

ESSENTIALS OF FUNCTIONAL MAGNETIC RESONANCE IMAGING

Tor D. Wager^{1*}
Martin A. Lindquist²

¹Columbia University, Department of Psychology

²Columbia University, Department of Statistics

Summary:

14741 words (main text)

1 table

10 figures

Draft of a chapter to appear in the Handbook of Social Neuroscience, J. Decety and J. Cacioppo, Editors.

Running head: FUNCTIONAL MRI

* Address correspondence to:

Dr. Tor D. Wager

Department of Psychology

Columbia University

1190 Amsterdam Ave.

New York, NY 10025

Phone: 212-854-5318

E-mail: tor@psych.columbia.edu

Acknowledgements

Parts of this chapter are adapted from: Wager, T. D., Hernandez, L., Jonides, J., & Lindquist, M. (2007). Elements of functional neuroimaging. In J. T. Cacioppo, L. G. Tassinary & G. G. Berntson (Eds.), *Handbook of Psychophysiology* (4th ed., pp. 19-55), Cambridge: Cambridge University Press. We would like to thank Dr. Doug Noll for providing Figure 5 for helpful comments on the manuscript, and Ted Yanagahara for creating Figure 1C.

In the last twenty years, neuroimaging techniques have opened up new frontiers in the study of the physical basis of the mind. As a result, there has been an expansion of the direct investigation of human brain activity in established disciplines such as cognitive psychology, personality and social psychology, psychophysiology, and psychiatry. In addition, new fields have emerged that blend existing approaches with neuroscientific techniques. For example, cognitive neuroscience, affective neuroscience, social cognitive neuroscience, and neuroeconomics all draw heavily on neuroimaging. In particular, functional magnetic resonance imaging (fMRI) is increasingly used as a primary research tool in these disciplines. It is particularly appealing because it is widely available to researchers' in psychological and other behavioral sciences and because it provides dynamic, second-by-second images of metabolic activity across the whole brain, with a balance of spatial and temporal resolution. Any effort to understand the brain substrates of psychology must involve converging evidence from multiple methods, and fMRI is an important tool in that it complements a number of other techniques. It complements lesion studies by providing images of dynamic activity changes in the intact, functioning brain. It complements electroencephalography (EEG) and magnetoencephalography (MEG) by providing enhanced ability to localize the sources of activity changes in the brain (Dale et al., 2000; Menon, Ford, Lim, Glover, & Pfefferbaum, 1997). It complements positron emission tomography (PET) studies by providing more information about the timing of brain activity changes in relation to task demands and behavior, and permitting analyses of individual subjects. It complements brain stimulation studies in humans by providing hypotheses about which regions play a necessary role in mental processes. Finally, it complements invasive electrophysiological studies in nonhuman animals by providing a molar, whole-brain view of areas likely to show altered electrophysiological activity. Currently, fMRI is the only technique that can provide whole-brain, second-by-second coverage of functional activity. In addition, it can be easily integrated with other MR-based methods, including imaging of brain structure and white-matter tract integrity, in the same imaging session. Images from each of these methods are shown in Figure 1.

This chapter is intended to provide an introductory overview of fMRI study design, analysis, and inference from start to finish. It complements previous, more detailed overviews (Wager, Hernandez, Jonides, & Lindquist, 2007; Wager, Lindquist, & Hernandez, in press) by providing a brief tour through the different stages of an fMRI study, highlighting key considerations at each step for those planning to conduct fMRI studies or just interested in reading and evaluating them. Each phase of an fMRI study, from conceiving the goals of the study to interpreting and describing the results, is like a link in a chain; the overall study is generally only as valuable to scientific progress as the weakest link. In Section 1, we discuss the kinds of questions that fMRI studies can productively answer. The ability to ask appropriate questions is the first "link" in obtaining useful results. We also address questions about the particulars of the study, such as appropriate sample size, and fMRI-specific experimental design. In Section 2, we discuss issues surrounding data acquisition, including how to choose scanning parameters that are appropriate for the specific scientific questions. In Section 3, we discuss fMRI data processing, and in Section 4 we discuss statistical analysis and inference. Finally, in Section 5, we discuss how to localize and interpret results, including which kinds of inferences are likely to be valid and which are not.

I. CONCEIVING AND DESIGNING AN FMRI EXPERIMENT

FMRI studies, like most other studies, start with a scientific question. As fMRI provides a measure of brain activity¹, questions that concern mind-brain relationships are most likely to be answerable. To ask whether there are different roles for different parts of prefrontal cortex in emotion generation, for example, is a challenging but addressable topic. To use fMRI (or PET, EEG, etc.) to ask what participants are thinking as they generate emotion, or what psychological processes are involved, is orders of magnitude more complex, and it is not clear that fMRI can provide much useful information on this point (Cacioppo & Tassinari, 1990; Coltheart, 2006; Poldrack, 2006; Sarter, Berntson, & Cacioppo, 1996), but those fMRI studies that do bear on psychological theory can be particularly powerful.

The ability to say something about the mind depends on having brain measures that are interpretable in terms of psychological processes. For example, part of the medial occipital cortex is often called “primary visual cortex” or V1, and because of a long history of research in neuroscience, activity in this region is often considered a marker for early visual sensory processes. Evidence for activity in V1 during mental imagery was critical in resolving a debate over whether imagery is a sensory process or a purely conceptual one (Kosslyn, Thompson, Kim, & Alpert, 1995). In this study, when participants imagined a larger letter, they showed more extensive activity in parts of V1 mapped to the peripheral visual field than when they imagined a smaller letter. Evidence from multiple modalities for attention effects in V1 and its thalamic inputs helped resolve a debate about whether attention involves “early” selection (O'Connor, Fukui, Pinsk, & Kastner, 2002; Roelfsema, Lamme, & Spekreijse, 1998). At this stage, plausible biological markers of psychological processes are extremely rare, and developing them is a major goal for the field of neuroimaging.

I.A. Goals of an fMRI study

Functional MRI studies are performed for a great variety of reasons, some geared towards basic understanding of brain-behavioral relationships and others for particular practical applications. However, they tend to fall within several broad families described in detail below. In general, it is useful to explicitly define the goals of your study, so that you can evaluate whether your design and implementation is likely to serve them at each step in the process.

¹ Brain “activity” or “activation” is typically used to refer to Blood Oxygen Level Dependent (BOLD) signal, a complex set of changes usually coupled to the ratio of oxygen demand and blood flow in local blood vessels. A number of more detailed descriptions of BOLD physiology have been published elsewhere (Buxton, Wong, & Frank, 1998; Etkin & Wager, 2007; Logothetis, Pauls, Augath, Trinath, & Oeltermann, 2001), so we will not cover that in detail here. In addition to BOLD, perfusion imaging can be used to obtain more direct measures of local cerebral blood flow (CBF), and spectroscopy and other non-BOLD techniques can provide other types of measures. However, BOLD signal is relatively robust and is almost universally used in psychological studies.

Structure-function mapping.

The most common use of fMRI to date is *structure-function mapping*, which is best defined as an exploratory, or hypothesis building, approach. “Brain mapping” studies are necessary for the development of biological markers and understanding the boundaries of the psychological categories they respect. This process is largely exploratory when an area of study is new, and few precise hypotheses are available, and moves towards testing of particular alternatives in particular regions or networks as evidence accumulates across many studies.

Process dissociation and association

A second use of fMRI is for *process dissociation or association*. The question of which mental processes are similar to one another, and which are different, is central to our understanding of the organization of the mind. Neuropsychologists and neuroscientists have long sought to “carve nature at its joints,” and identify the key types of mental processes and tasks that are of the same kind. Categorization of mental abilities has also been central in the field of psychometrics and aptitude testing. The basic logic in both areas is that performance measures that are inter-correlated share some common mechanisms, or capacity limitations, that are not shared by tasks in other categories. Neuroimaging activation has been used as a “molar” measure of association, much as correlations are used in psychometric testing: If two or more tasks produce dissociable activation patterns, they must have different mechanisms.

Dissociation occurs when a brain region is more active in Task A than Task B. A double dissociation occurs when each task activates one region more than the other task. Double dissociations are a powerful tool because they imply that the two tasks utilize different processes, and that one task is not a subset of the other. For example, in one study in our laboratory, we found that different types of task switching, or switching attention from one feature or object to another, differentially activate a set of regions thought to be involved in the control of attention (Wager, Jonides, Smith, & Nichols, 2005). Four types of switches were dissociable—each produced higher brain activity in some regions than the others—paralleling behavioral findings that performance switch costs are more highly correlated for similar types of switches than dissimilar types (Wager, Jonides, & Smith, 2006).

There is no perfect method for making inferences about mental structure, and dissociation logic is no exception. While we believe that it is a useful tool, we note two important caveats. First, double dissociations in neuroimaging data can be produced in trivial ways. Imagine a study of two “tasks:” (1) perceiving an intensely painful stimulus, and (2) perceiving a moderately painful stimulus. In this thought experiment, we will assume that (2) is a subset of (1), so there are no “different mechanisms.” Task 1 is likely to produce more activation in the insula and cingulate cortex, thereby producing a dissociation. It is also likely to produce greater de-activation in visual sensory cortex, so that Task 2 shows greater activity in these areas. In this case, we have observed a double dissociation, but not because there were unique, “low-pain” processes. Additional controls are necessary. Stronger evidence for unique processes would be provided if we could demonstrate ‘separate modifiability:’ regions activated in each task that are not affected at all by the other task (Sternberg, 2001).

The second important caveat is that questions about similarity and dissociation are always relative to the spatial sensitivity of the measure. Typical BOLD studies use “voxels” of about 27 mm³, which means that signal is

averaged across several million neurons at best. No matter where a voxel is in the brain, different neurons within this population are almost sure to encode different stimulus properties and respond to different tasks, as even a cursory look at the primate electrophysiology literature will reveal. Thus, studies that activate the “same” areas only do so at a particular spatial resolution. If physical pain and “pain empathy” activate overlapping voxels in the cingulate and insula, they engage the “same” processes only insofar as there is a single, general role for these structures across tasks—it does not mean that they activate the same neurons. And even if the same neurons are active, they may play different functional roles in different tasks because they participate in different networks. In one study, for example, we found that physical pain and pain empathy produced overlapping activation of the cingulate and insula, but that they could be distinguished by the connectivity of these regions with brainstem areas (for physical pain) or temporal and medial frontal regions (for pain empathy).

Prediction and psychological inference

A third class of goals concerns the use of fMRI data for a variety of practical and scientific purposes: Diagnosing early-onset Alzheimer’s, identifying subgroups of patients who will respond differently to treatment, predicting words, thoughts, and intentions from patterns of brain activity, and detecting deception are but a few examples. The incorporation of fMRI-based measures into medical diagnosis and treatment, law, and policy decisions at many levels is in its infancy. Though the potential is clearly evident, great care must be taken not to over generalize a limited set of results to broader contexts without rigorous validation.

A core issue is that fMRI is inherently suited to making inferences about brain responses given that participants are engaging in a particular task state that varies over time. It can say little about the reverse inference, inference about the psychological state given a particular pattern of brain response. For example, imagine a participant performing an ideal lie-detection task in an ideal scanner. Because conditions are ideal, the insula is activated when the participant is lying with 100% reliability. Would we then be able to use insula activity as a lie detector, to infer when the participant was lying? Not necessarily. Frustration, or anxiety, may also activate the insula—and as long as *anything besides lying* activates the insula, there is a plausible alternative inference. Formal analysis techniques geared towards prediction provide quantitative estimates of the accuracy of such “reverse” psychological inference, and this is an exciting area of rapid development in fMRI-based research. However, it is important to note that even in the ideal case, these techniques can only provide evidence that distinguishes a process from a limited set of well-defined alternatives. Thus, it might be possible to predict whether a given image is “familiar” or “unfamiliar” (the basis of many common tests of deception) given only those two alternatives, but it is much more difficult to rule out other, less well-specified alternatives—“unfamiliar but anxiety-provoking,” “unfamiliar but salient”—and to detect and avoid the use of counter-measures. Thus, the prospect of using fMRI for prediction and psychological inference is both an exciting and challenging one.

I.B. Design structure: Kinds of designs

Most published studies to date either focus on structure-function mapping or process association/dissociation, either across the brain or in specific anatomical regions of interest. The most common way of doing

this is through a “subtraction design.” Essentially, researchers construct two (or more) task conditions that are matched, ideally, on all sensory and psychological features except one process of interest. If we denote the experimental and control conditions as Conditions A and B, respectively, the subtraction contrast $[A - B]$ will isolate the process of interest. It is rarely possible to match features perfectly, and in this case areas showing a significant $[A - B]$ difference will have multiple interpretations. For example, some of our work compares moderately intense thermal pain with non-painful warmth (Wager, Scott, & Zubieta, 2007). The subtraction contrast identifies brain regions that respond to and encode the sensation of pain, but likely also includes areas involved in attentional orienting and action tendencies. The usefulness of this subtraction depends on the goals of the study. A number of variants expand on the subtraction concept, generally with the goal of permitting more specific inference.

In *parametric variation* designs, fMRI is used to test whether brain activity increases linearly or monotonically (or fits another expected pattern) as a psychological process is incrementally varied (Buchel, Holmes, Rees, & Friston, 1998). Figure 2 shows some examples. Such designs can provide more convincing inferences because they incorporate more specific predictions about a particular process. Parametric “modulators” can be task conditions such as task difficulty (Figure 2, left)(Dagher, Owen, Boecker, & Brooks, 1999) or time (Buchel, Morris, Dolan, & Friston, 1998), or measured variables such as reaction time (Grinband, Hirsch, & Ferrera, 2006), subjective value (Figure 2, right)(Hare, Camerer, & Rangel, 2009), or psychological model-derived processes such as prediction errors (McClure, Berns, & Montague, 2003). With *factorial designs*, multiple types of events are included in the design, and main effects and interactions between factors are tested. Factorial designs are particularly compatible with process association and dissociation inferences, described above, because they independently vary two or more factors that may influence brain activity. Finally, *multiple subtraction designs* compare a task with multiple control conditions, each designed to control for a separate, potentially confounded process (Price & Friston, 1997). The underlying goal of these techniques is essentially the same: to characterize which brain areas predict/respond to/differentiate one set of task conditions from another. The vast majority of studies use these types to provide more specific inferences in process association/dissociation designs, but they can do little to circumvent the limitation posed by the spatial resolution of fMRI, which integrates signal across millions of neurons into a single voxel.

A growing trend in the fMRI literature is to try to circumvent some of the spatial limitations by using two kinds of techniques: *fMRI adaptation designs* and *pattern classification-compatible designs*. The logic of *fMRI adaptation* studies is that repetition of a particular stimulus often produces neural “fatigue,” in the sense that the second repetition will produce a smaller response (Grill-Spector & Malach, 2001). The logic is that if two trial types, A and B, activate the same neurons, then there should be cross-trial habituation, i.e. the response to B when it follows A should be reduced. If they activate different neurons, there should be no cross-trial habituation. A full discussion of these designs is beyond the scope of this chapter, but briefly, it is important to consider three things: First, vascular responses show nonlinear “habituation” independent of neural activity or mental processes (Vazquez et al., 2006; Wager, Vazquez, Hernandez, & Noll, 2005). Second, even if two stimuli activate the *same neuron*, they

can produce different degrees of habituation (Sawamura, Orban, & Vogels, 2006). Finally, though interpretations of fMRI-adaptation effects are often cast in terms of neuronal firing, global processes related to memory may play an important role as well (Henson, 2003).

Pattern classification is an analysis technique that is naturally suited to prediction and reverse inference, and many dozens of studies now employ these techniques to predict a variety of psychological states or events from patterns of fMRI data (Mitchell et al., 2004). Predictive techniques are also increasingly being used to address questions about structure-function mapping and process overlap through **information-based mapping** (Haynes & Rees, 2006). The logic is that even if two tasks activate the same gross anatomical region (e.g., the insula), the activation may reflect different populations of inter-digitated neurons. If the spatial distribution of “clumps” of neurons that respond more to one task than the other are suitably distinct and the spatial resolution of the fMRI study is high enough, then the two tasks should produce different patterns of relative activity across local voxels. If so, then the machine learning techniques discussed above can be used to test association or dissociation. An algorithm might be trained to discriminate between Task and Control conditions for Task A (i.e., [Task A – Control A]). If the predictive map that results can subsequently be used to discriminate Task vs. Control for Task B, there is evidence for shared processes at a finer spatial scale.

