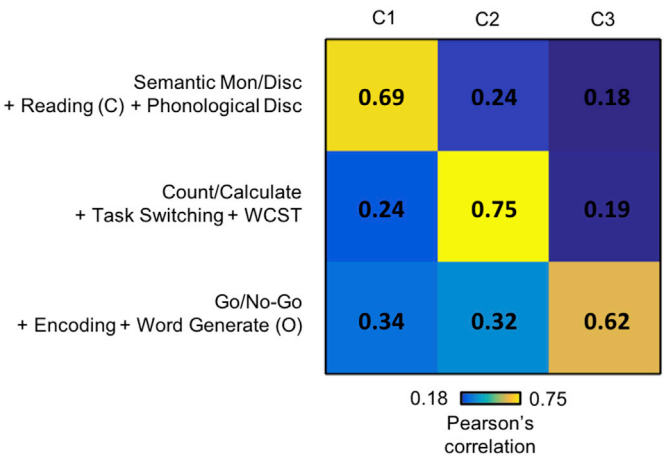


Fig. 12. Goodness of fit of the author-topic model for IFJ. The matrix represents the correlations between IFJ's co-activation patterns (columns) and the average activation maps of the top three tasks associated with each co-activation pattern (rows). The top tasks of each co-activation patterns are shown in Fig. 11. The diagonal values (average $r = 0.75$) were larger than off-diagonal values (average $r = 0.31$), suggesting a good model fit.

4. Discussion

The author-topic model encodes the intuitive notion that behavioral tasks recruited multiple cognitive components, supported by multiple brain regions (Mesulam, 1990; Poldrack, 2006; Barrett and Satpute,

2013). We have previously utilized the author-topic model for large-scale meta-analysis across functional domains (Yeo et al., 2015; Bertolero et al., 2015). By exploiting a recently developed CVB algorithm for the author-topic model (Ngo et al., 2016), we show that the model can also be utilized for small-scale meta-analyses focusing on discovering functional sub-domains or task-dependent co-activation patterns.

A dominant approach for small-scale meta-analyses is ALE, which seeks to find consistent activations across neuroimaging experiments within a functional domain or mental disorder or seed region (also known as MACM). ALE treats heterogeneity across experiments as noise. By contrast, the author-topic model evaluates whether the heterogeneity might be indicative of robust latent patterns within the data. We applied the author-topic model to two applications: one on fractionating a functional sub-domain and one on discovering multiple task-dependent co-activation patterns.

In the first application, the author-topic model encoded the notion that tasks involving self-generated thought might recruit one or more spatially overlapping cognitive components. The model revealed two cognitive components that appeared to delineate two overlapping default sub-networks, consistent with the hypothesized functional organization of the default network (Andrews-Hanna et al., 2014). In the second application, the author-topic model encoded the notion that experiments activating a brain region might recruit one or more co-activation patterns dependent on task contexts (McIntosh, 2000). In the current application, the model revealed that the IFJ participated in three co-activation patterns, suggesting that IFJ flexibly co-activate with different brain regions depending on the cognitive demands of different tasks. Overall, our work suggests that the author-topic model is a versatile tool suitable for both small-scale and large-scale meta-analyses.

4.1. Cognitive components of self-generated thought

Self-generated thought is a heterogeneous set of cognitive processes that includes inferring other people's mental states, dealing with challenging moral scenarios, understanding narratives, retrieving autobiographical memories, internalizing semantic information, and mind-wandering. These processes are characterized by an absence of external stimuli, self-related, and often involve simulation or inferential reasoning (Buckner et al., 2008; Spreng et al., 2009; Smallwood et al., 2011; Baird et al., 2011; Prebble et al., 2013; Smallwood, 2013). Studies of tasks involving self-generated thought have consistently found the activation of the default network, suggesting its functional importance (Buckner et al., 2008; Spreng et al., 2009; Andrews-Hanna et al., 2010a; Andrews-Hanna, 2012; Gorgolewski et al., 2014; Callard and Margulies, 2014). Additionally, the default network has been fractionated into sub-networks supporting different aspects of these stimulus independent cognitive processes (Buckner et al., 2008; Uddin et al., 2009; Sestieri et al., 2011; Andrews-Hanna et al., 2010b; Kim, 2012; Seghier and Price, 2012; Salomon et al., 2014; Bzdok et al., 2013).

The author-topic model revealed two cognitive components of self-generated thought that appeared to fractionate the default network (Fig. 7). The default network has been defined as the set of brain regions that are more active during passive task conditions relative to active task conditions (Shulman et al., 1997; Buckner et al., 2008). While there have been multiple studies fractionating the default network (Andrews-Hanna et al., 2010b; Mayer et al., 2010; Kim, 2011; Yeo et al., 2014; Humphreys et al., 2015), the specific patterns of fractionation have differed across studies. The spatial topography of components C1 and C2 in this paper corresponded well to the previously proposed “medial temporal subsystem” and “dorsal medial subsystem” respectively (Fig. 3A of Andrews-Hanna et al., 2014; Andrews-Hanna et al., 2010b).

The first cognitive component C1 was strongly recruited by navigation and autobiographical memory tasks, suggesting its involvement in constructive mental simulation based upon mnemonic content (Andrews-Hanna et al., 2014). Constructive mental simulation is the process of combining information from the past in order to create a novel

mental representation, such as imagining the future (Buckner and Carroll, 2007; Hassabis and Maguire, 2007; Schacter et al., 2007). “Navigation” tasks require constructive mental simulation to create a mental visualization (“scene construction”) for planning new routes and finding ways in unfamiliar contexts (Burgess et al., 2002; Byrne et al., 2007). On the other hand, “Autobiographical Memory” tasks require constructive mental simulation to project past experience (“constructive episodic simulation”; Atance and O'Neill, 2001; Schacter et al., 2007) or previously acquired knowledge (“semantic memory”; Irish et al., 2012; Brown et al., 2014) across spatiotemporal scale to enact novel perspectives. Overall, cognitive component C1 seems to support the projection of self, events, experiences, images and knowledge to a new temporal or spatial context based upon an associative constructive process, likely mediated by the hippocampus and connected brain structures (Moscovitch et al., 2016; Christoff et al., 2016).

The second cognitive component C2 was strongly recruited by narrative comprehension and theory of mind, suggesting its involvement in mentalizing, inferential, and conceptual processing (Andrews-Hanna et al., 2014). Mentalizing is the process of monitoring one's own mental states or predicting others' mental states (Frith and Frith, 2003), while conceptual processing involves internalizing and retrieving semantic or social knowledge (Binder and Desai, 2011; Van Overwalle, 2009). “Narrative Comprehension” engages conceptual processing to understand the contextual settings of the story, and requires mentalizing to follow and infer the characters' thoughts and emotions (Gernsbacher et al., 1998; Mason et al., 2008). “Theory of Mind” tasks require the recall of learned knowledge, social norms and attitudes to form a meta-representation of the perspectives of other people (Leslie, 1987; Frith and Frith, 2005; Binder and Desai, 2011). The grouping of Narrative Comprehension and Theory of Mind tasks echoes the link between the ability to comprehend narratives and the ability to understanding others' thoughts in developmental studies of children (Guajardo and Watson, 2002; Slaughter et al., 2007; Mason et al., 2008).

The two cognitive components had high probability of activating common and distinct brain regions. Both components engaged the posterior cingulate cortex and precuneus, which are considered part of the “core” sub-network that subserves personally relevant information necessary for both constructive mental simulation and mentalizing (Andrews-Hanna et al., 2014). The distinct brain regions supporting each cognitive component also corroborated the distinct functional role of each component. For instance, component C1, but not C2, had high probability of activating the medial temporal lobe and hippocampus. This is consistent with neuropsychological literature showing that patients with impairment of the medial temporal lobe and hippocampus retain theory of mind and narrative construction capabilities, while suffering deficits in episodic memories and imagining the future (Hassabis et al., 2007; Rosenbaum et al., 2007, 2009; Race et al., 2011).

The cognitive components of self-generated thought estimated by the author-topic model overlapped with default sub-networks A, B and C, as well as the temporal parietal network from a previously published resting-state parcellation (Yeo et al., 2011; Kong et al., 2018). The components loaded differentially on the resting-state networks, thus providing insights into the functions of distinct resting-state networks. Although resting-state fMRI is a powerful tool for extracting brain networks, participants do not actively perform a task during resting-state fMRI. Thus, coordinate-based meta-analysis can be used in conjunction with resting-state fMRI to discover new insights into brain networks and their functions (Seeley et al., 2007; Smith et al., 2009; Laird et al., 2011).

4.2. Co-activation patterns of the left IFJ

The inferior frontal junction (IFJ) is located in the prefrontal cortex at the intersection between the inferior frontal sulcus and the inferior precentral sulcus (Brass et al., 2005; Derrfuss et al., 2005). The IFJ has been suggested to be involved in a wide range of cognitive functions, including task switching (Brass and Von Cramon, 2002; Derrfuss et al.,

2004, 2005), attentional control (Asplund et al., 2010; Baldauf and Desimone, 2014), detection of conflicting responses (Chikazoe et al., 2009a, b; Levy and Wagner, 2011), short-term memory (Zanto et al., 2010; Sneve et al., 2013), construction of attentional episodes (Duncan, 2013) and so on. Using the author-topic model, we found that the IFJ participated in three task-dependent co-activation patterns.

Co-activation pattern C1 might be involved in some aspects of language processing, such as phonological processing for lexical understanding. Phonological processing is an important linguistic function, concerning the use of speech sounds in handling written or oral languages (Wagner and Torgesen, 1987; Poldrack et al., 1999; Friederici, 2002). The top tasks associated with C1 were “Semantic Monitoring/Discrimination”, “Covert Reading”, and “Phonological Discrimination” (Fig. 11B). Inspecting the top three experiments recruiting these three tasks (Tables S3–A) offered more insights into the functional characteristics of co-activation pattern C1. The top “Semantic Monitoring/Discrimination” experiments with the highest probability of recruiting co-activation pattern C1 examined retrieval of semantic meaning (Thompson-Schill et al., 1999; Wagner et al., 2001) and an experiment requiring lexical perception and not just perception of elementary sounds (Poeppel et al., 2004). The top “Covert Reading” experiments most strongly associated with C1 identified a common brain network activated by both reading and listening (Jobard et al., 2007), as well as language comprehension across different media (Small et al., 2009), suggesting the involvement of C1 in generic language comprehension. Among “Phonological Discrimination” experiments, C1 was most highly associated with experiments engaging transcoding of phonological representation for semantic perception (Xu et al., 2001; Démonet et al., 1994). The language and phonological processing interpretation was supported by C1’s strong left lateralization with high probability of activating classical auditory and language brain regions, including the left (but not right) inferior frontal gyrus and bilateral superior temporal cortex.

Co-activation pattern C2 might be engaged in attentional control, especially aspects of task maintenance and shifting of attentional set. Attentional set-shifting is the ability to switch between mental states associated with different reactionary tendencies (Omori et al., 1999; Konishi et al., 1998). The top three tasks most highly associated with C2 were “Counting/Calculation”, “Task Switching”, and “Wisconsin Card Sorting Test” (Fig. 11B). Inspecting the top three experiments under the top task paradigms provided further insights into the functional characteristics of co-activation pattern C2 (Tables S3–B). The top “Counting/Calculation” experiments most strongly recruiting co-activation pattern C2 involved switching of resolution strategies in executive function. For example, one experimental contrast seeks to isolate demanding mental calculation but not retrieval of numerical facts (Zago et al., 2001; Rivera et al., 2002), suggesting C2’s involvement in the selection and application of strategies to solve arithmetic problems. The top “Task Switching” experiments most strongly associated with C2 involved the switching of mental states to learn new stimulus-response or stimulus-outcome associations (Omori et al., 1999; Nagahama et al., 2001; Sylvester et al., 2003). C2 was also strongly expressed by “Wisconsin Card Sorting Test” (WCST) experiments, which required attentional set-shifting to change behavioral patterns in reaction to changes of perceptual dimension (color, shape, or number) upon which the target and reference stimuli were matched (Berman et al., 1995; Konishi et al., 2002; Konishi et al., 2003). Overall, the attentional control interpretation of co-activation pattern C2 is supported by C2’s high probability of activating classical attentional control regions, such as the superior parietal lobule and the intra-parietal sulcus, although there is a clear lack of DLPFC activation.

Co-activation pattern C3 might be engaged in inhibition or response conflict resolution. Conflict-response resolution is a central aspect of cognitive control, which involves monitoring and mediating incongruous response tendencies (Pardo et al., 1990; Braver et al., 2001; Barch et al., 2001). Co-activation pattern C3 is most strongly recruited by experiments utilizing “Go/No-Go”, “Encoding” and “Overt Word Generation” tasks (Fig. 11B). Closer examination of the top three experiments under

each task paradigm provided further insights into the functional characteristics of C3 (Tables S3–C). The top experiments utilizing “Go/No-Go” required the monitoring, preparing and reconciling of conflicting tendencies to either giving a “go” or “stop” (no-go) response (Chikazoe et al., 2009a, b; Simões-Franklin et al., 2010; Kawashima et al., 1996). It might be surprising at first glance that the “Go/No-Go” task was grouped together with “Encoding” and “Overt Word Generation” tasks. However, the top experiments utilizing the “Encoding” and “Overt Word Generation” task all required subjects to make competing decisions (Tables S3–C). The top “Encoding” experiments most strongly associated with C3 required selective association of to-be-learned items with existing memory or knowledge organization for effective enduring retention of new information (Kapur et al., 1996; Callan and Schweighofer, 2010; Mickley and Kensinger, 2009). The top experiments utilizing the “Overt Word Generation” task required subjects to make competing decision, such as inhibiting verbalization of wrong words in verbal fluency task (Baker et al., 1997; Ravnkilde et al., 2002) or inhibiting a predominant pattern (regular past-tense verbs) in favor of generating less conventional forms (irregular past-tense verbs) (Desai et al., 2006). Overall, the inhibition or response conflict interpretation of co-activation pattern C3 is supported by C3’s high probability of activating classical executive function regions, including the bilateral dorsal lateral prefrontal cortex.