To close this section, let’s return to the main goal of dissociation/association logic, which is to test whether mental processes share common and/or distinct neural substrates. Ultimately, fMRI can only provide inferences about brain signals, and their relationship to mental processes or mechanisms requires an additional conceptual leap, equating the similarity of *brain* mechanisms to similarity of the underlying *mental* process. Our study of physical pain and pain empathy highlights one aspect of this issue. Even if the same neurons respond to physical and social pain, if these neurons are part of different functional networks, are these shared neurons part of the same mental process? Imagine a conference call with some work colleagues. You are a “neuron,” and together you form a “network” that accomplishes a particular task. Now imagine that after this call, you have another conference call with some old college friends. You are involved in both; does that mean that the two “networks” accomplish, even to a small degree, the same task? Not necessarily.

The initial stages of any study involve grappling with these issues, and making critical decisions about what questions you will test and what fMRI will be able to tell you about your questions. Because fMRI can generally only provide *relative* measures activity differences across conditions (e.g., Task A vs. Task B), you will also need to decide which task or stimulus conditions to include and what specific comparisons, or **contrasts** across those conditions, you want to test. Once you have made some preliminary decisions on these issues, you are ready to implement a detailed design for your fMRI study. This is the subject of the next section.

I.C. Practical fMRI design

There are many considerations in designing an fMRI study, and they may vary with the study goals. How many subjects should be run, and for how long? Should stimuli or events be *blocked* together in time, or randomized in an *event-related* design? Here, we provide an overview of some of the general considerations. Some

principles of experimental design are not unique to fMRI—for example, it is always critical to avoid confounding the comparisons you are interested in with effects of time, practice, fatigue, task history, and other factors.

Functional MRI is uniquely challenging in this respect, however, because unlike traditional psychological outcome measures, fMRI activity is potentially sensitive to *any* differences among task conditions in what the participant sees, hears, thinks, or imagines.

How many subjects, and for how long? Power and sample size

It is always important to design a study that has a reasonable chance of detecting an effect—i.e., that is adequately powered. How many subjects are needed and how long should each be scanned? This seemingly simple question is actually quite complex. There is a tradeoff between scanning more subjects and the fMRI time spent per subject. The optimal balance depends on the relative balance of within-subjects noise (error) and between-subjects noise (individual differences) in each region of interest in the brain. The within-subjects effect size can be decomposed into multiple factors: (1) The effect size of the “neural” effect, including the neural effect magnitude, time on task, fatigue and habituation; (2) the inherent scanner noise, considering autocorrelation (slow drift) and physiological noise; (3) the number of independent data points collected, considering that sequential fMRI scans are non-independent (autocorrelated); (4) the accuracy of the model of neural signal and BOLD hemodynamics used in the analysis; and (5) correlations in the predicted response among conditions of interest. The between-subjects effect size depends on: (1) the within-subjects effect size; (2) the magnitude of individual differences in the comparison of interest; (3) the magnitude of individual differences in the BOLD response, including in the magnitude, shape, and fit to the assumed hemodynamic model; and (4) the accuracy of inter-subject registration (i.e., normalization or warping) or other method for matching brain areas across subjects for analysis. The bottom line is this: The greater the within-subjects noise, the more likely it is that it will pay off to scan fewer subjects for longer. The greater the between-subjects noise, the more likely it is that your money will be best spent on additional subjects.

It is a good idea to think about each of these issues when designing and analyzing an fMRI study. However, it can be daunting to attempt to consider each of these issues in detail, as information on many of them may not be obtainable until *after* you’ve run the study. The best way to determine how many subjects to run, and consequently how much money and time you will have to spend to run the study, is to estimate the within-subjects and between-subjects variation from existing data and perform a power analysis (Mumford & Nichols, 2008). If you can estimate the within- and between-subjects error, then the reference curves in Figure 3 can provide a guide to optimal choices (see the figure legend for more details). However, often these estimates are not available. As a rule of thumb, if a group analysis is the goal, scan subjects for at least 10 minutes in cognitive tasks. If additional time is free, scan them for longer—but if it is costly, the time and money is usually best spent scanning more subjects. If you care about obtaining significance maps within individual subjects, instead scan fewer subjects for longer. Two other points are worth mentioning in connection with the basics of designing a study. First, subjects get tired, they learn, and they tend to get sleepy (believe it or not!) in the scanner after 30 minutes or so of lying still and listening to the repetitive bleep-bleep-bleep of the scanner. After 60 minutes or so, many participants will

experience pain on pressure points, and they will move more than usual. These factors should be considered, and provide additional reasons to limit functional time to reasonably short periods. Second, it is a good idea to build positive controls into your task design—comparisons that are known to strongly activate particular brain regions, such as early visual cortex for [visual stimulation – rest] or motor cortex activation for [movement – rest]. These comparisons can serve as benchmarks for effect sizes in other studies and, if the expected results are not found, signal basic errors in data processing and analysis that might otherwise go undetected. Positive controls are a standard part of experimental research in biology, but sadly, too few functional imaging studies employ them. Those that do are usually more compelling.

Powerful vs. informative designs: making tradeoffs

Now, finally, we turn to the nuts and bolts of experimental design: Choosing the task conditions, when and for how long to present them, and what comparisons to make. The first neuroimaging designs, and still a workhorse, were ***block designs***. An example is shown in the top panel of Figure 4A. Here stimuli (or task conditions) of the same type are grouped together in time, and stimulation/performance within a block is typically as continuous as possible, to maximize the time on the task of interest. If stimulation can be continuous, blocking stimuli does not change the nature of the task in undesirable ways, and a design with two block types (e.g., visual stimulation and rest) alternated every 16-20 sec is maximally powerful (Liu, 2004; Skudlarski, Constable, & Gore, 1999; Wager & Nichols, 2003). In addition, they are relatively robust to mis-specification of the shape of the neural input and hemodynamic response functions (HRF), as we describe in more detail below.

Block designs with more conditions (including parametric variation and factorial designs) may be more informative, in that there are more conditions to compare and potentially more leverage on how to interpret activations, but they are also less powerful. There are two reasons for the loss in power. First, less time is spent imaging any one condition, so activation estimates will be less stable. Second, they cause alternations across blocks to occur more slowly in time. The noise in fMRI is high at low frequencies, meaning that the signal drifts slowly over time. High-pass filters are typically used to remove noise, but they cannot be used if the conditions of interest also vary slowly. Because a 16 sec alternation frequency closely matches the frequency of the hemodynamic response in most areas of the brain, it is nearly perfectly efficient. Thus, all things being equal, a two-condition block design with 16 sec blocks is the most powerful and robust design.

Event-related designs (Figure 4A, bottom) are a popular alternative, for several reasons. First, not all designs can be implemented as block designs without changing the nature of the task. For example, in a stop-signal task, the subject is asked to respond as fast as possible to a visual stimulus, but stop if an auditory tone is presented just after stimulus onset. Blocking “stop” trials would dramatically change how difficult it is to stop, because participants could anticipate and adapt to the trial structure. Second, event-related designs are more interpretable in one important, fundamental sense: They allow activity to be linked with more credibility to individual event types. Activation in block designs could be a “state” effect that occurs throughout a block, in between trials, or only on certain trials (e.g., errors), making interpretation more difficult. Event-related designs can isolate trials or even parts of trials (Ollinger, Shulman, & Corbetta, 2001). They can also permit analysis of activity modulated by

performance variables such as reaction time on a trial-by-trial basis, as discussed above, and estimation of the shape and timing of the HRF in response to each event type in each voxel (Bellgowan, Saad, & Bandettini, 2003). As with block designs, the most statistically powerful designs involve only two event types, in order to maximize time on task and minimize low-frequency variation in the design. A third reason why simpler event-related designs are more powerful is that fewer event types means less *multicollinearity*, a standard but largely insurmountable problem in statistical analysis that can dramatically reduce power and lead to unstable (even when “significant”) effects. Multicollinearity occurs when the linear model regressor (the predicted fMRI signal) for one event type can be predicted by a combination of the other regressors. Fewer event types mean, on average, less multicollinearity. Of course, designs with more event types may offer more interesting psychological inferences, if they are appropriately powered.

As mentioned above, event-related designs are also less robust. This is because the BOLD response to brief neural events is not an instantaneous one; it is determined by slow changes in blood oxygenation, flow, and volume, as illustrated in Figure 5. A brief burst of neural firing typically produces a slow hemodynamic response (HR) that peaks about 5-6 sec post-stimulus and returns to baseline after 20-30 sec, as shown in Figure 6. As discussed in the Analysis section below, fMRI analyses typically assume a fixed shape for the HR and use it to develop predictors for what BOLD responses to various types of events are expected to look like (shown in Figure 6B). These predictors are fit to the observed fMRI time series from each voxel in each subject. One problem is that the hemodynamic response varies across brain regions, subjects, and tasks (Aguirre, Zarahn, & D'Esposito, 1998; Schacter, Buckner, Koutstaal, Dale, & Rosen, 1997; Summerfield et al., 2006). Event-related designs are sensitive to the difference between the assumed and actual hemodynamic responses, whereas block designs are less sensitive.

The slow HR is an important factor that shapes other elements of fMRI design as well. It acts like a low-pass filter, smoothing out rapid changes in the signal. If you present events too close together, the HR will blur the rapid rise and fall in your experimental design into what looks like a flat line. Because the linear modeling approach depends on this rise and fall—it fits the shapes of the regressors to the shape of the rise and fall in the data—this blurring dramatically decreases power. On the other hand, if you present events too far apart, you create little variation in the regressors (and the data) across time because very little time is spent on-task. Thus, optimal designs are those that present events fairly close together. Alternating two event types every four seconds or so in a randomized fashion is optimal for a randomized event-related design.

The four-second rule of thumb takes into consideration another important factor: nonlinearity. In building predictors of fMRI signal, it is typically assumed that each event produces a discrete response that the same for all events, and that the total signal is the sum of responses to individual events. If this were strictly true, you could present 1000 spider pictures in 1 second and get a response 1000 times as large as if you had presented only one photo! Obviously, the response is not linear. In fact, the largest sources of nonlinearity are neural habituation and vascular saturation. Neurons often stop responding as strongly after a few hundred milliseconds (and some respond only to stimulus onset or offset), and blood vessels dilate and lose elasticity. The effect of both is that the BOLD

response reaches a ceiling as stimulation intensity or frequency increases (Vazquez & Noll, 1998)(Birn, Saad, & Bandettini, 2001; Friston, Mechelli, Turner, & Price, 2000; Wager, Vazquez et al., 2005). Four seconds delay between brief events is enough to prevent most nonlinear interactions between successive events (Miezin, Maccotta, Ollinger, Petersen, & Buckner, 2000).

Rest intervals and ‘jitter’. Standard BOLD fMRI can provide comparisons of relative activity between conditions, but it cannot provide absolute measures of blood flow. When designing an experiment, it is important to plan for all the comparisons that you want to make. This includes comparisons with rest or a comparable ‘low-level’ control condition. Let’s say you’re conducting an experiment comparing activity during viewing of faces and scenes (“places”). You’d like to compare [Faces – Places] and identify face-specific brain regions. But you would also like to know whether your face specific regions *increase* in response to faces. Some regions that show a [Face – Place] difference might do so because they *decrease* in response to places. To test for increases during face processing, you must test for increases relative to some other, third control condition, i.e., [Face – Rest]. If you are interested in this comparison, your design must contain rest intervals, so that time is spent sampling brain activity in this control condition. How much time you spend sampling the rest condition depends on how much you care about [Face – Rest] vs. [Face – Place]. An optimal design for the former involves spending as much time sampling face presentation and rest as possible, with equal allocation of time between them. For the latter, an optimal design involves heavy and balanced sampling of face and place presentation. Thus, the experimental design should consider the relative importance of different comparisons in your study.

In an event-related design, rest intervals are often included between events in the form of an inter-trial interval or inter-stimulus interval (ITI or ISI). Rapid and regular alternations between stimuli and rest will not provide an efficient design, for the same reason that rapid alternation between two event types (e.g., Face and Place) will not be efficient; the BOLD HR will essentially smooth out the differences between conditions and make it impossible to detect anything. For this reason, event-related designs usually use variable ITIs or temporal ‘jitter.’ How much jitter is appropriate? A rule of thumb, shown to be near-optimal, is to use short ITIs (e.g., 4 sec) for 50% the trials, and exponentially decrease the frequency as a function of ITI; thus, 25% of the trials might use a 6 sec ITI, 12.5% might use an 8 sec ITI, and so on.

Another important consideration is which type of low-level control condition to use. Many researchers believe that the ‘resting state’ is a privileged state against which to compare functional activation (Raichle et al., 2001). However, other researchers point out that rest involves active processes that likely include self-monitoring and memory retrieval (among others), and that individuals may vary dramatically in which processes they engage at rest. Stark and Squire (Stark & Squire, 2001), for example, found that activity in the medial temporal lobes was substantially higher during rest than during some low-level cognitive tasks. The memory tasks they used ‘activated’ relative to a low-level cognitive control condition, but ‘de-activated’ relative to rest. Thus, it may sometimes be more appropriate to compare tasks to low-level baseline tasks during which mental activity can be more precisely experimentally controlled (Johnson et al., 2005). There is no universally accepted “right answer.”

Optimized experimental designs. What constitutes an optimal experimental design depends on the

psychological nature of the task, as well as on the specific comparisons (contrasts) of interest in the study. In fMRI studies, the delay and shape of the BOLD response, scanner drift and nuisance factors such as physiological noise, and other factors also play important roles in design optimality. Not all designs with the same number of trials of a given set of conditions are equal, and the spacing and ordering of events is critical (Josephs & Henson, 1999; Liu, 2004; Smith, Jenkinson, Beckmann, Miller, & Woolrich, 2007; Wager & Nichols, 2003).

Several computer algorithms are available for constructing statistically optimized designs (e.g., T. D. Wager & Nichols, 2003), and they are particularly useful for event-related designs. These programs generally operate on the spacing, frequency, and order of stimulus presentation during the experiment. They incorporate assumptions about the HRF shape and linearity or non-linearity of the response, information about the comparisons of interest in the study, and other design and analysis choices such as the frequency of image collection and data filtering options. One important point is that there is a substantial tradeoff in power between *detecting* activation differences between conditions (i.e., an $[A - B]$ contrast) using an assumed HRF shape and estimating the *shape* of evoked activations with a more flexible model (Liu, Frank, Wong, & Buxton, 2001). For detecting activation differences, a block design alternating between two blocks (A and B) with a block length of 16 sec each is statistically most efficient. For HRF shape estimation, a useful toolbox for optimizing the event ordering is based on M-sequences, or sequences which are orthogonal to themselves shifted in time (Buracas & Boynton, 2002). Random event-related designs fall somewhere in between blocked and event-related designs, but are non-optimal on both. Computer-aided designs can produce substantially better results than random event-related designs on both measures.

II. ACQUIRING DATA

Functional MRI is a multidisciplinary venture, and considerable physics and engineering knowledge goes into the data acquisition. Thus, it is a good idea to make friends with a physicist or two and send them birthday cards (though we note that Dr. Doug Noll graciously provided us with Figure 5 without one). Image acquisition is rapidly developing, and cutting-edge research continues to push towards better image quality, lower noise and artifacts, higher temporal and spatial resolution, and, as a consequence, more meaningful results. Commercial manufacturers have increasingly made stable fMRI acquisition protocols available, and these standard acquisition procedures may be adequate for many studies, but often even seemingly “routine” imaging requires careful selection of acquisition parameters. How many slices will you collect? In what orientation? How large will the voxels be? How rapidly will images be acquired? Will you be able to obtain signal in key brain areas such as the amygdala, orbitofrontal cortex, and medial temporal cortices? This section is designed to help you to be able to talk to your physicist collaborator in order to design an acquisition scheme that is suitable for your study. Just as with the experimental design, the key to setting up an appropriate data acquisition protocol lies in understanding some basic tradeoffs and considering them in relation to the goals of your particular study. We first provide a bit of background, and then some considerations on these choices. In the basic exposition below, terms with specific meaning in MR physics are in bold italics.

MR basics

Both structural (*T1-* or *T2-weighted*) and functional (*T2*/susceptibility weighted* or *perfusion-weighted*) MRI images are obtained using the same scanner, though it is programmed differently depending on the image type. In all cases, a subject is placed in a strong magnetic field, usually 1.5 to 7 Tesla, generated by a coil of wire wrapped around a tube. The person being imaged lies in the center of the tube, or the **bore**. The field is *always on*, and anything attracted to magnets can be pulled into the bore of the magnet with irresistible force, even when the imaging control console is shut down. Therefore, it is extremely important not to bring metal objects, or anyone who might have metal in their body, near the scanner.