The intriguing location of the left IFJ and its functional heterogeneity suggests the role of IFJ as an integrative hub for different cognitive functions. For example, the IFJ has been suggested to consolidate information streams for cognitive control from its bordering brain regions (Brass et al., 2005). The involvement of the IFJ in three task-dependent co-activation patterns supported the view that the IFJ orchestrates different cognitive mechanisms to allow their operations in harmony.

5. Conclusion

Heterogeneities across neuroimaging experiments are often treated as noise in coordinate-based meta-analyses. Here we demonstrate that the author-topic model can be utilized to determine if the heterogeneities can be explained by a small number of latent patterns. In the first application, the author-topic model revealed two overlapping cognitive components subserving self-generated thought. In the second application, the author-topic revealed the participation of the left IFJ in three task-dependent co-activation patterns. These applications exhibited the broad utility of the author-topic model, ranging from discovering functional subdomains or task-dependent co-activation patterns.

Acknowledgements

This work was supported by Singapore MOE Tier 2 (MOE2014-T2-2-016), NUS Strategic Research (DPRT/944/09/14), NUS SOM Aspiration Fund (R185000271720), Singapore NMRC (CBRG/0088/2015), NUS YIA and the Singapore National Research Foundation Fellowship (Class of 2017). Simon Eickhoff is supported by the National Institute of Mental Health (R01-MH074457), the Helmholtz Portfolio Theme “Supercomputing and Modeling for the Human Brain” and the European Union’s Horizon 2020 Research and Innovation Programme under Grant Agreement No. 7202070 (HBP SGA1). R. Nathan Spreng is supported by the Natural Sciences and Engineering Research Council of Canada, the Canadian Institutes of Health Research, and received salary support from the Fonds de la Recherche du Québec – Santé (FRQS). Comprehensive access to the BrainMap database was authorized by a collaborative-use license agreement (<http://www.brainmap.org/collaborations.html>). BrainMap database development is supported by NIH/NIMH R01 MH074457. Our research also utilized resources provided by the Center for Functional Neuroimaging Technologies, NIH P41EB015896 and instruments supported by NIH 1S10RR023401, NIH 1S10RR019307, and NIH 1S10RR023043 from the Athinoula A. Martinos Center for Biomedical Imaging at the Massachusetts General Hospital. Our computational work was partially performed on resources of the National

Supercomputing Centre, Singapore (<https://www.nsc.sg>).

Appendix A. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.neuroimage.2019.06.037>.

Appendix

Pseudo-code	Function
(1) Read the input activation foci and task labels	CBIG_AuthorTopic_PreprocessInput
(2) Initialize the model's hyperparameters	CBIG_AuthorTopic_SetupParameters
(3) Repeat for N = 1000 re-initializations	
(a) Initialize the variational parameters ϕ	CBIG_AuthorTopic_InitializeParams
(b) Update the variational parameters ϕ (Eq. 14 in Supplemental S1)	
- Approximate $E_q(N_{\cdot\cdot}^{-ef})$ and $Var_q(N_{\cdot\cdot}^{-ef})$	CBIG_AuthorTopic_ComputeVariationalTerm_N_c
- Approximate $E_q(N_{\cdot\cdot cv}^{-ef})$ and $Var_q(N_{\cdot\cdot cv}^{-ef})$	CBIG_AuthorTopic_ComputeVariationalTerm_N_cv
- Approximate $E_q(N_{\cdot\cdot t}^{-ef})$ and $Var_q(N_{\cdot\cdot t}^{-ef})$	CBIG_AuthorTopic_ComputeVariationalTerm_N_t
- Approximate $E_q(N_{\cdot\cdot ic}^{-ef})$ and $Var_q(N_{\cdot\cdot ic}^{-ef})$	CBIG_AuthorTopic_ComputeVariationalTerm_N_tc
- Update the variational parameters ϕ (Eq. 14 in Supplemental S1)	
- Recompute the variational log likelihood. If it converges, go to step (3c), otherwise repeat from step (3b)	
(c) Update the model parameters θ_{ic} and β_{cv} (Eq. 16 and 17 in Supplemental S1)	CBIG_AuthorTopic_EstimateParams

Pseudo-code of the Collapsed Variational Bayes (CVB) algorithm for estimating the author-topic model's parameters. The left column outlines the main steps of the algorithm. The right column denotes the functions in the source code that correspond to each step. The source code of the CVB algorithm and the input file of the self-generated thought dataset are available at https://github.com/ThomasYeoLab/CBIG/tree/master/stable_projects/meta-analysis/Ngo2019_AuthorTopic. Note that steps (2) and (3) can be called by a single function *CBIG_AuthorTopic_RunInference*.

References

Abraham, A., Pedregosa, F., Eickenberg, M., Gervais, P., Mueller, A., Kossaifi, J., Gramfort, A., Thirion, B., Varoquaux, G., 2014. Machine learning for neuroimaging with scikit-learn. *Front. Neuroinf.* 8, 14.

Andrews-Hanna, J.R., Reidler, J.S., Huang, C., Buckner, R.L., 2010a. Evidence for the default networks role in spontaneous cognition. *J. Neurophysiol.* 104, 322–335.

Andrews-Hanna, J.R., Reidler, J.S., Sepulcre, J., Poulin, R., Buckner, R.L., 2010b. Functional-anatomic fractionation of the brains default network. *Neuron* 65, 550–562.

Andrews-Hanna, J.R., Smallwood, J., Spreng, R.N., 2014. The default network and self-generated thought: component processes, dynamic control, and clinical relevance. *Ann. N. Y. Acad. Sci.* 1316, 29–52.

Andrews-Hanna, J.R., 2012. The brain's default network and its adaptive role in internal mentation. *Neuroscientist* 18, 251–270.

Asplund, C.L., Todd, J.J., Snyder, A.P., Marois, R., 2010. A central role for the lateral prefrontal cortex in goal-directed and stimulus-driven attention. *Nat. Neurosci.* 13, 507–512.

Atance, C.M., O'Neill, D.K., 2001. Episodic future thinking. *Trends Cognit. Sci.* 5, 533–539.

Baird, B., Smallwood, J., Schooler, J.W., 2011. Back to the future: autobiographical planning and the functionality of mind-wandering. *Conscious. Cognit.* 20, 1604–1611.

Baker, S.C., Frith, C.D., Dolan, R.J., 1997. The interaction between mood and cognitive function studied with PET. *Psychol. Med.* 27, 565–578.

Baldauf, D., Desimone, R., 2014. Neural mechanisms of object-based attention. *Science* 344, 424–427.

Barch, D.M., Braver, T.S., Akbudak, E., Conturo, T., Ollinger, J., Snyder, A., 2001. Anterior cingulate cortex and response conflict: effects of response modality and processing domain. *Cerebr. Cortex* 11, 837–848.

Barrett, L.F., Satpute, A.B., 2013. Large-scale brain networks in affective and social neuroscience: towards an integrative functional architecture of the brain. *Curr. Opin. Neurobiol.* 23, 361–372.

Beckmann, C.F., Smith, S.M., 2004. Probabilistic independent component analysis for functional magnetic resonance imaging. *IEEE Trans. Med. Imaging* 23, 137–152.

Beissner, F., Meissner, K., Bär, K.J., Napadow, V., 2013. The autonomic brain: an activation likelihood estimation meta-analysis for central processing of autonomic function. *J. Neurosci.* 33, 10503–10511.

Berman, K.F., Ostrem, J.L., Randolph, C., Gold, J., Goldberg, T.E., Coppola, R., Carson, R.E., Herscovitch, P., Weinberger, D.R., 1995. Physiological activation of a cortical network during performance of the Wisconsin Card Sorting Test: a positron emission tomography study. *Neuropsychologia* 33, 1027–1046.

Bertolero, M.A., Yeo, B.T.T., D'Esposito, M., 2015. The modular and integrative functional architecture of the human brain. *Proc. Natl. Acad. Sci. Unit. States Am.* 112, E6798–E6807.

Bertolero, M.A., Yeo, B.T., D'Esposito, M., 2017. The diverse club. *Nat. Commun.* 8, 1277.

Bertolero, M.A., Yeo, B.T.T., Bassett, S.D., D'Esposito, M., 2018. A mechanistic model of connector hubs, modularity, and cognition. *Nature Human Behaviour* 112, E6798.

Binder, J.R., Desai, R.H., Graves, W.W., Conant, L.L., 2009. Where is the semantic system? A critical review and meta-analysis of 120 functional neuroimaging studies. *Cerebr. Cortex* 19, 2767–2796.

Binder, J.R., Desai, R.H., 2011. The neurobiology of semantic memory. *Trends Cognit. Sci.* 15, 527–536.

Braga, R.M., Buckner, R.L., 2017. Parallel interdigitated distributed networks within the individual estimated by intrinsic functional connectivity. *Neuron* 9, 457–471.

Brass, M., Derrfuss, J., Forstmann, B., von Cramon, D.Y., 2005. The role of the inferior frontal junction area in cognitive control. *Trends Cognit. Sci.* 9, 314–316.

Brass, M., von Cramon, D.Y., 2002. The role of the frontal cortex in task preparation. *Cerebr. Cortex* 12, 908–914.

Braver, T.S., Barch, D.M., Gray, J.R., Molfese, D.L., Snyder, A., 2001. Anterior cingulate cortex and response conflict: effects of frequency, inhibition and errors. *Cerebr. Cortex* 11, 825–836.

Brown, A.D., Addis, D.R., Romano, T.A., Marmar, C.R., Bryant, R.A., Hirst, W., Schacter, D.L., 2014. Episodic and semantic components of autobiographical memories and imagined future events in post-traumatic stress disorder. *Memory* 22, 595–604.

Buckner, R.L., Andrews-Hanna, J.R., Schacter, D.L., 2008. The brains default network. *Ann. N. Y. Acad. Sci.* 1124, 1–38.

Buckner, R.L., Carroll, D.C., 2007. Self-projection and the brain. *Trends Cognit. Sci.* 11, 49–57.

Buckner, R.L., Krienen, F.M., Castellanos, A., Diaz, J.C., Yeo, B.T.T., 2011. The organization of the human cerebellum estimated by intrinsic functional connectivity. *J. Neurophysiol.* 106, 2322–2345.

Burgess, N., Maguire, E.A., O'Keefe, J., 2002. The human hippocampus and spatial and episodic memory. *Neuron* 35, 625–641.

Button, K.S., Ioannidis, J.P., Mokrysz, C., Nosek, B.A., Flint, J., Robinson, E.S., Munafò, M.R., 2013. Power failure: why small sample size undermines the reliability of neuroscience. *Nat. Rev. Neurosci.* 14, 365.

Byrne, P., Becker, S., Burgess, N., 2007. Remembering the past and imagining the future: a neural model of spatial memory and imagery. *Psychol. Rev.* 114, 340.

Bzdok, D., Langner, R., Schilbach, L., Engemann, D.A., Laird, A.R., Fox, P.T., Eickhoff, S., 2013. Segregation of the human medial prefrontal cortex in social cognition. *Front. Hum. Neurosci.* 7, 232.

Calhoun, V.D., Adali, T., Pearlson, G.D., Pekar, J.J., 2001. A method for making group inferences from functional MRI data using independent component analysis. *Human Brain Mmapping* 14, 140–151.

Callan, D.E., Schweighofer, N., 2010. Neural correlates of the spacing effect in explicit verbal semantic encoding support the deficient-processing theory. *Hum. Brain Mapp.* 31, 645–659.

Callard, F., Margulies, D.S., 2014. What we talk about when we talk about the default mode network. *Front. Hum. Neurosci.* 8, 619.

Carp, J., 2012. The secret lives of experiments: methods reporting in the fMRI literature. *Neuroimage* 63, 289–300.

Chikazoe, J., Jimura, K., Asari, T., Yamashita, K., Morimoto, H., Hirose, S., Miyashita, Y., Konishi, S., 2009a. Functional dissociation in right inferior frontal cortex during performance of go/no-go task. *Cerebr. Cortex* 19, 146–152.