During imaging, the participant is exposed to a radiofrequency (RF) electromagnetic field pulse delivered through a second coil (commonly called “*the head coil*”) that typically surrounds the participant’s head. Unlike PET, this does not involve exposure to radioactive material, and there are no known risks of repeating this procedure. The main magnetic field (B_0 , “*b-naught*”) aligns the spins of protons in hydrogen atoms (“*the spins*”) in the brain along its axis. When the RF pulse is delivered, these protons absorb the energy at a very specific frequency band that depends on the field strength and become *excited*. The nuclei then emit the energy at the same frequency as they *relax*. Different tissue types have different *relaxation times*, which creates the contrast between gray matter, white matter, and cerebrospinal fluid (CSF) in the images.

There are several kinds of relaxation that are used to create contrast in images: T_1 , T_2 , and T_2^* . These letters are used to describe the relaxation rate associated with each. T_1 is the rate at which spins relax back to the main magnetic field. T_1 -weighted images have values across the image sensitive to the differences in T_1 across gray matter, white matter, and CSF, and they provide excellent detailed images of brain anatomy. T_2 refers to how quickly the energy applied from the RF pulse dissipates. T_2 -weighted images also provide excellent anatomical structure, and additional detail in some subcortical and brainstem nuclei. Regions such as the substantia nigra, with high iron content (which is magnetic) have very fast T_2 relaxation and appear as dark spots in T_2 -weighted images, providing anatomical definition in these regions. Finally, T_2^* the rate of attenuation of the magnetic field applied by the RF pulse. T_2^* depends additionally on local inhomogeneities in magnetic susceptibility that are caused by changes in blood flow and oxygenation.

As the spins relax, the returned energy is detected by the *receiver coil*, generally the same coil that produced the RF field. This energy is detected as a one-dimensional series of fluctuations over time; however, due to some brilliant work that resulted in a Nobel Prize to Paul Lauterbur and Peter Mansfield, it is possible to reconstruct a three-dimensional image from these signals. This is done by applying *gradients*, magnets that change the strength of the magnetic field in systematic ways across space so that the frequency and the phase of the signals detected by the receive coil encode the location of the signal in the brain. *Pulse sequences*, or software programs that implement particular patterns of RF and gradient manipulations, are designed to acquire data that can be used to reconstruct images of the brain. There are various types of pulse sequences, including *echo-planar imaging*

(EPI) and *spiral imaging*. With a standard EPI sequence, 3-D localization is done by selectively exciting tissue in one 2-D slice of brain at a time (*slice selection*) and applying gradients to encode the location of the signal within the 2-D slice. It should be noted that the MR signal is not acquired voxel-by-voxel; rather the acquired data is collected in the frequency domain, or *k-space*. The data can be converted from k-space to *image space* (an image that looks like a brain) by applying a *Fourier transform*.

BOLD signal

T_2^* -weighted imaging is used to obtain measures of regional brain activity. The most popular method uses the Blood Oxygenation Level Dependent (BOLD) signal (Kwong et al., 1992; Ogawa et al., 1992), which takes advantage of the difference in T_2^* between oxygenated and deoxygenated hemoglobin. As neural activity increases, so does metabolic demand for oxygen and nutrients. As oxygen is extracted from the blood, the hemoglobin becomes paramagnetic, which creates small distortions in the B_0 field that cause a T_2^* decrease (i.e. a faster decay of the signal). Increases in deoxyhemoglobin can lead to a decrease in BOLD signal, often referred to as the “initial dip” (Figure 6A) The initial decrease in signal is followed by an increase, due to an over-compensation in blood flow that tips the balance towards oxygenated hemoglobin (and less signal loss due to dephasing). This ultimately leads to a higher BOLD signal that gives the positive rise that characterizes the hemodynamic response.

Arterial Spin Labeling (ASL)

ASL is a family of techniques that provides an alternative to BOLD imaging (Williams, Detre, Leigh, & Koretsky, 1992). ASL uses pulse sequences sensitive to blood volume or cerebral perfusion (or cerebral blood flow, CBF), a quantifiable physiological measure that is more directly related to neuronal metabolism than BOLD. In typical ASL experiments, two images are collected: one following a tagging period, and the other after a control period during which the tag clears and the system is allowed to return to equilibrium. Subtraction of these two images yields a signal that is proportional to the amount of tag that is in the tissue, and can be used to quantify the perfusion rate. ASL may be used to quantify and study baseline cerebral blood flow without comparison to an active state, which is desirable for many kinds of studies of special populations (aging populations, children, those with schizophrenia, mood disorders, etc.). Because the measurements are stable over time, ASL can be used to compare conditions across long time periods, facilitating the study of phenomena such as emotional states, drug effects, and learning.

Study-specific acquisition choices

There are a number of ways to optimize fMRI acquisition to maximize the chances of achieving the goals of a particular study, but each option comes with associated costs, as summarized in Table 1. Below, we discuss several options and their costs and benefits. If you have access to multiple scanners, one choice may involve which *field strength* to use, typically either 1.5 T or 3.0 T. Higher field strength means lower thermal noise, the inherent noise in the excitation and measurement of spins, and a higher BOLD contrast to noise ratio (CNR). But higher fields are

not always better, because *susceptibility artifacts* are greater at higher field strengths. Susceptibility artifacts are decreases and/or spatial distortions in the signal that occur due to local inhomogeneity in the magnetic field, often at the boundaries between tissue and sinus spaces at the bottom of the brain, including orbitofrontal cortex, ventromedial prefrontal cortex, amygdala, and the medial temporal lobes. The T2*-weighted imaging sequence produces BOLD contrast by being sensitive to inhomogeneity, so susceptibility artifacts are difficult to avoid in BOLD imaging. Minimizing susceptibility artifacts is an important consideration if signal in these areas is of interest. Some approaches are to use “z-shimming” during acquisition to make the field more homogenous (Constable & Spencer, 1999), improved reconstruction algorithms (Noll, Fessler, & Sutton, 2005), “unwarping” algorithms that measure EPI distortion by measuring inhomogeneity in the magnetic field and attempting to correct it (Andersson, Hutton, Ashburner, Turner, & Friston, 2001), and use of magnetic inserts in the mouths of participants that act as shims (Wilson & Jezzard, 2003). However, choices of standard acquisition parameters can also have a great impact on susceptibility and other signal characteristics. We describe some of these choices, and their impact, below.

One way of minimizing susceptibility is to shorten the TE, the *echo time* between delivery of the RF excitation pulse and reading out the signal. There is an optimal TE for T2* sensitivity (typically about 40 msec for 1.5 T, and 25 msec for 3.0 T), and it grows shorter as the field strength increases, because signal loss (*dephasing*) due to inhomogeneity occurs at a faster rate. Shortening the TE, say to 25 msec at 1.5 T, will decrease susceptibility, but also decrease BOLD contrast and thus the BOLD signal to noise ratio. For high-susceptibility areas such as the amygdala and orbitofrontal cortex, this is likely to be advantageous, but it will reduce sensitivity for other brain areas. This tradeoff, and others described below, must be evaluated with respect to the particular goals of your study.

Another way to reduce susceptibility is to scan using smaller voxels, and, in particular, thinner slices. Current typical voxel sizes are 3 x 3 x 3 mm. Collecting thinner slices will reduce the inhomogeneity gradient across space within a voxel, and lower inhomogeneity means less susceptibility artifact. This is obviously advantageous for increasing the spatial resolution, which improves localization ability. However, it comes at a cost to CNR, which decreases proportional to the voxel volume, and thereby reduces sensitivity to experimental effects. Thus, the CNR for a 1 x 1 x 1 voxel acquisition scheme is 27 times lower than the standard 3 mm acquisition. Another problem with smaller voxels is that it takes longer to acquire an image across the whole brain. Researchers interested in high-resolution fMRI studies may opt for increased spatial resolution and reduced susceptibility, but these benefits come at a cost to statistical power and brain coverage.

These tradeoffs may be less extreme if the repetition time (TR) to acquire images is lengthened. A typical TR is on the order of 2 sec, meaning that an image volume is acquired every 2 sec. Longer TRs decrease the temporal resolution, but permit more brain coverage (more slices) with smaller voxel sizes. For TRs in the range of 0.5 – 2.5 sec, the cost of reduced temporal resolution is minimal, as the BOLD hemodynamic response lag is on the order of 6 sec. However, there is a danger in event-related designs: If trials are locked to the TR, the peak of the BOLD response will fall somewhere in between the samples. This issue can be avoided by ‘jittering’ the events

with respect to image acquisition. Other problems with longer TRs, however, are harder to overcome. A chief one is that even small head movements (or brain movement due to breathing and heart-beat) during the acquisition of a volume can induce inhomogeneity that creates large artifacts in the data. In addition, motion during volume acquisition causes inaccuracies in image preprocessing, as slice timing acquisition and motion correction algorithms do not adjust for the shift in slice locations within the volume.

A relatively recent development is the use of *accelerated imaging*, which involves sub-sampling k-space, i.e., acquiring a volume more quickly by only sampling some frequencies. This reduces the BOLD signal/noise ratio, but allows image acquisition to be accelerated. Some common acceleration schemes are SENSE and GRAPPA. The acceleration factor (i.e., the *SENSE factor*) determines how much faster images are acquired: An acceleration factor of 2 decreases the TR by $\frac{1}{2}$, 3 by $\frac{1}{3}$, etc. This can be done without loss in spatial resolution due to advances in *parallel imaging*, the simultaneous use of multiple RF receive coils to collect data. Many fMRI systems are now equipped with 8-channel head coils, and 32-channel coils are also commercially available. Because the coils are smaller, they must be positioned closer to the head surface, and they produce a signal/noise benefit primarily on the brain surface.

In summary, tradeoffs between all these factors make different acquisition schemes optimal for different study goals (Table 1). A good strategy is to work with a physicist to test different combinations of parameters before the study, to balance the coverage, spatial and temporal resolution, and susceptibility artifacts. As a starting point for a customized acquisition scheme, two examples are below. A typical whole-brain acquisition for our group is ~ 35 slices with TR = 2 sec, $3 \times 3 \times 3$ mm voxels, and an 8-channel coil with SENSE factor 2. For experiments in which signal in the orbitofrontal cortex and amygdala is of interest, an alternate scheme might be 15 slices at TR = 1 sec, with $1.5 \times 1.5 \times 1.5$ mm voxels but only coverage of a slab of brain tissue including the areas of interest.

III. ANALYZING DATA

There are several common objectives in the statistical analysis of fMRI data. These vary from simply localizing regions of the brain activated by a certain task, to determining distributed networks that correspond to brain function and making predictions about psychological or disease states. In the following section we discuss the standard approach towards preprocessing fMRI data, the general linear model (GLM), and emerging techniques for modeling brain connectivity.

III.A Preprocessing fMRI data

Prior to analysis, fMRI data typically undergoes a series of preprocessing steps aimed at reconstructing images from the raw data, identifying and removing artifacts and validating model assumptions. A common pipeline is shown in Figure 7. The goals of preprocessing are to minimize the influence of data acquisition and physiological artifacts, validate statistical assumptions and standardize the locations of brain regions across subjects in order to achieve increased validity and sensitivity in group analysis. When analyzing fMRI data it is

typically assumed that all voxels in a brain volume were acquired simultaneously. Further, it is assumed that each data point in a specific voxel's time series only consists of signal from that voxel (i.e., that the participant did not move in between measurements). Finally, when performing group analysis and making population inference, all individual brains are assumed to be registered, so that each voxel is located in the same anatomical region for all subjects. Without pre-processing the data prior to analysis, none of these assumptions would hold and the resulting statistical analysis would be invalid. The major steps in fMRI preprocessing are reconstruction, slice acquisition timing correction, realignment, co-registration of structural and functional images, registration or nonlinear warping to a template (also called normalization), and smoothing.

Image Reconstruction

The first stage of preprocessing is image reconstruction. The raw MR signal is acquired in k-space and the data is subsequently transformed into brain volumes; each consisting of a number of uniformly spaced volume elements, or voxels, whose intensity represents the spatial distribution of the nuclear spin density within that particular voxel. There are a variety of approaches towards sampling the data, each with its own benefits relating to speed, signal-to-noise ratio and ease of reconstruction. For a more complete discussion see Lindquist (Lindquist, 2008).

Visualization and artifact reduction

Neuroimaging data contain artifacts that arise from a number of sources, including head movement, reconstruction and interpolation processes, and brain movement and vascular effects related to periodic physiological fluctuations. Functional MRI data often contains transient spike artifacts and slow drift over time related to a variety of sources, including magnetic gradient instability, RF interference, and movement-induced inhomogeneities in the magnetic field. These artifacts will likely violate the assumption of normally and identically distributed errors; unless dealt with, the consequences include reduced power in group analysis and increased false positives in single-subject inference.

A first line of defense is to examine the data using exploratory techniques and diagnose problems. There are many possible sources of artifacts, some that signal fundamental problems with the scanner setup (e.g., RF leaks) and others that are inevitable and whose impact can be minimized during analysis. However, because fMRI data involves hundreds of thousands of measurement over time, visualizing the data is not trivial. A popular approach is to extract principal components or independent components from the whole-brain time series and visualize them. Multiple packages, such as FSL's Melodic software, the Group ICA of Functional MRI (GIFT) toolbox, fMRISTAT software, and others provide ways of doing this.

Some ICA/PCA packages also provide the capability to remove components deemed artefactual from the data, and there is increasing interest in using such methods to "denoise" fMRI data (Nakamura et al., 2006; Tohka et al., 2007). One central difficulty is determining which components to remove. Many components are

likely to involve a mix of task-related and artifact-related signals, and there is no guarantee that removing a component will increase the accuracy or precision of estimates of task-related effects.

Another popular approach is the estimation and inclusion of physiological artifacts from known sources, including participant head motion, heart-beat, respiration, and fluctuations in CO² concentration with breathing. Several recent papers show benefits of estimating a canonical “cardiac response function” “respiration response function” and de-noising of fMRI data by combining heart-beat and respiration measures with these functions during scanning (Chang, Cunningham, & Glover, 2009).

One all-too-common artifact that cannot be corrected by model-based physiological correction is gradient artifacts or “spikes” in the data that can affect one or more slices. A growing number of laboratories have developed procedures for identifying such artifacts and minimizing their influence, either through robust estimation techniques or simply removing, interpolating, or Winsorizing them. A sensible approach, which we prefer, is to identify spikes and create dummy regressors (which model only that time point) for inclusion in the statistical analysis. Toolboxes such as AFNI’s 3d despike, ArtRepair (Mazaika, Hoefft, Glover, & Reiss, 2009), and others may also be useful.

A potential issue with all of these approaches is that single-subject inferences may be inaccurate (i.e., p-values may be wrong) if variance is removed from the data without accounting for it in the subsequent statistical analysis. Any removal of components or known physiological effects should be accompanied by an appropriate reduction in the degrees of freedom in the statistical model. If covariates are entered in a regression analysis, this reduction is automatically performed. But if filtering or denoising is done *before* analysis, it is more difficult to produce valid single-subject p-values, particularly if components are removed rather than covariates specified *a priori*.

Slice timing correction

When analyzing 3D fMRI data it is typically assumed that the entire brain volume is measured simultaneously. In reality, because the volume consists of multiple slices that are sampled sequentially, and therefore at different time points, similar time courses from different slices will be temporally shifted relative to one another. Thus, it is common practice to correct the signal intensity in all voxels, so they can be considered to be sampled at the same point in the acquisition period. This can be done by interpolating the signal intensity at the chosen time point from the same voxel in previous and subsequent acquisitions. A number of interpolation techniques exist (e.g., bilinear and sinc interpolation) with varying degrees of accuracy and speed. Sinc interpolation is the slowest, but generally the most accurate. Some researchers do not use slice timing, as it adds interpolation error to the data, and instead use more flexible hemodynamic models to account for variations in acquisition time.

Motion Correction

An important issue involved in any fMRI study is the proper handling of any subject movement that may have taken place during data acquisition. Even small amounts of head motion can be a major source of error if not treated correctly. All typical analyses assume that each voxel contains signal in the same brain area from image to image. When movement occurs, the signal from a specific voxel will come from different areas of the brain, blurring the signal and creating edge artifacts. Even worse, head movement creates alterations in the magnetic field that can result in large susceptibility artifacts. The first step in correcting for motion, and a partial solution for the first problem above, is to find the best possible alignment between the input image and some target image (e.g., the first image or the mean image). A rigid body transformation involving 6 variable parameters is used. This allows the input image to be translated (shifted in the x, y, and z directions) and rotated (altered roll, pitch, and yaw) to match the target image. Usually, the matching process is performed by minimizing some cost function (e.g., the sums of squared differences) that assesses similarity between the two images. Once the parameters that achieve optimal realignment are determined, the image is re-sampled using interpolation to create new motion corrected voxel values. This procedure is repeated for each brain volume.

Realignment corrects adequately for small movements of the head, but it does not correct for the more complex spin-history artifacts created by the motion. The parameters at each time point are saved for later inspection and are often included in the analysis as covariates of no interest; however, even this additional step does not completely remove artifacts created by head motion. Residual artifacts remain in the data and contribute to noise. Sometimes this noise is correlated with task contrasts of interest, which can create false results in single-subject analyses. Because these artifacts typically differ in sign and magnitude across subjects, group analyses are usually robust to such artifacts in terms of false positives (though there are exceptions), but movement artifacts can severely compromise power.

Co-registration

Functional MRI data is typically of low spatial resolution and provides relatively little anatomical detail. Therefore, high-resolution structural images are used for warping and localization. The same transformations (warps) are applied to the functional images, which produce the activation statistics, so accurate registration of structural and functional images is critical. Co-registration aligns structural and functional images, or in general, different types of images of the same brain. Because functional and structural images are collected using different sequences and different tissue classes have different average intensities, using a least squares difference method to match images is typically not appropriate. Instead, an affine transformation can be used that maximizes the *mutual information* among the two images, or the degree that knowing the intensity of one can be used to predict the intensity of the other (Cover & Thomas, 1991). Typically, a single structural image is co-registered to the first or mean functional image.

Normalization

For group analysis, each voxel must lie within the same brain structure for each individual subject. Of course individual brains differ in their shapes and features, but there are similarities shared by every non-pathological brain. Normalization attempts to register each subject's anatomy to a standardized stereotaxic space defined by a template brain (e.g., the Talairach or Montreal Neurological Institute [MNI] brain). In this scenario, it is common to use nonlinear transformations to match local features. One begins by estimating a smooth continuous mapping between the points in an input image with those in the target image. Next, the mapping is used to re-sample the input image so that it is warped onto the target image, resulting in a normalized image that can be compared with similarly normalized images obtained from other subjects. The main benefits of normalizing data is that spatial locations can be reported and interpreted in a consistent manner, results can be generalized to a larger population and results can be compared across studies and subjects. The drawbacks are that it reduces spatial resolution and may introduce errors due to interpolation.

Smoothing

It is common practice to spatially smooth fMRI data prior to analysis. Smoothing typically involves convolving the functional images with a Gaussian kernel, described by the full width of the kernel at half its maximum height (FWHM). Common values vary between 4-12mm FWHM. There are several reasons why fMRI data is smoothed. First, it may improve inter-subject registration and overcome limitations in the spatial normalization by blurring any residual anatomical differences. Second, it ensures that the assumptions of random field theory (RFT), commonly used to correct for multiple-comparisons, are valid. A rough estimate of the amount of smoothing required to meet the assumptions of RFT is a FWHM of 3 times the voxel size (e.g., 9mm for 3mm voxels), but the required value may be much higher for small samples (e.g., < 20 subjects in a group analysis; Nichols & Hayasaka, 2003). Third, if the spatial extent of a region of interest is larger than the spatial resolution, smoothing may reduce random noise in individual voxels and increase the signal-to-noise ratio within the region.

The process of spatially smoothing an image is equivalent to applying a low-pass filter to the sampled k-space data prior to reconstruction. This implies that much of the acquired data is discarded as a byproduct of smoothing and temporal resolution is sacrificed without gaining any benefits. Additionally, acquiring an image with high spatial resolution and then smoothing the image is not the same as directly acquiring a low resolution image. The signal-to-noise ratio during acquisition increases as the square of the voxel volume, so acquiring small voxels means that signal is lost that can never be recovered. Hence, it is optimal in terms of sensitivity to acquire images at the desired resolution and not employ smoothing.

Previously, many investigators also *temporally smoothed* their data. This procedure removes high-frequency signals from the data. However, this approach has largely been replaced by more standard time series models (e.g., autoregressive modeling). There is no expected benefit to temporal smoothing on sensitivity, as it further decreases the temporal resolution of the data, and it is not recommended.

III.B. The General Linear Model

The general linear model (GLM) is a linear analysis method that subsumes many basic analysis techniques, including t-tests, ANOVA, and multiple regression (Worsley & Friston, 1995). It is by far the most common approach for assessing task–brain activity relationships in the neuroimaging literature. The GLM can be used to estimate whether the brain responds to a single type of event, to compare different types of events, to assess correlations between brain activity and behavioral performance or other psychological variables, and perform several types of tests of functional and effective connectivity. The GLM is appropriate when multiple predictor variables are used to explain variability in a single, continuously distributed outcome variable where the predictors are related to psychological events, and the outcome variable is signal in a brain voxel or region of interest. The analysis is typically ‘massively univariate,’ meaning that the analyst performs a separate GLM analysis at every voxel in the brain, and summary statistics are saved in maps of statistic values across the brain.

GLM model basics

Functional MRI analysis typically proceeds in three stages: The “first level” analysis, in which activation parameter estimates for each condition (event type), voxel, and individual are obtained. Differences in activation parameter estimates across conditions, or contrasts, are then estimated for each individual in a second stage. Finally, a group analysis is conducted on contrast values to test whether the observed activation differences between conditions are likely to exist in the population. This is typically referred to as “second level” analysis.

First-level analysis. In the first level analysis, separate GLM models are estimated on the time series data for each voxel within each individual. For a single subject, the fMRI time course from one voxel is the outcome variable (y). Activity is modeled as the sum of a series of independent predictors (x_1, x_2 , etc.) related to task conditions and other nuisance covariates of no interest (e.g., head movement estimates). For each task condition or event type of interest, a time series of the predicted shape of the signal response is constructed, using prior information about the shape of the vascular response to a brief impulse of neural activity. An example is shown in Figure 6B. The vectors of predicted time series values for each task condition are collated into the columns of the design matrix, X . The GLM fitting procedure estimates the amplitude for each column of X , so that the sum of fitted values across all predictors’ best fits the data. These amplitudes are regression slopes, denoted $\hat{\mathbf{a}}$ (the “hat” denotes an estimate of a theoretical constant value). It also estimates a time series of error values, $\hat{\mathbf{a}}$, that cannot be explained by the model. The model is thus described by the equation:

$$\mathbf{y} = \mathbf{X}\hat{\mathbf{a}} + \hat{\mathbf{a}} \quad (5)$$

where $\hat{\mathbf{a}}$ is a $k \times 1$ vector of regression slopes, $\mathbf{X}$ is an $n \times k$ model matrix, $\mathbf{y}$ is an $n \times 1$ vector containing

the observed data, and $\hat{\mathbf{a}}$ is an $n \times 1$ vector of unexplained error values. Error values are assumed to be independent and follow a normal distribution with mean 0 and standard deviation σ . The estimated $\hat{\mathbf{a}}$ s correspond to the estimated magnitude of activation for each psychological condition described in the columns of $\mathbf{X}$.

One of the advantages of the GLM is that there exists an algebraic solution for $\hat{\mathbf{a}}$ that minimizes the squared error. The estimation yields a series of maps of $\hat{\mathbf{a}}$ s across voxels for each subject. Those $\hat{\mathbf{a}}$ s are used as estimates of the BOLD response amplitude for each condition or event type.

Contrasts. Differences among conditions can be assessed using contrasts, or linear combinations of the $\hat{\mathbf{a}}$ s. For example, imagine that Condition A and Condition B shown in Figure 4 reflect periods of auditory and visual stimulation, respectively. The contrast $[A - B]$ estimates the activation difference for auditory – visual stimulation, and it can be tested by applying a linear contrast vector $\mathbf{c}$ (i.e., calculating $\mathbf{c}\hat{\mathbf{a}}$). In this case $\mathbf{c} = [1 \ -1]$. In general, multiple contrasts can be specified, and they are usually calculated after estimating the model and saved as subject-specific contrast images, with one image per contrast.

Contrasts can be used to estimate signal magnitudes in response to a single condition (i.e., $\mathbf{c} = [1 \ 0]$), an average over multiple conditions ($\mathbf{c} = [1 \ 1]$), or a difference in magnitude across conditions ($\mathbf{c} = [1 \ -1]$). Hypothesis testing is performed in the usual manner by testing individual model parameters or contrasts using a t-test. Contrast images are usually entered as data for a group analysis, as described below. In addition, another kind of contrast, the F-contrast, can be used to test the significance of a subset of parameters (e.g., do a set of nuisance covariates explain a significant amount of variation in the data?). However, F-contrast images are not suitable for group analysis.

Group analysis. Though making inferences about single subject is common in some areas (i.e., fMRI of early visual processes), most researchers are interested in generalizing their results to a population of unobserved individuals--science is, after all, about generalizable knowledge. To achieve this goal requires a group analysis. The group analysis consists of specifying a second-level GLM model, this time using as the data ($\mathbf{y}$) a set of contrast images for a single contrast, one per subject. The simplest design matrix is an intercept term (all values constant), which performs a one-sample t-test of whether contrast values are non-zero at each voxel. However, it is possible to specify designs that implement a two-sample t-test (e.g., regressors with 1 or -1), regression relating individual differences in a continuous predictor to contrast values, and ANOVA and other designs. One common error that it is critical to avoid is that typically only one image per subject is entered into the group analysis. If there is more than one, then the group analysis has repeated measures, and correlated errors across repeated measures must be dealt with appropriately.

The two-stage first- and second-level analysis scheme implements a particular form of a mixed-effects

model, a model containing both fixed and random effects. Levels of variables treated as fixed effects are assumed to be invariant (e.g., males and females; visual and auditory modalities), and no inferences are made about previously unobserved levels. Levels of variables treated as random effects, by comparison, are assumed to be randomly sampled from a larger population, and the model permits inferences to new, unobserved levels. The variable “subject identity” is a classic example of a variable that should be treated as a random effect, and this is what the analysis described above does.

However, the two-stage analysis makes an assumption that is often violated: it assumes that the contrast standard error is the same across all subjects, which implies identical design matrices and residual variances. This is rarely if ever true in practice, which has led to the development and increasing popularity of full mixed-effects models. Such models incorporate both first-level (within-subjects) and second-level (between-subjects) effects into one model, and in so doing can relax the assumption of homogenous standard errors and appropriately weight the contribution of each subject’s estimates to the group estimates in proportion to their precision. The larger a subject’s standard error, the less reliable their estimate, and the less that subject should contribute to the group results. Fitting this type of model requires estimating variance components at each level. Typically, separate error variances are estimated for (1) measurement error within-subjects (and model mis-fitting), and (2) the inter-individual differences among subjects. Restricted maximum likelihood (ReML) is a popular type of estimate of variance components based on the residuals. Since variance estimates and model fits (s) are inter-dependent, iterative algorithms such as EM are used to estimate ReML variance components. Though the full mixed-effects model is correct in more cases and is advantageous, Mumford and Nichols have recently reported that the simple two-stage analysis is valid and performs reasonably well (Mumford & Nichols, 2009).

Model-building

The most challenging aspect of using the GLM is the creation of realistic predictions of task-related signals to include in the columns of X . To build the model, researchers start with an ‘indicator’ vector representing the neuronal activity for each condition sampled at the resolution of the fMRI experiment (Figure 6B). This vector has zero value except during hypothesized neural activation periods, where the signal is assigned a value of 1. Each indicator vector is convolved with a canonical HRF to yield a predicted time course related to that event, which forms a column of the X . If the canonical HRF fits the shape of the BOLD response to psychological events, this simple approach has great sensitivity to detect differences. However, if the canonical HRF does not fit, there is at best a drop in power, and at worst false positives and mis-interpretation of results (Lindquist & Wager, 2007; Loh, Lindquist, & Wager, 2008). As an alternative, the same ‘neural’ indicator vector can be convolved with multiple canonical waveforms and entered into multiple columns of X for a single event type (Figure 8). These reference waveforms are called basis functions, and the predictors for an event type constructed using different basis

functions can combine linearly to better fit the evoked BOLD responses. The ability of a basis set to capture variations in hemodynamic responses depends on both the number and shape of the reference waveforms. There is a fundamental tradeoff between flexibility to model variations and power. This is because each parameter is estimated with error, and flexible models can tend to model noise and thus produce noisier parameter estimates. For a critical evaluation of various basis sets see Lindquist and Wager (2007b) and Lindquist, Loh, Atlas, and Wager (2008c).

Additional predictors are typically added to account for other sources of variation in the data. One kind of predictor is the parametric modulator (Figure 9). These are increasingly used to model trial-to-trial variation in a psychological process, as discussed above. Another kind is nuisance covariates, which are included to reduce noise and to prevent signal changes related to head movement and physiological (e.g., respiration) artifacts from influencing the contrast estimates. In fMRI, the signal typically drifts slowly over time, so that the most power is in the lowest temporal frequencies and the noise is autocorrelated. Un-modeled physiological artifacts may give rise to much of this autocorrelated noise (Lund, Madsen, Sidaros, Luo, & Nichols, 2005). In addition to modeling known physiological artifacts, covariates that implement high-pass filtering, or removal of signal frequencies below a specified cutoff, can also be added at this stage. This characteristic has prompted the widespread use of high-pass filters that removes fluctuations below a specified frequency cutoff from the data. High-pass filtering is often performed in the GLM analysis by adding covariates of no interest (e.g. low-frequency cosines). Of course, care must be taken to ensure that the fluctuations induced by the task design are not in the range of frequencies removed by the filter.

III.C Multiple Comparisons

The results of fMRI studies are usually summarized in a statistical parametric map (SPM). These maps describe brain activation by color-coding voxels whose t-values (or comparable statistics) exceed a certain statistical threshold for significance. The implication is that these voxels are activated by the experimental task. When constructing such a map it is important to carefully consider the appropriate threshold to use when declaring a voxel active. In a typical experiment up to 100,000 hypothesis tests (one for each voxel) are performed simultaneously, and it is crucial to correct for multiple comparisons. Several approaches towards controlling the false positive rate have been used; the fundamental difference between methods lies in whether they control the family-wise error rate (FWER) or the false discovery rate (FDR).

Methods that control the FWER based on the effect size in each voxel--e.g., random field theory, "height" thresholds and Bonferroni correction--tend to give over-conservative results (Hayasaka and Nichols, 2004). They tend to correct for more spatial comparisons that are actually made, because voxels are, as a whole, are positively inter-correlated across the brain. These methods provide the gold standard for performing a whole-brain analysis

and interpreting any region that is significant with FWER correction can typically be considered truly “activated.” (However, the commonly used “extent-based” correction in SPM is under-conservative and should be used with caution.) They are not optimal for estimating which voxels show true effects and which do not--i.e., balancing false positives and false negatives--because there is typically extremely low power, or a low chance of detecting an effect in any single truly activated voxel. Basic power calculations suggest that even with large effect sizes ($d = 1$), typical sample-sizes yield voxel-wise FWER-corrected power rates of only around 5% (see Figure 10).

The false discovery rate (FDR) controls the proportion of false positives among all rejected tests (Genovese, Lazar, and Nichols, 2002). Controlling the FDR at $q < .05$ (by convention, q is used instead of p) implies that on average, only 5% of the significant results are expected to be false positives. The FDR controlling procedure is adaptive in the sense that the larger the signal, the lower the threshold. If there are no effects anywhere in the brain except in one voxel, the FDR-corrected threshold will be equivalent to the FWER-corrected threshold for that voxel. If there are highly significant effects everywhere in the brain except one voxel in question, the FDR-corrected threshold will be equal to the uncorrected p -value threshold (e.g., $p < .05$). Researchers are increasingly using FDR correction because it provides an often reasonable balance between true and false positives in small-sample studies. However, its major limitation is that it does not allow one to infer that any specific findings are not false positives. It is reasonable to infer that most of the findings are true positives, but does not provide information about which ones are not.