- Chikazoe, J., Jimura, K., Hirose, S., Yamashita, K., Miyashita, Y., Konishi, S., 2009b. Preparation to inhibit a response complements response inhibition during performance of a stop-signal task. *J. Neurosci.* 29, 15870–15877.
- Christoff, K., Irving, Z.C., Fox, K.C.R., Spreng, R.N., Andrews-Hanna, J.R., 2016. Mind-wandering as spontaneous thought: a dynamic framework. *Nat. Rev. Neurosci.* 17, 718–731.
- Cole, M.W., Reynolds, J.R., Power, J.D., Repovs, G., Anticevic, A., Braver, T.S., 2013. Multi-task connectivity reveals flexible hubs for adaptive task control. *Nat. Neurosci.* 16, 1348.
- Cortese, S., Kelly, C., Chabernaud, C., Proal, E., Di Martino, A., Milham, M.P., Castellanos, F.X., 2012. Toward systems neuroscience of ADHD: a meta-analysis of 55 fMRI studies. *Am. J. Psychiatry* 169, 1038–1055.
- Costafreda, S.G., Brammer, M.J., David, A.S., Fu, C.H.Y., 2008. Predictors of amygdala activation during the processing of emotional stimuli: a meta-analysis of 385 PET and fMRI studies. *Brain Res. Rev.* 58, 57–70.
- Crossley, N.A., Mechelli, A., Scott, J., Carletti, F., Fox, P.T., McGuire, P., Bullmore, E.T., 2014. The hubs of the human connectome are generally implicated in the anatomy of brain disorders. *Brain* 137, 2382–2395.
- Démonet, J.F., Price, C., Wise, R., Frackowiak, R.S., 1994. A PET study of cognitive strategies in normal subjects during language tasks: influence of phonetic ambiguity and sequence processing on phoneme monitoring. *Brain* 117, 671–682.
- Derrfuss, J., Brass, M., Neumann, J., von Cramon, D.Y., 2005. Involvement of the inferior frontal junction in cognitive control: meta-analyses of switching and Stroop studies. *Hum. Brain Mapp.* 25, 22–34.
- Derrfuss, J., Brass, M., Von Cramon, D.Y., 2004. Cognitive control in the posterior frontolateral cortex: evidence from common activations in task coordination, interference control, and working memory. *Neuroimage* 23, 604–612.
- Desai, R., Conant, L.L., Waldron, E., Binder, J.R., 2006. fMRI of past tense processing: the effects of phonological complexity and task difficulty. *J. Cogn. Neurosci.* 18, 278–297.
- Di Martino, A., Ross, K., Uddin, L.Q., Sklar, A.B., Castellanos, F.X., Milham, M.P., 2009. Functional brain correlates of social and nonsocial processes in autism spectrum disorders: an activation likelihood estimation meta-analysis. *Biol. Psychiatry* 65, 63–74.
- Duncan, J., 2010. The multiple-demand (MD) system of the primate brain: mental programs for intelligent behaviour. *Trends Cognit. Sci.* 14, 172–179.
- Duncan, J., 2013. The structure of cognition: attentional episodes in mind and brain. *Neuron* 80, 35–50.
- Eickhoff, S.B., Bzdok, D., Laird, A.R., Kurth, F., Fox, P.T., 2012. Activation likelihood estimation meta-analysis revisited. *Neuroimage* 59, 2349–2361.
- Eickhoff, S.B., Jbabdi, S., Caspers, S., Laird, A.R., Fox, P.T., Zilles, K., Behrens, T.E.J., 2010. Anatomical and functional connectivity of cytoarchitectonic areas within the human parietal operculum. *J. Neurosci.* 30, 6409–6421.
- Eickhoff, S.B., Laird, A.R., Grefkes, C., Wang, L.E., Zilles, K., Fox, P.T., 2009. Coordinate-based activation likelihood estimation meta-analysis of neuroimaging data: a random-effects approach based on empirical estimates of spatial uncertainty. *Hum. Brain Mapp.* 30, 2907–2926.
- Fedorenko, E., Duncan, J., Kanwisher, N., 2013. Broad domain generality in focal regions of frontal and parietal cortex. *Proc. Natl. Acad. Sci. Unit. States Am.* 110, 16616–16621.
- Fischl, B., 2012. FreeSurfer. *Neuroimage* 62, 774–781.
- Fitzgerald, P.B., Laird, A.R., Maller, J., Daskalakis, Z.J., 2008. A meta-analytic study of changes in brain activation in depression. *Hum. Brain Mapp.* 29, 683–695.
- Fox, P.T., Lancaster, J.L., Laird, A.R., Eickhoff, S.B., 2014. Meta-analysis in human neuroimaging: computational modeling of large-scale databases. *Annu. Rev. Neurosci.* 37, 409–434.
- Fox, P.T., Lancaster, J.L., 2002. Mapping context and content: the BrainMap model. *Nat. Rev. Neurosci.* 3, 319–321.
- Friederici, A.D., 2002. Towards a neural basis of auditory sentence processing. *Trends Cognit. Sci.* 6, 78–84.
- Frith, C., Frith, U., 2005. Theory of mind. *Curr. Biol.* 15, R644–R645.
- Frith, U., Frith, C.D., 2003. Development and neurophysiology of mentalizing. *Phil. Trans. Biol. Sci.* 358, 459–473.
- Gernsbacher, M.A., Hallada, B.M., Robertson, R.R.W., 1998. How automatically do readers infer fictional characters emotional states? *Sci. Stud. Read.* 2, 271–300.
- Gorgolewski, K.J., Lurie, D., Urchs, S., Kipping, J.A., Craddock, R.C., Milham, M.P., Margulies, D.S., Smallwood, J., 2014. A correspondence between individual differences in the brain's intrinsic functional architecture and the content and form of self-generated thoughts. *PLoS One* 9, e97176.
- Gorgolewski, K.J., Varoquaux, G., Rivera, G., Schwartz, Y., Ghosh, S.S., Maumet, C., Sochat, V.V., Nichols, T.E., Poldrack, R.A., Poline, J.-B., Yarkoni, T., Margulies, D.S., 2015. NeuroVault.org: a web-based repository for collecting and sharing unthresholded statistical maps of the brain. *Front. Neuroinf.* 9, 8.
- Guajardo, N.R., Watson, A.C., 2002. Narrative discourse and theory of mind development. *J. Genet. Psychol.* 163, 305–325.
- Hassabis, D., Kumaran, D., Vann, S.D., Maguire, E.A., 2007. Patients with hippocampal amnesia cannot imagine new experiences. *Proc. Natl. Acad. Sci. Unit. States Am.* 104, 1726–1731.
- Hassabis, D., Maguire, E.A., 2007. Deconstructing episodic memory with construction. *Trends Cognit. Sci.* 11, 299–306.
- Humphreys, G.F., Hoffman, P., Visser, M., Binney, R.J., Ralph, M.A.L., 2015. Establishing task-and modality-dependent dissociations between the semantic and default mode networks. *Proc. Natl. Acad. Sci. Unit. States Am.* 112, 7857–7862.