III.B. Assessing brain connectivity

To date, human brain mapping has been primarily used to provide maps that show which regions of the brain are activated by specific tasks. Recently, there has been an increased interest in augmenting this type of analysis with connectivity studies that describe how various brain regions interact and how these interactions depend on experimental conditions. It is common practice in the analysis of neuroimaging data to make the distinction between *functional* and *effective connectivity* {Friston, 1994 #1}. Functional connectivity is defined as the undirected association between two or more fMRI time series, while effective connectivity is the directed influence of one brain region on the physiological activity recorded in other brain regions; it implies both causality and directness. It implies causality because the models used to assess effective connectivity are usually directional, and directness in the sense that effective connectivity measures attempt to partial out indirect influences from other regions.

Functional connectivity is a statement about observed associations among regions and/or other performance and physiological variables—for example, the correlation between time series in two regions (bivariate connectivity). Simple functional connectivity analyses usually compare correlations between ROIs,

sometimes in a task-dependent fashion, or between a ‘seed’ region of interest and voxels throughout the brain. Multivariate analysis methods are also used to reveal networks of multiple interconnected regions. Popular methods include Principal Components Analysis (PCA) (Andersen AH, 1999), Partial Least Squares (PLS) (McIntosh, 1996) and Independent Components Analysis (ICA) (Calhoun, 2001; McKeown, 1998). Connectivity between two or more regions may result from direct influences (i.e., functional links between regions) or indirect effects due to common input from a third variable. None of these methods are able to address issues of causality or the common influences of other variables.

Functional connectivity methods can be applied at different levels of analysis, with different interpretations at each level. Connectivity across time series data can reveal networks that are dynamically co-activated over time (either ‘intrinsically,’ regardless of task state, or in a task-dependent fashion), and is closest to the concept of communication among regions, though it does not conclusively demonstrate that. Connectivity across single-trial response estimates (Rissman, Gazzaley, & D’Esposito, 2004) can identify coherent networks of task-related activations. Whereas these levels are only accessible to fMRI and EEG/MEG, which provide relatively rich time series data, other levels of analysis may be examined in PET studies as well. Connectivity across subjects can reveal patterns of coherent individual differences, which may result from communication among regions but also from differences in strategy use or other genetically determined or learned differences among individuals. Finally, connectivity across studies can reveal tendencies for studies to co-activate within sets of regions, which may be influenced by any of the factors mentioned above, and also differences among tasks or other study-level variables.

Effective connectivity analysis, on the other hand, is model-dependent. Typically, a small set of regions and a proposed set of connections are specified a priori, and tests of fit are used to compare a small number of alternative models and assess the statistical significance of individual connections. Because connections may be specified directionally (with hypothesized causal influences of one area on another), the model implies causal relationships. Because there are many possible models, the choice of regions and connections must be anatomically motivated. Most effective connectivity depends on two models: a neuroanatomical model that describes which areas are connected, and a mathematical model that describes how areas are connected. Common methods include Structural Equation Modeling (SEM) (McIntosh, 1994) and Dynamic Causal Modeling (DCM) (Friston, 2003). While ‘effective connectivity’ methods have become increasingly popular, it is important to keep in mind that the conclusions about direct influences and causality obtained using these models are only as good as the specified models. Any misspecification of the underlying model will almost certainly lead to erroneous conclusions. In particular, the exclusion of important lurking variables (brain regions involved in the network but not included in the model) can completely change the fit of the model and thereby affect both the direction and strength of the connections. Great care always needs to be taken when interpreting the results of these methods.

The distinction between functional and effective connectivity is not entirely clear (Horwitz, 2003). If the

discriminating features are a directional model in which causal influences are specified; and the willingness to make claims about direct vs. indirect connections, then many analyses, including multiple regression, might count as effective connectivity. While the reason many researchers use both SEM and DCM is to obtain the goal of ascribing causality between different brain regions, it is important to keep in mind that the tests performed in both techniques are based on model fit rather than on the causality of the effect. Similarly, Granger causality (Roebroeck, Formisano, & Goebel, 2005) is another approach that is typically considered to test effective connectivity, though neither causal influences nor direct vs. indirect effects are tested within the basic model framework. Causality is tested strictly in the sense of temporal relationships, rather than on whether activity in a brain region is necessary or sufficient for activity in another. In the end, it is not the label of “functional” or “effective” that is important, but the specific assumptions and robustness and validity of inference afforded by each method.

When performing connectivity and correlation studies it is tempting to make statements regarding causal links between different brain regions. The idea of causality is a very deep and important philosophical issue (Pearl, 2000; Rubin, 1974). Often a cavalier attitude is taken in attributing causal effects and the differentiation between explanation and causation is often blurred. Properly randomized experimental designs permit causal inferences of task manipulations on brain activity. However, in neuroimaging and EEG/MEG studies, all the brain variables are observed, and none are manipulated. Therefore, we do not recommend making strong conclusions about causality and ‘direct’ influences among brain regions using these methods, because the validity of such conclusions is very difficult to verify. The combination of neuroimaging and TMS or related forms of brain stimulation (Bohning et al., 1997) may provide more reliable causal inferences about the effects of activating one brain region on another. By stimulating the brain, experimental manipulation of one brain area can be achieved and its causal effects on other brain regions thus examined. However, the problem remains of assessing which effects are ‘direct’ as opposed to mediated by other intervening regions.

One way of circumventing some of the issues with making claims about directness and causality is not to make them. In recent work, we have taken the approach of fitting mediation models, which are inherently directional, but using them as a means for discovery and modeling of functional pathways rather than as a basis for claims that one brain region is directly influencing another one (Wager, Davidson, Hughes, Lindquist, & Ochsner, 2008). We have used formal mediation tests from path analysis to identify brain voxels that mediate the relationship between activity in one brain area (an “initial variable”, X) and either another area or a behavioral or physiological outcome of interest (Y , e.g., performance). The areas that are significant mediators are likely to be part of a functional pathway or network that relates the initial variable to the outcome--i.e., there is an X - M - Y pathway. If all the variables are observed rather than experimentally manipulated, we feel more comfortable avoiding inferences about the direction of causality.

The mediation effect is related to, but not the same as, another effect commonly localized in fMRI analyses: a moderation or modulation effect. Whereas mediation provides evidence for a pathway connecting more than two variables, and tests the product of a test of moderation is an interaction test: The level of a moderator (e.g., high vs. low) influences connectivity between Region 1 and Region 2.

The path modeling framework can also be used to localize brain mediators or moderators of experimental independent variables, which permits stronger claims about causality. For example, a mediation analysis could be used to localize voxels that mediate effects of an experimental design variable (e.g., Task A - Control B) on a behavioral or physiological outcome of interest (e.g., performance). The causality of the task-to-brain relationship, at least, can be established with relative confidence.

A particular kind of moderation effect, the “psychophysiological interaction” (Friston et al., 1997) has been popularized in the fMRI literature and is included in the SPM software. In this analysis, the interaction between an independent task variable and a ‘seed’ brain region is calculated, and a search over all voxel time-series is conducted for correlates of that interaction. A common interpretation is that the task condition moderates the connectivity between the ‘seed’ and other significant regions. While the effect of the task on brain is likely causal, it is difficult to establish which of the two brain regions (‘seed’ or result) is the moderator, and which is moderated. The interaction can be interpreted in multiple ways.

Many researchers have now attempted to use the timing of responses across the brain to infer causality, and many such approaches are actively being developed. Variants of Granger causal models (Granger, 1969) are an example, as they can localize relationships between brain regions that have different relative timing. However, because of the potentially substantial variability in hemodynamic responses across the brain, these differences may be due to vascular differences rather than neural timing differences. In addition, temporal precedent does not necessitate causality: A rooster crowing every morning “Granger causes” the sun to rise. Interactions between task variables and timing (e.g., demonstrating that Region 1 precedes Region 2 more in Task A than in Task B) provide potentially stronger inferences, but there are alternative explanations. First, if both regions respond to Task A, but not Task B, any hemodynamic latency difference is likely to create a latency $\times$ task interaction in a Granger model. In addition, with all of the path and Granger models discussed above, one cannot know whether there is not another, unmeasured variable (e.g., system-wide release of a neuromodulator such as dopamine or opioids) that is responsible for creating the observed functional/effective connectivity.

IV. CONCLUSIONS

In this chapter, we have reviewed the essential elements of conducting an fMRI study, including conceiving and designing an experiment, acquiring data, and analyzing the data with univariate and multivariate statistical

models. In closing, we would like to emphasize a few points. First, fMRI-based cognitive and affective neuroscience is a rapidly growing area. There are many perspectives and techniques, and these are changing continuously. As of this writing, there are literally hundreds of published, freely available toolboxes for fMRI data analysis. Many of these are collected on the NIH-funded Neuroimaging Informatics Tools and Resources Clearinghouse (NITRC; <http://www.nitrc.org/>). We hope that this brief introduction will provide some pointers towards further study and some practical advice. Second, fMRI is a collaborative endeavor, and there are very few individuals (if any!) who fully understand every aspect of the process. An fMRI “dream team” would likely include a psychologist, physicist, statistician, computer scientist, engineer, neuroscientist, and neuroanatomist, each of whom would contribute unique and complimentary knowledge. We believe that fMRI should be played like a team sport, with consultation obtained and credit given on each of the various aspects to the extent practicable.

The challenges inherent in the complexity of data collection and analysis are formidable, but progress is being made on every front. Over the past 10 years, we have seen a marked increase in the quality of fMRI studies. The data quality is increasing because scanners are becoming more sensitive and more stable. Best-practice procedures for study design and analysis are beginning to emerge, and more groups use valid analysis techniques on larger samples. And finally, data analysis techniques from several fields are being brought to bear in creative and powerful ways, providing qualitatively new views on the function of the human brain. It is exciting to think of what developments lie ahead.

References

- Aguirre, G. K., Zarahn, E., & D'Esposito, M. (1998). The variability of human, BOLD hemodynamic responses. *Neuroimage*, 8(4), 360-369.
- Andersson, J. L., Hutton, C., Ashburner, J., Turner, R., & Friston, K. (2001). Modeling geometric deformations in EPI time series. *Neuroimage*, 13(5), 903-919.
- Bellgowan, P. S., Saad, Z. S., & Bandettini, P. A. (2003). Understanding neural system dynamics through task modulation and measurement of functional MRI amplitude, latency, and width. *Proc Natl Acad Sci U S A*, 100(3), 1415-1419.
- Birn, R. M., Saad, Z. S., & Bandettini, P. A. (2001). Spatial heterogeneity of the nonlinear dynamics in the fMRI BOLD response. *Neuroimage*, 14(4), 817-826.
- Bohning, D. E., Pecheny, A. P., Epstein, C. M., Speer, A. M., Vincent, D. J., Dannels, W., et al. (1997). Mapping transcranial magnetic stimulation (TMS) fields in vivo with MRI. *Neuroreport*, 8(11), 2535-2538.
- Buchel, C., Holmes, A. P., Rees, G., & Friston, K. J. (1998). Characterizing stimulus-response functions using nonlinear regressors in parametric fMRI experiments. *Neuroimage*, 8(2), 140-148.
- Buchel, C., Morris, J., Dolan, R. J., & Friston, K. J. (1998). Brain systems mediating aversive conditioning: an event-related fMRI study. *Neuron*, 20(5), 947-957.
- Buracas, G. T., & Boynton, G. M. (2002). Efficient design of event-related fMRI experiments using M-sequences. *Neuroimage*, 16(3 Pt 1), 801-813.
- Buxton, R. B., Wong, E. C., & Frank, L. R. (1998). Dynamics of blood flow and oxygenation changes during brain activation: the balloon model. *Magn Reson Med*, 39(6), 855-864.
- Cacioppo, J. T., & Tassinary, L. G. (1990). Inferring psychological significance from physiological signals. *Am Psychol*, 45(1), 16-28.
- Calhoun, V. D., Adali, T., Pearlson, G. D., and Pekar, J. J. (2001). Spatial and temporal independent component analysis of functional MRI data containing a pair of task-related waveforms. *Human Brain Mapping*, 13, 43-53.
- Chang, C., Cunningham, J. P., & Glover, G. H. (2009). Influence of heart rate on the BOLD signal: the cardiac response function. *Neuroimage*, 44(3), 857-869.
- Coltheart, M. (2006). What has functional neuroimaging told us about the mind (so far)? *Cortex*, 42(3), 323-331.
- Constable, R. T., & Spencer, D. D. (1999). Composite image formation in z-shimmed functional MR imaging. *Magn Reson Med*, 42(1), 110-117.
- Cover, T. M., & Thomas, J. A. (1991). Elements of Information Theory. In (pp. 18-26). New York: Wiley.
- Dagher, A., Owen, A. M., Boecker, H., & Brooks, D. J. (1999). Mapping the network for planning: a correlational PET activation study with the Tower of London task. *Brain*, 122 (Pt 10), 1973-1987.
- Dale, A. M., Liu, A. K., Fischl, B. R., Buckner, R. L., Belliveau, J. W., Lewine, J. D., et al. (2000). Dynamic Statistical Parametric Mapping Combining fMRI and MEG for High-Resolution Imaging of Cortical Activity. *Neuron*, 26(1), 55-67.
- Etkin, A., & Wager, T. D. (2007). Functional neuroimaging of anxiety: a meta-analysis of emotional processing in PTSD, social anxiety disorder, and specific phobia. *Am J Psychiatry*, 164(10), 1476-1488.
- Friston, K., Harrison, L., Penny, W. (2003). Dynamic causal modelling. *Neuroimage*, 19, 1273-1302.
- Friston, K. J., Buechel, C., Fink, G. R., Morris, J., Rolls, E., & Dolan, R. J. (1997). Psychophysiological and modulatory interactions in neuroimaging. *Neuroimage*, 6(3), 218-229.
- Friston, K. J., Mechelli, A., Turner, R., & Price, C. J. (2000). Nonlinear responses in fMRI: the Balloon model, Volterra kernels, and other hemodynamics. *Neuroimage*, 12(4), 466-477.
- Granger, C. W. J. (1969). Investigating causal relations by econometric models and cross-spectral methods. *Econometrica* 37, 424-438.
- Grill-Spector, K., & Malach, R. (2001). fMR-adaptation: a tool for studying the functional properties of human cortical neurons. *Acta Psychol (Amst)*, 107(1-3), 293-321.
- Grinband, J., Hirsch, J., & Ferrera, V. P. (2006). A neural representation of categorization uncertainty in the human brain. *Neuron*, 49(5), 757-763.

- Hare, T. A., Camerer, C. F., & Rangel, A. (2009). Self-control in decision-making involves modulation of the vmPFC valuation system. *Science*, 324(5927), 646-648.
- Haynes, J. D., & Rees, G. (2006). Decoding mental states from brain activity in humans. *Nature Reviews Neuroscience*, 7(7), 523-534.
- Henson, R. N. (2003). Neuroimaging studies of priming. *Prog Neurobiol*, 70(1), 53-81.
- Horwitz, B. (2003). The elusive concept of brain connectivity. *Neuroimage*, 19, 466-470.
- Johnson, M. K., Raye, C. L., Mitchell, K. J., Greene, E. J., Cunningham, W. A., & Sanislow, C. A. (2005). Using fMRI to investigate a component process of reflection: prefrontal correlates of refreshing a just-activated representation. *Cogn Affect Behav Neurosci*, 5(3), 339-361.
- Josephs, O., & Henson, R. N. (1999). Event-related functional magnetic resonance imaging: modelling, inference and optimization. *Philos Trans R Soc Lond B Biol Sci*, 354(1387), 1215-1228.
- Kosslyn, S. M., Thompson, W. L., Kim, I. J., & Alpert, N. M. (1995). Topographical representations of mental images in primary visual cortex. *Nature*, 378(6556), 496-498.
- Kwong, K. K., Belliveau, J. W., Chesler, D. A., Goldberg, I. E., Weisskoff, R. M., Poncelet, B. P., et al. (1992). Dynamic magnetic resonance imaging of human brain activity during primary sensory stimulation. *Proc Natl Acad Sci U S A*, 89(12), 5675-5679.
- Lindquist, M., Zhang, C.-H., Glover, G. H., & Shepp, L. (2008). Rapid Three-Dimensional Functional Magnetic Resonance Imaging of the Negative BOLD Response. *Journal of Magnetic Resonance*, 191, 100-111.
- Lindquist, M. A. (2008). The Statistical Analysis of fMRI Data. *Statistical Science*, 23(4), 439-464.
- Lindquist, M. A., & Wager, T. D. (2007). Validity and power in hemodynamic response modeling: a comparison study and a new approach. *Hum Brain Mapp*, 28(8), 764-784.
- Liu, T. T. (2004). Efficiency, power, and entropy in event-related fMRI with multiple trial types. Part II: design of experiments. *Neuroimage*, 21(1), 401-413.
- Liu, T. T., Frank, L. R., Wong, E. C., & Buxton, R. B. (2001). Detection power, estimation efficiency, and predictability in event-related fMRI. *Neuroimage*, 13(4), 759-773.
- Logothetis, N. K., Pauls, J., Augath, M., Trinath, T., & Oeltermann, A. (2001). Neurophysiological investigation of the basis of the fMRI signal. *Nature*, 412(6843), 150-157.
- Loh, J. M., Lindquist, M. A., & Wager, T. D. (2008). Residual Analysis for Detecting Mis-modeling in fMRI. *Statistica Sinica*, To appear.
- Lund, T. E., Madsen, K. H., Sidaros, K., Luo, W. L., & Nichols, T. E. (2005). Non-white noise in fMRI: Does modelling have an impact? *Neuroimage*.
- Mazaika, P. K., Hoeft, F., Glover, G. H., & Reiss, A. L. (2009). Methods and Software for fMRI Analysis of Clinical Subjects. *NeuroImage*, 47, 58-58.
- McClure, S. M., Berns, G. S., & Montague, P. R. (2003). Temporal prediction errors in a passive learning task activate human striatum. *Neuron*, 38(2), 339-346.
- McIntosh, A., Gonzalez-Lima, F. (1994). Structural equation modeling and its application to network analysis in functional brain imaging. *Human Brain Mapping*, 2, 2-22.
- McIntosh, A. R., Bookstein, F.L., Haxby, J.V., Grady, C.L. (1996). Spatial Pattern Analysis of Functional Brain Images Using Partial Least Squares. *NeuroImage*, 3, 143-157.
- McKeown, M. J., Makeig, S. (1998). Analysis of fMRI data by blind separation into independent spatial components. *Human Brain Mapping*, 6, 160-188.
- Menon, V., Ford, J. M., Lim, K. O., Glover, G. H., & Pfefferbaum, A. (1997). Combined event-related fMRI and EEG evidence for temporal-parietal cortex activation during target detection. *Neuroreport*, 8(14), 3029-3037.
- Miezin, F. M., Maccotta, L., Ollinger, J. M., Petersen, S. E., & Buckner, R. L. (2000). Characterizing the hemodynamic response: effects of presentation rate, sampling procedure, and the possibility of ordering brain activity based on relative timing. *Neuroimage*, 11(6 Pt 1), 735-759.
- Mitchell, T. M., Hutchinson, R., Niculescu, R. S., Pereira, F., Wang, X., Just, M., et al. (2004). Learning to decode cognitive states from brain images. *Machine Learning*, 57(1), 145-175.
- Mumford, J. A., & Nichols, T. (2009). Simple group fMRI modeling and inference. *Neuroimage*, 47(4), 1469-1475.

- Mumford, J. A., & Nichols, T. E. (2008). Power calculation for group fMRI studies accounting for arbitrary design and temporal autocorrelation. *Neuroimage*, 39(1), 261-268.
- Nakamura, W., Anami, K., Mori, T., Saitoh, O., Cichocki, A., & Amari, S. (2006). Removal of ballistocardiogram artifacts from simultaneously recorded EEG and fMRI data using independent component analysis. *IEEE Trans Biomed Eng*, 53(7), 1294-1308.
- Noll, D. C., Fessler, J. A., & Sutton, B. P. (2005). Conjugate phase MRI reconstruction with spatially variant sample density correction. *IEEE Trans Med Imaging*, 24(3), 325-336.
- O'Connor, D. H., Fukui, M. M., Pinsk, M. A., & Kastner, S. (2002). Attention modulates responses in the human lateral geniculate nucleus. *Nature Neuroscience*, 5(11), 1203-1209.
- Ogawa, S., Tank, D. W., Menon, R., Ellermann, J. M., Kim, S. G., Merkle, H., et al. (1992). Intrinsic signal changes accompanying sensory stimulation: functional brain mapping with magnetic resonance imaging. *Proc Natl Acad Sci U S A*, 89(13), 5951-5955.
- Ollinger, J. M., Shulman, G. L., & Corbetta, M. (2001). Separating processes within a trial in event-related functional MRI. *Neuroimage*, 13(1), 210-217.
- Pearl, J. (2000). *Causality : models, reasoning, and inference*. Cambridge, U.K. ; New York: Cambridge University Press.
- Poldrack, R. A. (2006). Can cognitive processes be inferred from neuroimaging data? *Trends Cogn Sci*, 10(2), 59-63.
- Price, C. J., & Friston, K. J. (1997). Cognitive conjunction: a new approach to brain activation experiments. *Neuroimage*, 5(4 Pt 1), 261-270.
- 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 U S A*, 98(2), 676-682.
- Rissman, J., Gazzaley, A., & D'Esposito, M. (2004). Measuring functional connectivity during distinct stages of a cognitive task. *Neuroimage*, 23(2), 752-763.
- Roebroeck, A., Formisano, E., & Goebel, R. (2005). Mapping directed influence over the brain using Granger causality and fMRI. *Neuroimage*, 25(1), 230-242.
- Roelfsema, P. R., Lamme, V. A. F., & Spekreijse, H. (1998). Object-based attention in the primary visual cortex of the macaque monkey. *Nature*, 395(6700), 376-381.
- Rubin, D. B. (1974). Estimating Causal Effects of Treatments in Randomized and Nonrandomized Studies. *Journal of Educational Psychology*, 66(5), 688-701.
- Sarter, M., Berntson, G. G., & Cacioppo, J. T. (1996). Brain imaging and cognitive neuroscience. Toward strong inference in attributing function to structure. *Am Psychol*, 51(1), 13-21.
- Sawamura, H., Orban, G. A., & Vogels, R. (2006). Selectivity of neuronal adaptation does not match response selectivity: a single-cell study of the FMRI adaptation paradigm. *Neuron*, 49(2), 307-318.
- Schacter, D. L., Buckner, R. L., Koutstaal, W., Dale, A. M., & Rosen, B. R. (1997). Late onset of anterior prefrontal activity during true and false recognition: an event-related fMRI study. *Neuroimage*, 6(4), 259-269.
- Skudlarski, P., Constable, R. T., & Gore, J. C. (1999). ROC analysis of statistical methods used in functional MRI: individual subjects. *Neuroimage*, 9(3), 311-329.
- Smith, S., Jenkinson, M., Beckmann, C., Miller, K., & Woolrich, M. (2007). Meaningful design and contrast estimability in FMRI. *Neuroimage*, 34(1), 127-136.
- Stark, C. E., & Squire, L. R. (2001). When zero is not zero: the problem of ambiguous baseline conditions in fMRI. *Proc Natl Acad Sci U S A*, 98(22), 12760-12766.
- Sternberg, S. (2001). Separate modifiability, mental modules, and the use of pure and composite measures to reveal them. *Acta Psychol (Amst)*, 106(1-2), 147-246.
- Summerfield, C., Greene, M., Wager, T., Egner, T., Hirsch, J., & Mangels, J. (2006). Neocortical connectivity during episodic memory formation. *PLoS Biol*, 4(5), e128.
- Tohka, J., Foerde, K., Aron, A. R., Tom, S. M., Toga, A. W., & Poldrack, R. A. (2007). Automatic independent component labeling for artifact removal in fMRI. *Neuroimage*.
- Vazquez, A. L., Cohen, E. R., Gulani, V., Hernandez-Garcia, L., Zheng, Y., Lee, G. R., et al. (2006). Vascular dynamics and BOLD fMRI: CBF level effects and analysis considerations. *Neuroimage*, 32(4), 1642-1655.
- Wager, T. D., Davidson, M. L., Hughes, B. L., Lindquist, M. A., & Ochsner, K. N. (2008). Prefrontal-subcortical pathways mediating successful emotion regulation. *Neuron*, 59(6), 1037-1050.

- Wager, T. D., Hernandez, L., Jonides, J., & Lindquist, M. (2007). Elements of functional neuroimaging. In J. T. Cacioppo, L. G. Tassinary & G. G. Berntson (Eds.), *Handbook of Psychophysiology* (4th ed., pp. 19-55). Cambridge: Cambridge University Press.
- Wager, T. D., Jonides, J., & Smith, E. E. (2006). Individual differences in multiple types of shifting attention. *Memory & Cognition*, 34(8), 1730-1743.
- Wager, T. D., Jonides, J., Smith, E. E., & Nichols, T. E. (2005). Towards a taxonomy of attention-shifting: Individual differences in fMRI during multiple shift types. *Cogn Affect Behav Neurosci*, 5(2), 127-143.
- Wager, T. D., Lindquist, M., & Hernandez, L. (in press). Essentials of functional neuroimaging. In J. Cacioppo & G. G. Berntson (Eds.), *Handbook of Neuroscience for the Behavioral Sciences*.
- Wager, T. D., & Nichols, T. E. (2003). Optimization of experimental design in fMRI: a general framework using a genetic algorithm. *Neuroimage*, 18(2), 293-309.
- Wager, T. D., Scott, D. J., & Zubieta, J. K. (2007). Placebo effects on human mu-opioid activity during pain. *Proc Natl Acad Sci U S A*, 104(26), 11056-11061.
- Wager, T. D., Vazquez, A., Hernandez, L., & Noll, D. C. (2005). Accounting for nonlinear BOLD effects in fMRI: parameter estimates and a model for prediction in rapid event-related studies. *Neuroimage*, 25(1), 206-218.
- Williams, D. S., Detre, J. A., Leigh, J. S., & Koretsky, A. P. (1992). Magnetic resonance imaging of perfusion using spin inversion of arterial water. *Proc Natl Acad Sci U S A*, 89(1), 212-216.
- Wilson, J. L., & Jezzard, P. (2003). Utilization of an intra-oral diamagnetic passive shim in functional MRI of the inferior frontal cortex. *Magn Reson Med*, 50(5), 1089-1094.
- Worsley, K. J., & Friston, K. J. (1995). Analysis of fMRI time-series revisited--again. *Neuroimage*, 2(3), 173-181.

Table 1.

	Higher field strength	shorter TE	Smaller voxels	longer TR	Accelerated imaging
Brain coverage			-	+	+
Contrast/noise	+	-	-	+	-
BOLD signal/noise	+	-	-	+	-*
Spatial resolution			+	+	+
Temporal resolution			-	-	+
Susceptibility artifact reduction	-	+		-	+
Resistance to motion/physiological artifacts	-			-	+

Note. Analysis options are shown in columns, and their effects shown in rows. Characteristics listed in rows are advantageous, but there are inherent tradeoffs in the data acquisition process that make it impossible to have one's cake and eat it too. Minus signs (-) mean "disadvantageous" and plus signs (+) mean "advantageous". *: there is reduced signal to noise but more images, so the benefits and costs may balance.

Figure Captions

Figure 1. Different kinds of MRI-based images. (A) The same slice of brain imaged with two different structural imaging methods: T1-weighted (left), which appear brighter for tissues with relatively little water content, and T2-weighted, with brighter signal in tissues with high T2 decay. (B) Three kinds of functional images: echo-planar (left) and spiral (center) Blood Oxygen Level Dependent (BOLD) images, and a perfusion-related image acquired with Continuous Arterial Spin Labeling. (C) Diffusion tensor imaging allows researchers to measure directional diffusion and reconstruct the fiber tracts of the brain. This provides a way to study how different brain areas are connected.

Figure 2. Parametric variation designs can provide more convincing inferences because they incorporate more specific predictions about a particular process. Parametric “modulators” can be task conditions such as task difficulty (Dagher et al., 1999) (left), subjective value (Hare et al., 2009) (right) or other measured or experimental variables.

Figure 3. The optimal number of subjects for a fixed total cost varies systematically based on the ratio of within-subjects error and between-subjects differences. (A) Statistical power (on the y-axis) as a function of the sample size/scan time per subject tradeoff (on the x-axis). The lines in each graph show different levels of between-subjects variation (inter-individual differences) and the three panels show three different levels of within-subjects noise. When within-subjects noise is low (left panel), it is always better to run more subjects, even when very little time is spent imaging each subject. However, as the within-subjects noise increases (middle and right panels), it becomes optimal to increase the acquisition time for each participant. As the inter-individual differences increase (shown by increasingly light-shaded lines), the optimal ratio shifts towards more subjects with less time per subject. The points marked on the graphs show approximately optimal tradeoff points. Of course, the usefulness of these curves for planning future studies depends on having reasonable estimates for within-subjects and between-subjects error, which vary depending on the scanner noise, experimental design, paradigm, and brain region. (B) An example of the sample size/run time tradeoff from a working memory study in our laboratory. We identified regions of interest showing reliable increases or decreases relative to rest (left panel), and then calculated the average effect size and power across voxels for an independent 3-back vs. 2-back comparison. The average effect size was $d = 0.23$. The average within-subject standard deviation per observation was 66, and the average between-subject standard deviation was 2.48. The optimal allocation of 60 scan hours for these values was $N = 62$, with 28 min of functional time per subject in a single session (in our calculations, 30 min per session was allocated for subject setup and structural imaging.) The power curve is shown in the right panel.

Figure 4. The two most common classes of experimental design are block designs and event-related designs. In a block design (top) experimental conditions are separated into extended time intervals, or blocks, of the same type. In an event-related design (bottom) the stimulus consists of short discrete events whose timing and order can be randomized.

Figure 5. Influences on T_2^* -weighted signal in BOLD fMRI imaging. Courtesy Dr. Doug Noll.

Figure 6. (A) Empirical hemodynamic response functions (HRFs) from primary visual and motor cortices, adapted from Lindquist et al. (Lindquist, Zhang, Glover, & Shepp, 2008) Data were sampled at a high time-resolution using a recently developed acquisition technique (100 ms, with whole-brain coverage at 12 mm spatial resolution), permitting visualization of fine-grained details of the HRF, including the initial dip in signal due to blood oxygenation decreases. Participants saw a contrast-reversing checkerboard (visual) for 100 ms and made a button-press response an average of 250 ms later. The signal in the visual cortex proceeds the signal in the motor cortex throughout the duration of the HRF. (B) In an fMRI experiment with two conditions (A and B), the stimulus function is convolved with a canonical HRF to obtain two sets of predicted BOLD responses. The responses are placed into the columns of a design matrix X and used to compute whether there is significant signal corresponding to the two conditions in a particular time course.

Figure 7. A data analysis pipeline that shows the major processing steps for a functional MRI study. This is the pipeline currently used in our laboratory.

Figure 8. Design matrices using basis sets. The same ‘neural’ indicator vector can be convolved with multiple canonical waveforms and entered into multiple columns of X for a single event type. These reference waveforms are basis functions, and the predictors for an event type constructed using different basis functions can combine linearly to better fit the evoked BOLD responses. The ability of a basis set to capture variations in hemodynamic responses depends on both the number and shape of the reference waveforms. There is a fundamental tradeoff between flexibility to model variations and power. This is because each parameter is estimated with error, and flexible models can tend to model noise and thus produce noisier parameter estimates. For a critical evaluation of various basis sets see Lindquist and Wager (2007b) and Lindquist, Loh, Atlas, and Wager (2008c). Here, the top panel shows an idealized data time series and event onsets (vertical lines). The panels below depict regressors for three different kinds of models (with time on the x-axis): One assuming a canonical hemodynamic response (top), one using a common basis set with a canonical waveform and two derivatives (middle), and a finite impulse response (FIR) basis set that assumes almost nothing about the shape of the response. At the right is shown an idealized, but non-canonical, hemodynamic response (solid lines), with the fitted response for each model

superimposed (dashed gray lines).

Figure 9. Example of a parametric modulator design. The top panel shows an idealized data time series and event onsets (vertical lines). The strength or intensity of a psychological process varies from trial to trial, as depicted by the relative heights of the lines. The middle panels show the predictors for a single event type. The upper panel shows the regressor for the average response, and the lower panel shows the modulatory effect, created by convolving with a variable-amplitude “neural” (impulse) function after centering the amplitude. Centering makes the modulator more orthogonal to the mean response, and gives each parameter a unique psychological interpretation (mean trial effect and effects of the modulating variable.)

Figure 10. (A) Power curves - calculated for effect sizes of 0.5, 1, and 2 - for a whole-brain search with 200,000 voxels and FWE correction at $p < .05$ using the Bonferroni method. (B) Same results using nonparametric permutation testing which takes into account the spatial smoothness in the data.

Fig 1. T1 and T2 images & DTI

Figure 1.