- Irish, M., Addis, D.R., Hodges, J.R., Piguet, O., 2012. Considering the role of semantic memory in episodic future thinking: evidence from semantic dementia. *Brain* 135, 2178–2191.
- Jobard, G., Vigneau, M., Mazoyer, B., Tzourio-Mazoyer, N., 2007. Impact of modality and linguistic complexity during reading and listening tasks. *Neuroimage* 34, 784–800.
- Kapur, S., Tulving, E., Cabeza, R., McIntosh, A.R., Houle, S., Craik, F.I., 1996. The neural correlates of intentional learning of verbal materials: a PET study in humans. *Cogn. Brain Res.* 4, 243–249.
- Kawashima, R., Satoh, K., Itoh, H., Ono, S., Furumoto, S., Gotoh, R., Koyama, M., Yoshioka, S., Takahashi, T., Takahashi, K., Yanagisawa, T., 1996. Functional anatomy of GO/NO-GO discrimination and response selection—a PET study in man. *Brain Res.* 728, 79–89.
- Kernbach, J.M., Yeo, B.T.T., Smallwood, J., Margulies, D.S., Thiebaut de Schotten, M., Walter, H., Sabuncu, M.R., Holmes, A.J., Gramfort, A., Varoquaux, G., Thirion, B., Bzdok, D., 2018. Subspecialization within default mode nodes characterized in 10,000 UK Biobank participants. In: *Proceedings of the National Academy of Sciences*, p. 201804876.
- Kim, C., Johnson, N.F., Cilles, S.E., Gold, B.T., 2011. Common and distinct mechanisms of cognitive flexibility in prefrontal cortex. *J. Neurosci.* 31, 4771–4779.
- Kim, H., 2012. A dual-subsystem model of the brain's default network: self-referential processing, memory retrieval processes, and autobiographical memory retrieval. *Neuroimage* 61, 966–977.
- Kong, R., Li, J., Orban, C., Sabuncu, M.R., Liu, H., Schaefer, A., Sun, N., Zuo, X.N., Holmes, A., Eickhoff, S.B., Yeo, B.T.T., 2018. Spatial topography of individual-specific cortical networks predicts human cognition, personality and emotion. *Cerebr. Cortex* 1, 19.
- Konishi, S., Hayashi, T., Uchida, I., Kikyo, H., Takahashi, E., Miyashita, Y., 2002. Hemispheric asymmetry in human lateral prefrontal cortex during cognitive set shifting. *Proc. Natl. Acad. Sci. Unit. States Am.* 99, 7803–7808.
- Konishi, S., Jimura, K., Asari, T., Miyashita, Y., 2003. Transient activation of superior prefrontal cortex during inhibition of cognitive set. *J. Neurosci.* 23, 7776–7782.
- Konishi, S., Nakajima, K., Uchida, I., Kameyama, M., Nakahara, K., Sekihara, K., Miyashita, Y., 1998. Transient activation of inferior prefrontal cortex during cognitive set shifting. *Nat. Neurosci.* 1, 80–84.
- Koski, L., Paus, T., 2000. Functional Connectivity of the Anterior Cingulate Cortex within the Human Frontal Lobe: a Brain-Mapping Meta-Analysis. *Executive Control and the Frontal Lobe: Current Issues*, pp. 55–65.
- Laird, A.R., Eickhoff, S.B., Li, K., Robin, D.A., Glahn, D.C., Fox, P.T., 2009a. Investigating the functional heterogeneity of the default mode network using coordinate-based meta-analytic modeling. *J. Neurosci.* 29, 14496–14505.
- Laird, A.R., Eickhoff, S.B., Kurth, F., Fox, P.M., Uecker, A.M., Turner, J.A., Robinson, J.L., Lancaster, J.L., Fox, P.T., 2009b. ALE meta-analysis workflows via the Brainmap database: progress towards a probabilistic functional brain atlas. *Front. Neuroinf.* 3, 23.
- Laird, A.R., Fox, P.M., Eickhoff, S.B., Turner, J.A., Ray, K.L., McKay, D.R., Glahn, D.C., Beckmann, C.F., Smith, S.M., Fox, P.T., 2011. Behavioral interpretations of intrinsic connectivity networks. *J. Cogn. Neurosci.* 23, 4022–4037.
- Laird, A.R., Fox, P.M., Price, C.J., Glahn, D.C., Uecker, A.M., Lancaster, J.L., Turkeltaub, P.E., Kochunov, P., Fox, P.T., 2005. ALE meta-analysis: Controlling the false discovery rate and performing statistical contrasts. *Hum. Brain Mapp.* 25, 155–164.
- Lancaster, J.L., Tordesillas-Gutiérrez, D., Martínez, M., Salinas, F., Evans, A., Zilles, K., Mazziotta, J.C., Fox, P.T., 2007. Bias between MNI and Talairach coordinates analyzed using the ICBM-152 brain template. *Hum. Brain Mapp.* 28, 1194–1205.
- Leech, R., Braga, R., Sharp, D.J., 2012. Echoes of the brain within the posterior cingulate cortex. *J. Neurosci.* 32, 215–222.
- Leslie, A.M., 1987. Pretense and representation: the origins of "theory of mind". *Psychol. Rev.* 94, 412.
- Levy, B.J., Wagner, A.D., 2011. Cognitive control and right ventrolateral prefrontal cortex: reflexive reorienting, motor inhibition, and action updating. *Ann. N. Y. Acad. Sci.* 1224, 40–62.
- Li, J., Kong, R., Liegeois, R., Orban, C., Sun, N., Holmes, A.J., Sabuncu, M.R., Ge, T., Yeo, B.T.T., 2019. Global signal regression strengthens association between resting-state functional connectivity and behavior. *NeuroImage* 196, 126–141.
- Mar, R.A., 2011. The neural bases of social cognition and story comprehension. *Annu. Rev. Psychol.* 62, 103–134.
- Mason, R.A., Williams, D.L., Kana, R.K., Minshew, N., Just, M.A., 2008. Theory of mind disruption and recruitment of the right hemisphere during narrative comprehension in autism. *Neuropsychologia* 46, 269–280.
- Mayer, J.S., Roebroeck, A., Maurer, K., Linden, D.E.J., 2010. Specialization in the default mode: task-induced brain deactivations dissociate between visual working memory and attention. *Hum. Brain Mapp.* 31, 126–139.
- McIntosh, A.R., 2000. Towards a network theory of cognition. *Neural Networks* 13, 861–870.
- Mesulam, M., 1990. Large-scale neurocognitive networks and distributed processing for attention, language, and memory. *Ann. Neurol.* 28, 597–613.
- Mickley Steinmetz, K.R., Kensinger, E.A., 2009. The effects of valence and arousal on the neural activity leading to subsequent memory. *Psychophysiology* 46, 1190–1199.
- Minzenberg, M.J., Laird, A.R., Thelen, S., Carter, C.S., Glahn, D.C., 2009. Meta-analysis of 41 functional neuroimaging studies of executive function in schizophrenia. *Arch. Gen. Psychiatr.* 66, 811–822.
- Moscovitch, M., Cabeza, R., Winocur, G., Nadel, L., 2016. Episodic memory and beyond: the Hippocampus and neocortex in transformation. *Annu. Rev. Psychol.* 67, 105–134.