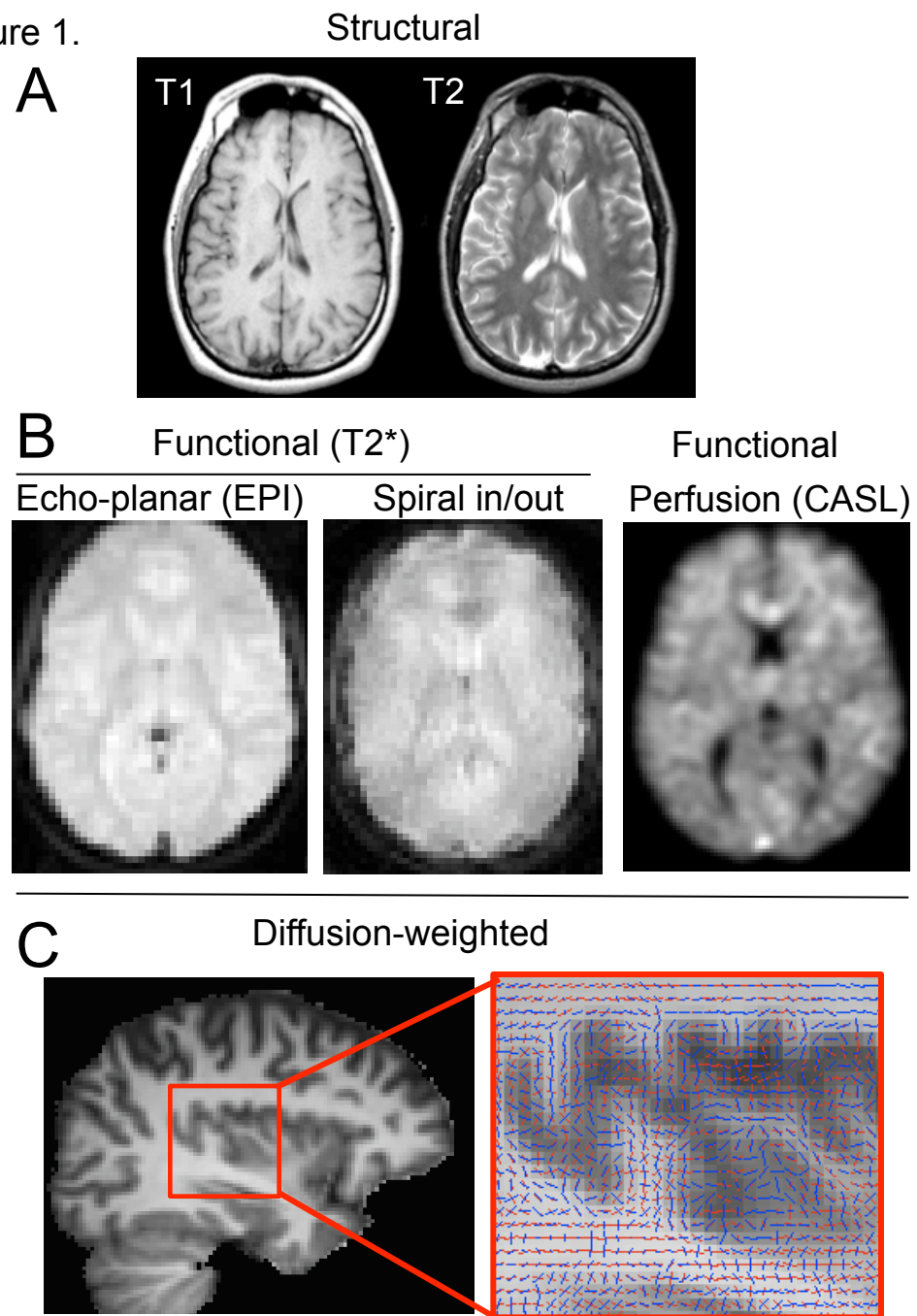
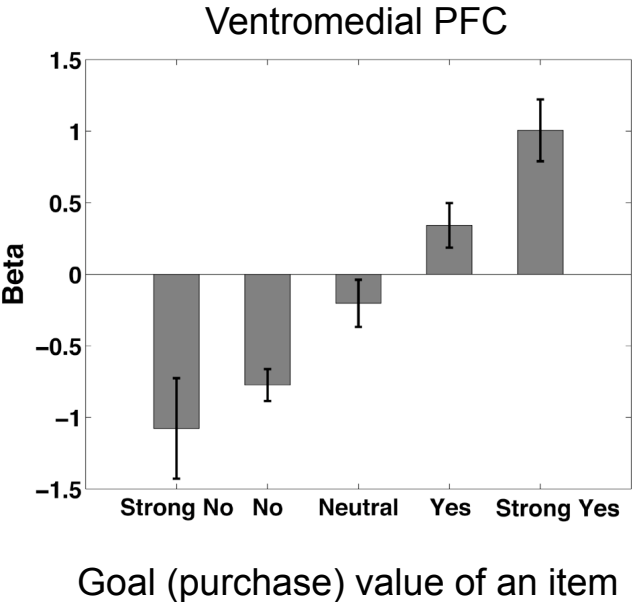
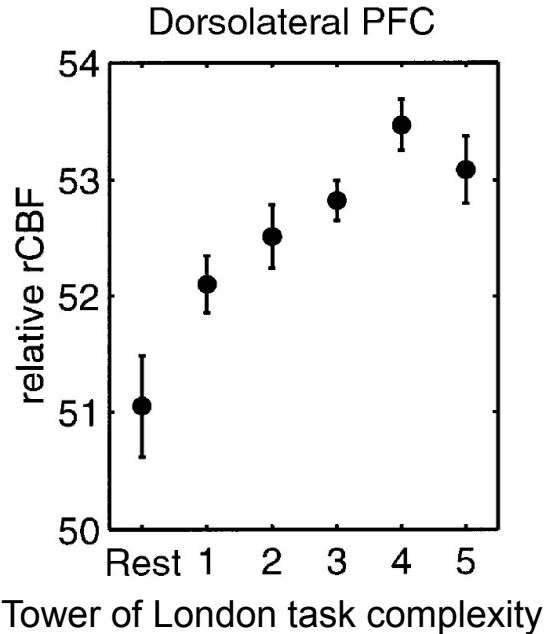
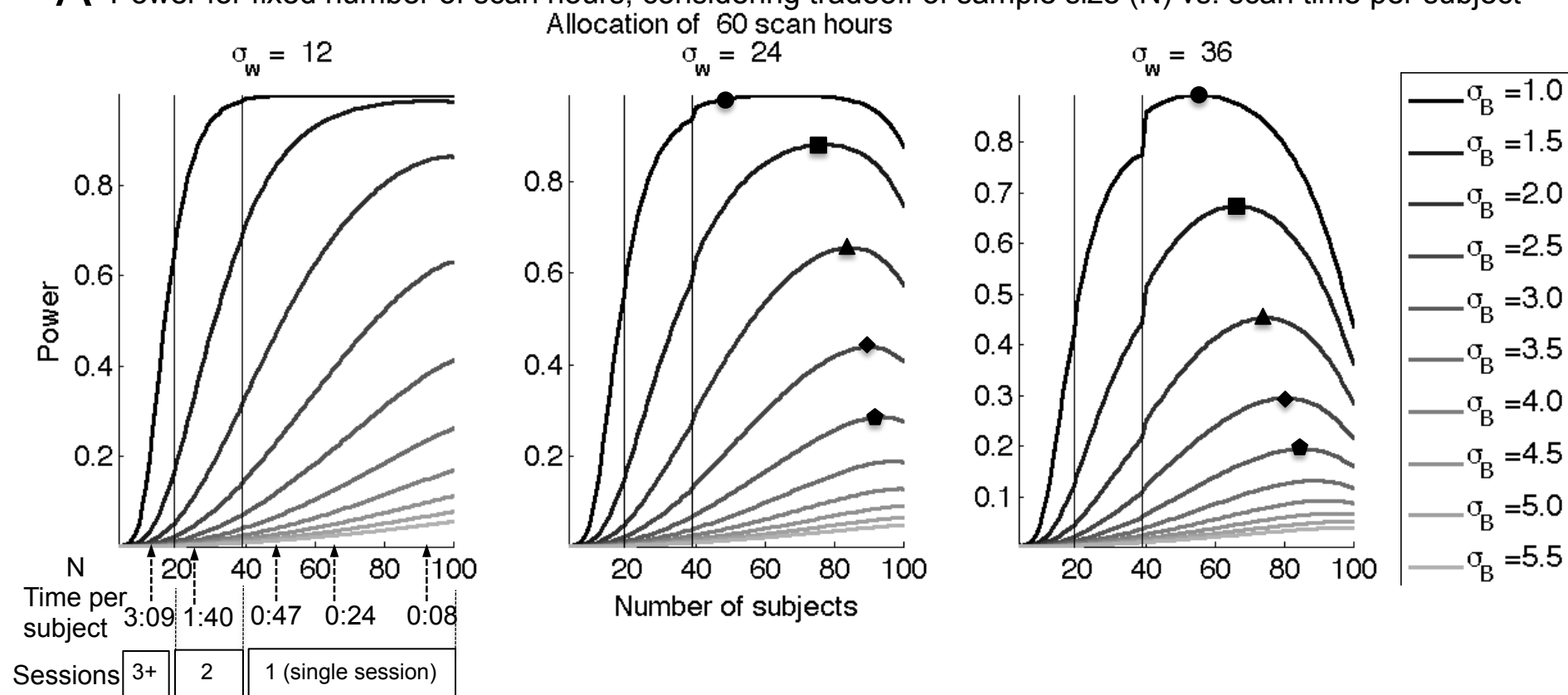


Figure 2.



A Power for fixed number of scan hours, considering tradeoff of sample size (N) vs. scan time per subject



B Empirical example: working memory

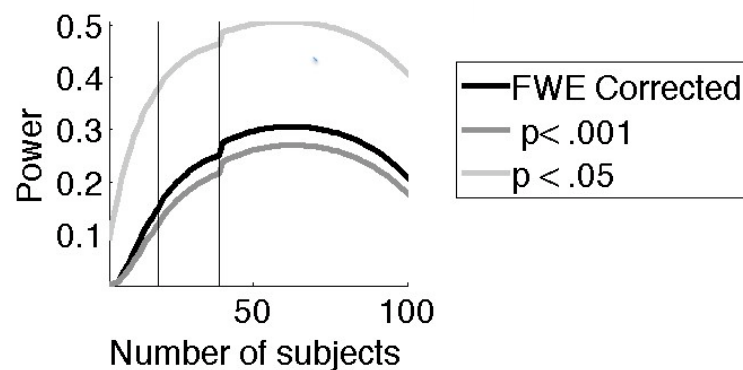


Figure 3.

Fig 8. Building a design matrix

Figure 4.

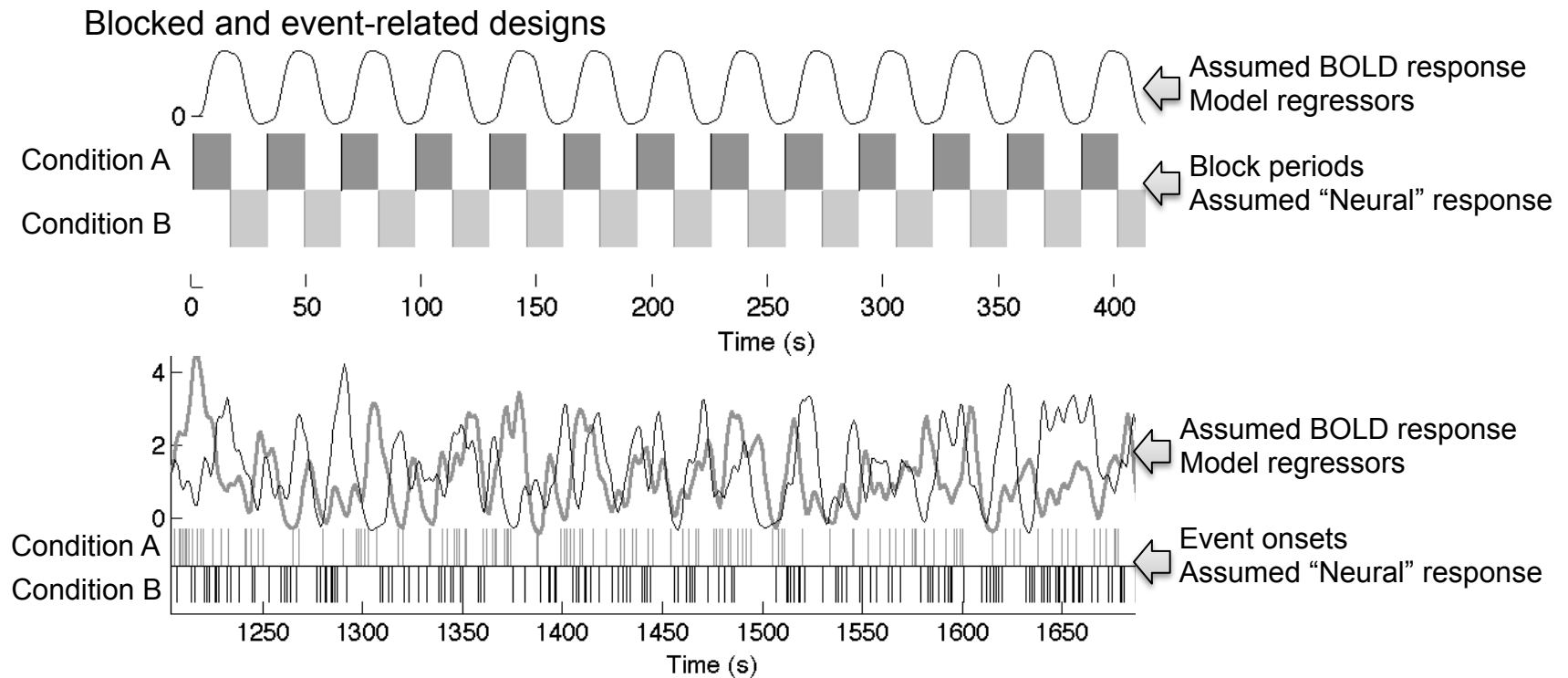
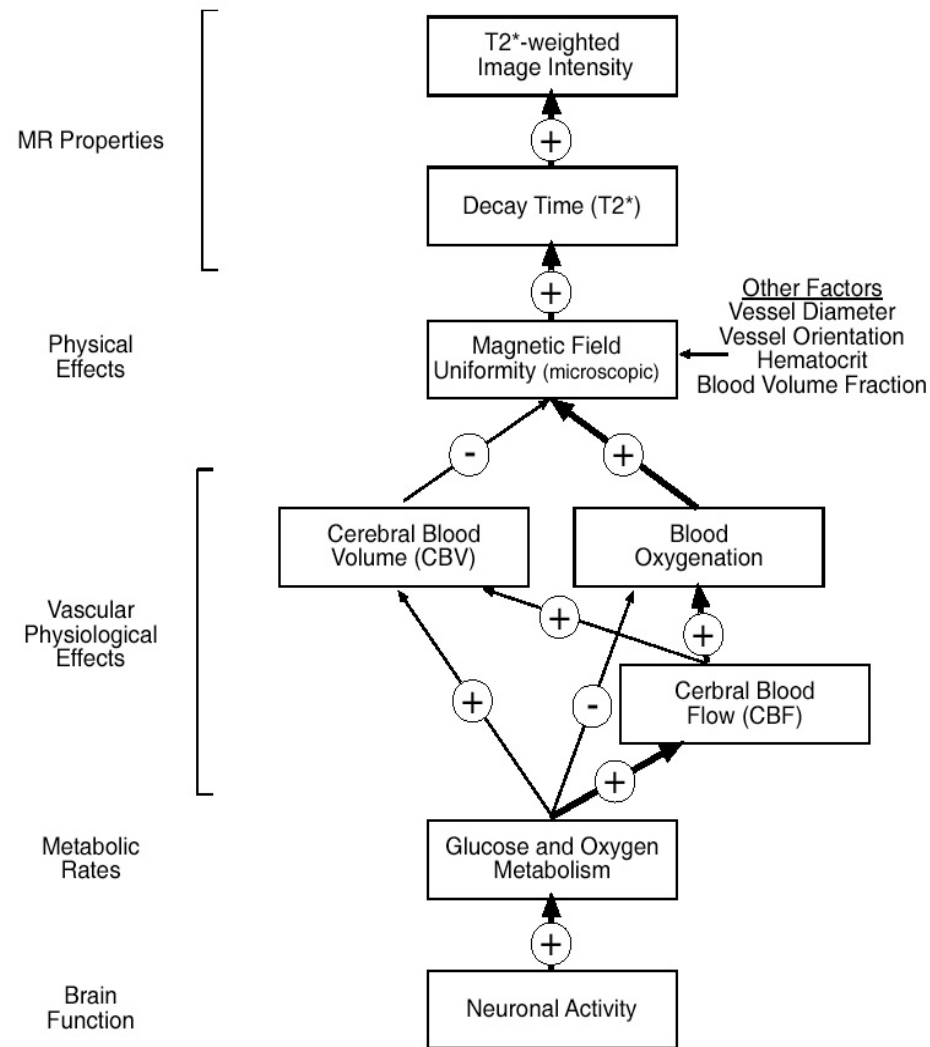