- Muhle-Karbe, P.S., Derrfuss, J., Lynn, M.T., Neubert, F.X., Fox, P.T., Brass, M., Eickhoff, S.B., 2015. Co-activation-based parcellation of the lateral prefrontal cortex delineates the inferior frontal junction area. *Cerebr. Cortex* 26, 2225–2241.
- Nagahama, Y., Okada, T., Katsumi, Y., Hayashi, T., Yamauchi, H., Oyanagi, C., Konishi, J., Fukuyama, H., Shibasaki, H., 2001. Dissociable mechanisms of attentional control within the human prefrontal cortex. *Cerebr. Cortex* 11, 85–92.
- Ngo, G.H., Eickhoff, S.B., Fox, P.T., Yeo, B.T.T., 2016. Collapsed variational Bayesian inference of the author-topic model: application to large-scale coordinate-based meta-analysis. In: *Proceedings of the 2016 International Workshop in Pattern Recognition in Neuroimaging (PRNI)*.
- Omori, M., Yamada, H., Murata, T., Sadato, N., Tanaka, M., Ishii, Y., Isaki, K., Yonekura, Y., 1999. Neuronal substrates participating in attentional set-shifting of rules for visually guided motor selection: a functional magnetic resonance imaging investigation. *Neurosci. Res.* 33, 317–323.
- Pardo, J.V., Pardo, P.J., Janer, K.W., Raichle, M.E., 1990. The anterior cingulate cortex mediates processing selection in the Stroop attentional conflict paradigm. *Proc. Natl. Acad. Sci. Unit. States Am.* 87, 256–259.
- Poeppl, D., Guillemin, A., Thompson, J., Fritz, J., Bavelier, D., Braun, A.R., 2004. Auditory lexical decision, categorical perception, and FM direction discrimination differentially engage left and right auditory cortex. *Neuropsychologia* 42, 183–200.
- Poldrack, R.A., Mumford, J.A., Schonberg, T., Kalar, D., Barman, B., Yarkoni, T., 2012. Discovering relations between mind, brain, and mental disorders using topic mapping. *PLoS Comput. Biol.* 8 e1002707.
- Poldrack, R.A., Wagner, A.D., Prull, M.W., Desmond, J.E., Glover, G.H., Gabrieli, J.D.E., 1999. Functional specialization for semantic and phonological processing in the left inferior prefrontal cortex. *Neuroimage* 10, 15–35.
- Poldrack, R.A., Yarkoni, T., 2016. From brain maps to cognitive ontologies: informatics and the search for mental structure. *Annu. Rev. Psychol.* 67, 587–612.
- Poldrack, R.A., 2006. Can cognitive processes be inferred from neuroimaging data? *Trends Cognit. Sci.* 10, 59–63.
- Poline, J., Breeze, J.L., Ghosh, S.S., Gorgolewski, K., Halchenko, Y.O., Hanke, M., Helmer, K.G., Marcus, D.S., Poldrack, R.A., Schwartz, Y., others, 2012. Data sharing in neuroimaging research. *Front. Neuroinf.* 6, 9.
- Prebble, S.C., Addis, D.R., Tippett, L.J., 2013. Autobiographical memory and sense of self. *Psychol. Bull.* 139, 815.
- Race, E., Keane, M.M., Verfaellie, M., 2011. Medial temporal lobe damage causes deficits in episodic memory and episodic future thinking not attributable to deficits in narrative construction. *J. Neurosci.* 31, 10262–10269.
- Raichle, M.E., MacLeod, A.M., Snyder, A.Z., Powers, W.J., Gusnard, D.A., Shulman, G.L., 2001. A default mode of brain function. *Proc. Natl. Acad. Sci. Unit. States Am.* 98, 676–682.
- Ravnkilde, B., Videbech, P., Rosenberg, R., Gjedde, A., Gade, A., 2002. Putative tests of frontal lobe function: a PET-study of brain activation during Stroop's Test and verbal fluency. *J. Clin. Exp. Neuropsychol.* 24, 534–547.
- Reid, A.T., Bzdok, D., Genon, S., Langner, R., Müller, V.I., Eickhoff, C.R., Hoffstaedter, F., Cieslik, E.-C., Fox, P.T., Laird, A.R., others, 2016a. ANIMA: a data-sharing initiative for neuroimaging meta-analyses. *Neuroimage* 124, 1245–1253.
- Rivera, S.M., Menon, V., White, C.D., Glaser, B., Reiss, A.L., 2002. Functional brain activation during arithmetic processing in females with fragile X Syndrome is related to FMR1 protein expression. *Hum. Brain Mapp.* 16, 206–218.
- Robinson, J.L., Laird, A.R., Glahn, D.C., Lovello, W.R., Fox, P.T., 2010. Metaanalytic connectivity modeling: delineating the functional connectivity of the human amygdala. *Hum. Brain Mapp.* 31, 173–184.
- Rosen-Zvi, M., Chemudugunta, C., Griffiths, T., Smyth, P., Steyvers, M., 2010. Learning author-topic models from text corpora. *ACM Trans. Inf. Syst.* 28, 4.
- Rosenbaum, R.S., Gilboa, A., Levine, B., Winocur, G., Moscovitch, M., 2009. Amnesia as an impairment of detail generation and binding: evidence from personal, fictional, and semantic narratives in KC. *Neuropsychologia* 47, 2181–2187.
- Rosenbaum, R.S., Stuss, D.T., Levine, B., Tulving, E., 2007. Theory of mind is independent of episodic memory. *Science* 318, 1257–1257.
- Rottschy, C., Langner, R., Dogan, I., Reetz, K., Laird, A.R., Schulz, J.B., Fox, P.T., Eickhoff, S.B., 2012. Modelling neural correlates of working memory: a coordinate-based meta-analysis. *Neuroimage* 60, 830–846.
- Salimi-Khorshidi, G., Smith, S.M., Keltner, J.R., Wager, T.D., Nichols, T.E., 2009. Meta-analysis of neuroimaging data: a comparison of image-based and coordinate-based pooling of studies. *Neuroimage* 45, 810–823.
- Salomon, R., Levy, D.R., Malach, R., 2014. Deconstructing the default: cortical subdivision of the default mode/intrinsic system during self-related processing. *Hum. Brain Mapp.* 35, 1491–1502.
- Schacter, D.L., Addis, D.R., Buckner, R.L., 2007. Remembering the past to imagine the future: the prospective brain. *Nat. Rev. Neurosci.* 8, 657–661.
- Seeley, W.W., Menon, V., Schatzberg, A.F., Keller, J., Glover, G.H., Kenna, H., Reiss, A.L., Greicius, M.D., 2007. Dissociable intrinsic connectivity networks for salience processing and executive control. *J. Neurosci.* 27, 2349–2356.
- Seghier, M.L., Price, C.J., 2012. Functional heterogeneity within the default network during semantic processing and speech production. *Front. Psychol.* 3, 281.
- Sestieri, C., Corbetta, M., Romani, G.L., Shulman, G.L., 2011. Episodic memory retrieval, parietal cortex, and the default mode network: functional and topographic analyses. *J. Neurosci.* 31, 4407–4420.
- Sevinc, G., Spreng, R.N., 2014. Contextual and perceptual brain processes underlying moral cognition: a quantitative meta-analysis of moral reasoning and moral emotions. *PLoS One* 9 e87427.
- Shackman, A.J., Salomons, T.V., Slagter, H.A., Fox, A.S., Winter, J.J., Davidson, R.J., 2011. The integration of negative affect, pain and cognitive control in the cingulate cortex. *Nat. Rev. Neurosci.* 12, 154–167.
- Shulman, G.L., Fiez, J.A., Corbetta, M., Buckner, R.L., Miezin, F.M., Raichle, M.E., Petersen, S.E., 1997. Common blood flow changes across visual tasks: II. Decreases in cerebral cortex. *J. Cogn. Neurosci.* 9, 648–663.