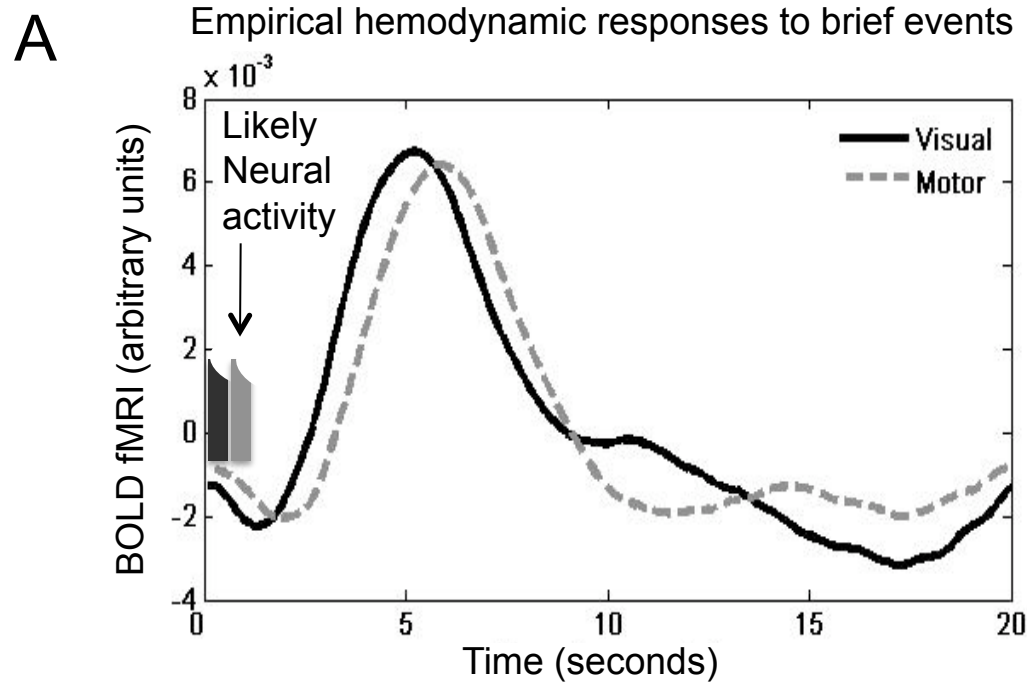
Fig. 3. BOLD physiology

Figure 5.

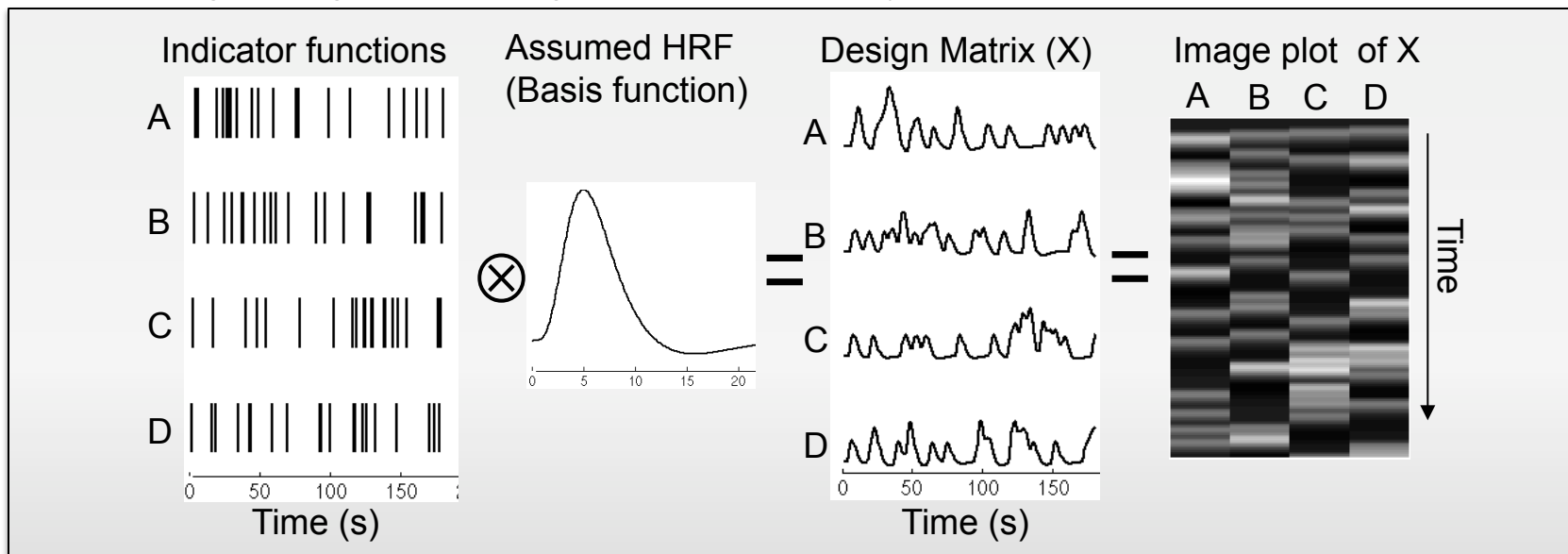


Comparing the HRFs

Figure 6.



B Constructing a design matrix using an assumed hemodynamic response shape



Overview flowchart

Figure 7.

A standard processing pipeline

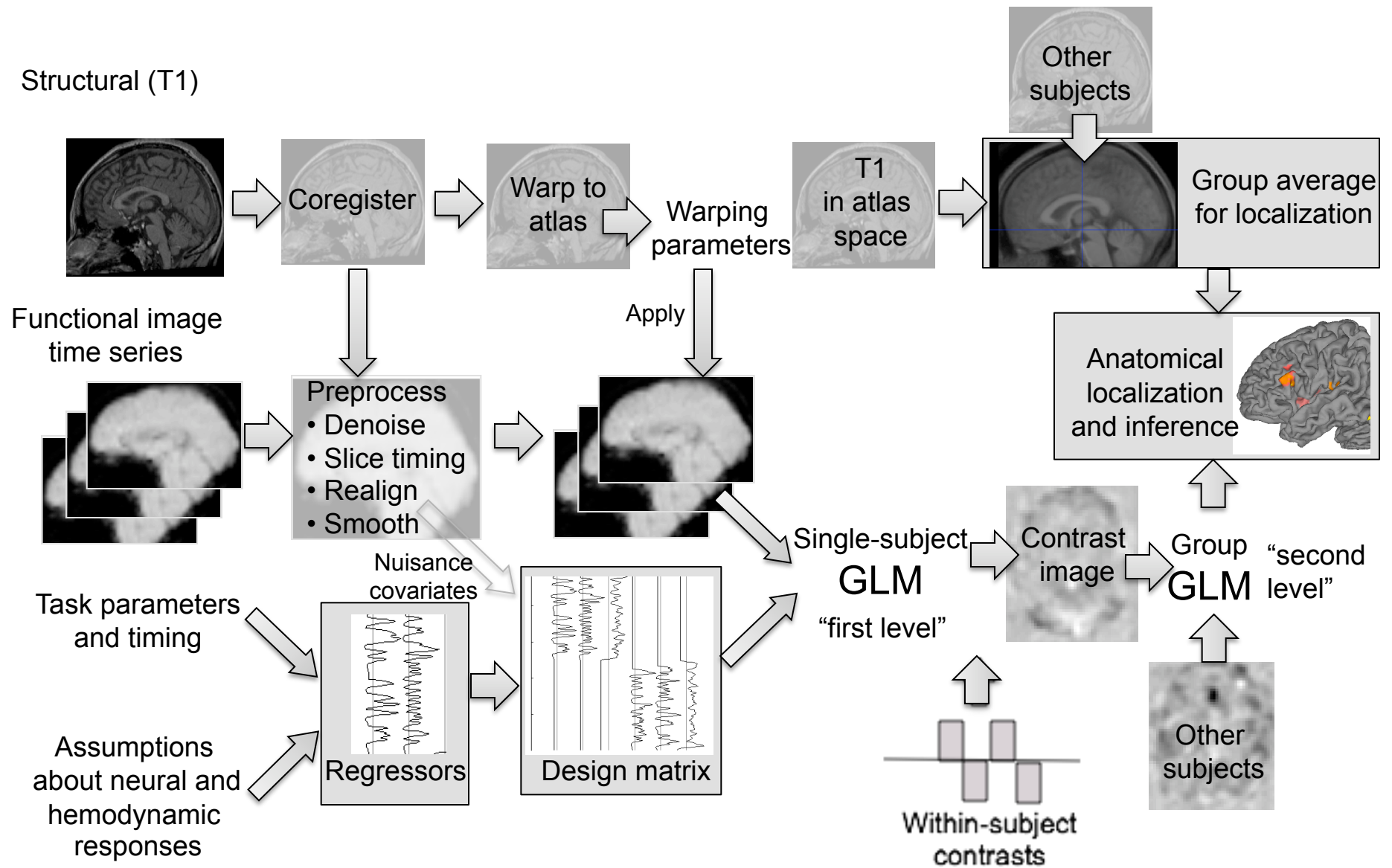


Figure 8.

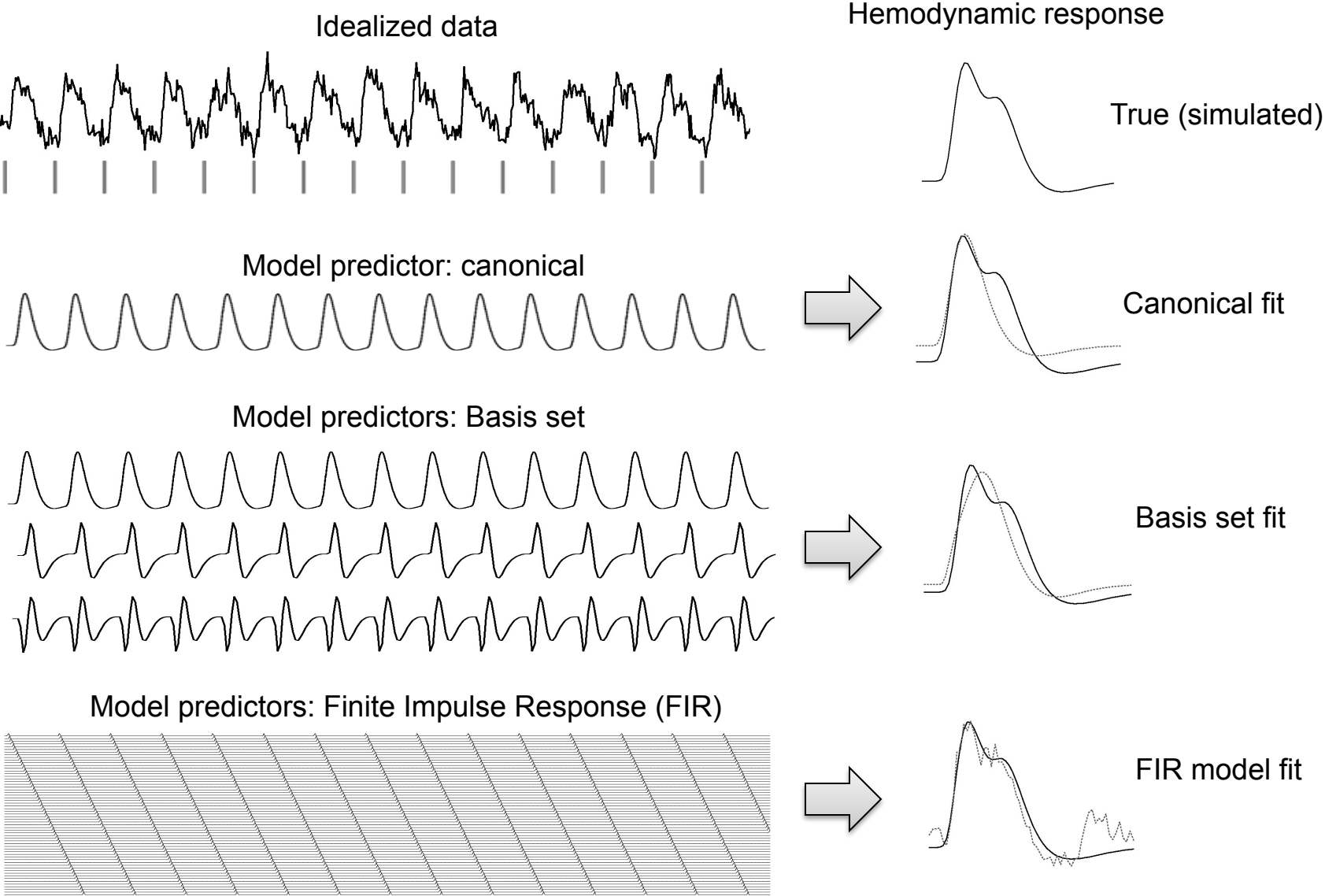


Figure 9.

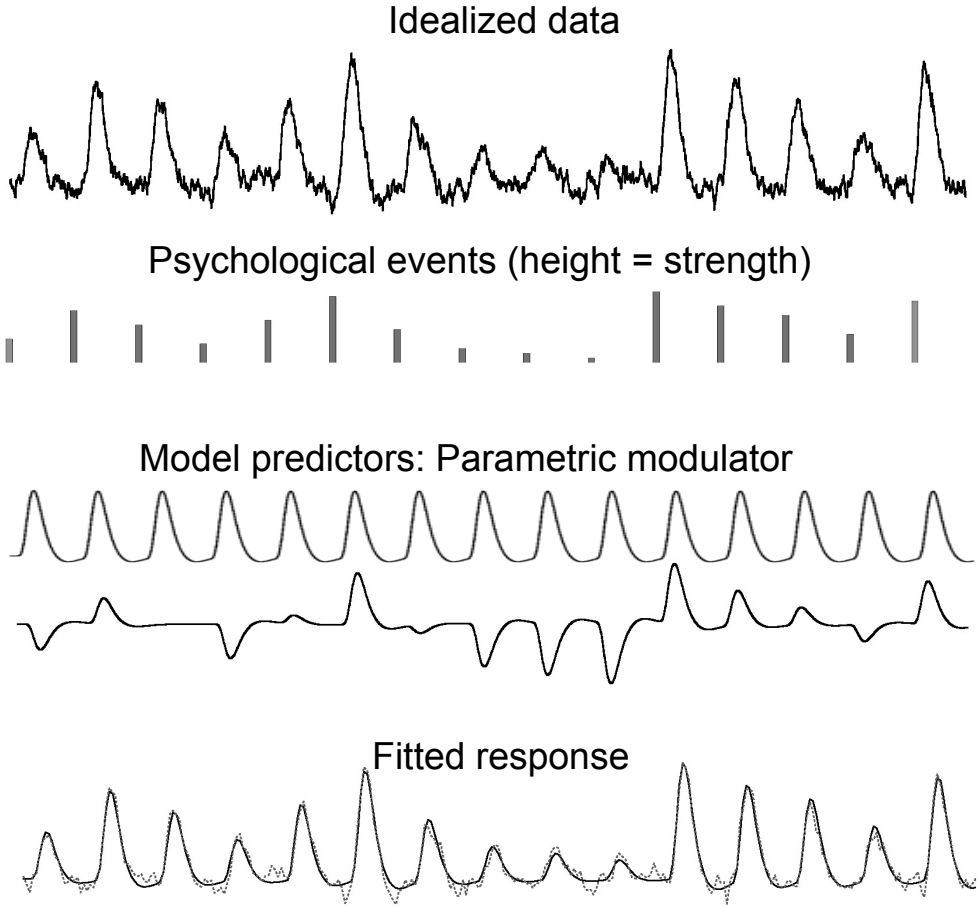


Figure 10.