- Simões-Franklin, C., Hester, R., Shpaner, M., Foxe, J.J., Garavan, H., 2010. Executive function and error detection: the effect of motivation on cingulate and ventral striatum activity. *Hum. Brain Mapp.* 31, 458–469.
- Slaughter, V., Peterson, C.C., Mackintosh, E., 2007. Mind what mother says: narrative input and theory of mind in typical children and those on the autism spectrum. *Child Dev.* 78, 839–858.
- Small, G.W., Moody, T.D., Siddarth, P., Bookheimer, S.Y., 2009. Your brain on Google: patterns of cerebral activation during internet searching. *Am. J. Geriatr. Psychiatry* 17, 116–126.
- Smallwood, J., 2013. Distinguishing how from why the mind wanders: a process-occurrence framework for self-generated mental activity. *Psychol. Bull.* 139, 519.
- Smallwood, J., Schooler, J.W., Turk, D.J., Cunningham, S.J., Burns, P., Macrae, C.N., 2011. Self-reflection and the temporal focus of the wandering mind. *Conscious. Cognit.* 20, 1120–1126.
- Smith, S.M., Fox, P.T., Miller, K.L., Glahn, D.C., Fox, P.M., Mackay, C.E., Filippini, N., Watkins, K.E., Toro, R., Laird, A.R., others, 2009. Correspondence of the brains functional architecture during activation and rest. *Proc. Natl. Acad. Sci. Unit. States Am.* 106, 13040–13045.
- Smith, S.M., Jenkinson, M., Woolrich, M.W., Beckmann, C.F., Behrens, T.E., Johansen-Berg, H., Bannister, P.R., De Luca, M., Drobnjak, I., Flitney, D.E., Niazy, R.K., 2004. Advances in functional and structural MR image analysis and implementation as FSL. *Neuroimage* 2004, S208–S219.
- Sneve, M.H., Magnussen, S., Alnæs, D., Endestad, T., D'Esposito, M., 2013. Top-down modulation from inferior frontal junction to FEFs and intraparietal sulcus during short-term memory for visual features. *J. Cogn. Neurosci.* 25, 1944–1956.
- Spaniol, J., Davidson, P.S.R., Kim, A.S.N., Han, H., Moscovitch, M., Grady, C.L., 2009. Event-related fMRI studies of episodic encoding and retrieval: meta-analyses using activation likelihood estimation. *Neuropsychologia* 47, 1765–1779.
- Spreng, R.N., Mar, R.A., Kim, A.S.N., 2009. The common neural basis of autobiographical memory, prospection, navigation, theory of mind, and the default mode: a quantitative meta-analysis. *J. Cogn. Neurosci.* 21, 489–510.
- Spreng, R.N., Andrews-Hanna, J.R., 2015. The Default Network and Social Cognition. *Brain Mapping: an Encyclopedic Reference*, pp. 165–169.
- Sylvester, C.C., Wager, T.D., Lacey, S.C., Hernandez, L., Nichols, T.E., Smith, E.E., Jonides, J., 2003. Switching attention and resolving interference: fMRI measures of executive functions. *Neuropsychologia* 41, 357–370.
- Thompson-Schill, S.L., Aguirre, G.K., Desposito, M., Farah, M.J., 1999. A neural basis for category and modality specificity of semantic knowledge. *Neuropsychologia* 37, 671–676.
- Toro, R., Fox, P.T., Paus, T., 2008. Functional coactivation map of the human brain. *Cerebr. Cortex* 18, 2553–2559.
- Turkeltaub, P.E., Eden, G.F., Jones, K.M., Zeffiro, T.A., 2002. Meta-analysis of the functional neuroanatomy of single-word reading: method and validation. *Neuroimage* 16, 765–780.
- Turkeltaub, P.E., Eickhoff, S.B., Laird, A.R., Fox, M., Wiener, M., Fox, P., 2012. Minimizing within-experiment and within-group effects in activation likelihood estimation meta-analyses. *Hum. Brain Mapp.* 33, 1–13.
- Uddin, L.Q., Clare Kelly, A.M., Biswal, B.B., Xavier Castellanos, F., Milham, M.P., 2009. Functional connectivity of default mode network components: correlation, anticorrelation, and causality. *Hum. Brain Mapp.* 30, 625–637.
- Uddin, L.Q., 2015. Salience processing and insular cortical function and dysfunction. *Nat. Rev. Neurosci.* 16, 55.
- Van Essen, D.C., Smith, S.M., Barch, D.M., Behrens, T.E.J., Yacoub, E., Ugurbil, K., Consortium, WU-Minn HCP Consortium, others, 2013. The Wu-Minn human connectome project: an overview. *Neuroimage* 80, 62–79.
- Van Overwalle, F., 2009. Social cognition and the brain: a meta-analysis. *Hum. Brain Mapp.* 30, 829–858.
- Varoquaux, G., Sadaghiani, S., Pinel, P., Kleinschmidt, A., Poline, J.B., Thirion, B., 2010. A group model for stable multi-subject ICA on fMRI datasets. *Neuroimage* 2010 (51), 288–299.
- Wager, T.D., Lindquist, M.A., Nichols, T.E., Kober, H., Van Snellenberg, J.X., 2009. Evaluating the consistency and specificity of neuroimaging data using meta-analysis. *Neuroimage* 45, S210–S221.
- Wager, T.D., Phan, K.L., Liberzon, I., Taylor, S.F., 2003. Valence, gender, and lateralization of functional brain anatomy in emotion: a meta-analysis of findings from neuroimaging. *Neuroimage* 19, 513–531.
- Wagner, A.D., Paré-Blagoev, E.J., Clark, J., Poldrack, R.A., 2001. Recovering meaning: left prefrontal cortex guides controlled semantic retrieval. *Neuron* 31, 329–338.
- Wagner, R.K., Torgesen, J.K., 1987. The nature of phonological processing and its causal role in the acquisition of reading skills. *Psychol. Bull.* 101, 192.
- Xu, B., Grafman, J., Gaillard, W.D., Ishii, K., Vega-Bermudez, F., Pietrini, P., Reeves-Tyer, P., DiCamillo, P., Theodore, W., 2001. Conjoint and extended neural networks for the computation of speech codes: the neural basis of selective impairment in reading words and pseudowords. *Cerebr. Cortex* 11, 267–277.
- Yarkoni, T., Poldrack, R.A., Nichols, T.E., Van Essen, D.C., Wager, T.D., 2011. Large-scale automated synthesis of human functional neuroimaging data. *Nat. Methods* 8, 665–670.
- Yeo, B.T.T., Krienen, F.M., Chee, M.W.L., Buckner, R.L., 2014. Estimates of segregation and overlap of functional connectivity networks in the human cerebral cortex. *Neuroimage* 88, 212–227.

- Yeo, B.T.T., Krienen, F.M., Eickhoff, S.B., Yaakub, S.N., Fox, P.T., Buckner, R.L., Asplund, C.L., Chee, M.W.L., 2015. Functional specialization and flexibility in human association cortex. *Cerebr. Cortex* 25, 3654–3672.
- Yeo, B.T.T., Krienen, F.M., Sepulcre, J., Sabuncu, M.R., Lashkari, D., Hollinshead, M., Roffman, J.L., Smoller, J.W., Zöllei, L., Polimeni, J.R., Fischl, B., 2011. The organization of the human cerebral cortex estimated by intrinsic functional connectivity. *J. Neurophysiol.* 106, 1125–1165.
- Zago, L., Pesenti, M., Mellet, E., Crivello, F., Mazoyer, B., Tzourio-Mazoyer, N., 2001. Neural correlates of simple and complex mental calculation. *Neuroimage* 13.
- Zanto, T.P., Rubens, M.T., Bollinger, J., Gazzaley, A., 2010. Top-down modulation of visual feature processing: the role of the inferior frontal junction. *Neuroimage* 53, 736–745.