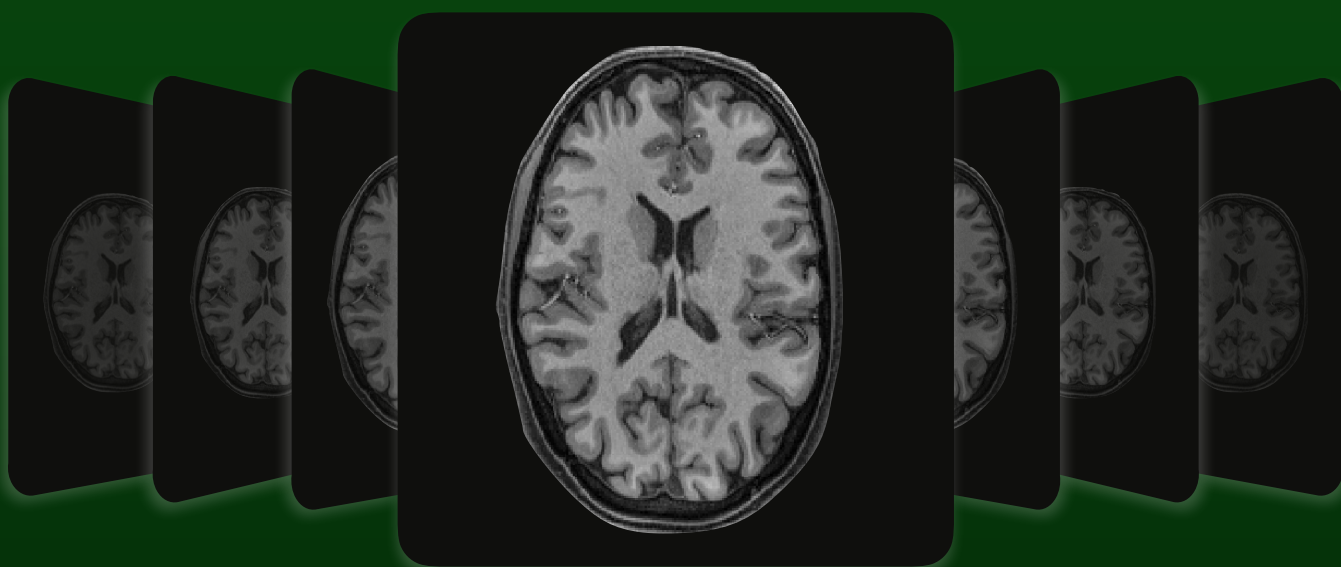




Short introduction to
MRI Physics
for Neuroimaging



Michael Chappell
Thomas Okell
Mark Jenkinson

Series editors:
Mark Jenkinson and Michael Chappell

PRIMER
APPENDIX

List of Primers

Series Editors: Mark Jenkinson and Michael Chappell

Introduction to Neuroimaging Analysis

Mark Jenkinson

Michael Chappell

Introduction to Perfusion Quantification using Arterial Spin Labelling

Michael Chappell

Bradley MacIntosh

Thomas Okell

Introduction to Resting State fMRI Functional Connectivity

Janine Bijsterbosch

Stephen Smith

Christian Beckmann

List of Primer Appendices

Short Introduction to Brain Anatomy for Neuroimaging

Short Introduction to MRI Physics for Neuroimaging

Short Introduction to MRI Safety for Neuroimaging

Copyright

Portions of this material were originally published in *Introduction to Neuroimaging Analysis* authored by Mark Jenkinson and Michael Chappell, and have been reproduced by permission of Oxford University Press: <https://global.oup.com/academic/product/introduction-to-neuroimaging-analysis-9780198816300>. For permission to reuse this material, please visit: <http://www.oup.co.uk/academic/rights/permissions>.

Preface

This text is one of a number of appendices to the Oxford Neuroimaging Primers, designed to provide extra details and information that someone reading one of the primers might find helpful, but where it is not crucial to the understanding of the main material. This appendix specifically addresses the physical principles that underpin MRI, as it is used in neuroimaging. In it we seek to go into more detail than we might in one of the primers, such as the Introduction to Neuroimaging Analysis, for those who want to understand more about what is actually going on when we acquire MRI images of the brain. In turn, this appendix also provides more context on some of the decisions that are made when designing the most suitable acquisition for any given MRI modality for a particular imaging application or study design.

We hope that this appendix, in keeping with the series as a whole, will be an accessible introduction to the topic of MRI physics for those without a background in the physical sciences. Hence, we have concentrated on concepts rather than delving into any detailed mathematics of image acquisition. However, we also hope it is a good introduction to physical scientists meeting MRI for the first time, perhaps before going on to more technical texts, such as those we include in the Further Reading at the end of the Appendix.

This appendix contains several different types of boxes in the text that are designed to help you navigate the material or find out more information for yourself. To get the most out of this appendix, you might find the description of each type of box below helpful.

Boxes

These boxes contain more technical or advanced descriptions of some topics covered in this appendix. None of the material in the rest of the appendix assumes that you have read these boxes, and they are not essential for understanding any of the other material. If you are new to the field and are reading this appendix for the first time, you may prefer to skip the material in these boxes and come back to them later.

Box 4.1: Frequency and

So far we have considered

Further Reading

At the end, we include a list of suggestions for further reading, including both articles and books. A brief summary of the contents of each suggestion is included, so that you can choose the most relevant references for you. None of the material in this appendix assumes that you have read anything from the further reading. Rather, this list suggests a starting point for diving deeper, but is by no means an authoritative survey of all the relevant material you might want to consult.

FURTHER READING

■ Huettel, S. A., Song, A. ...

Whilst the principles of MRI physics are well established and thus the material in this appendix will, we hope, be relevant for many years to come. Advances in the field of MRI acquisition continue and new techniques that acquire images with higher resolution, more quickly and with new information, appear all the time. Hence, all we hope for as authors is that this will be a useful introduction to what is a large and fascinating field of research that extends well beyond purely neuroimaging applications.

Michael Chappell, Thomas Okell, and Mark Jenkinson

Contents

List of Primers	i
List of Primer Appendices	i
Preface	ii
Contents	iii
1 Introduction	1
2 M is for Magnetic	2
3 R is for Resonance	2
3.1 Excitation	3
4 I is for Imaging	4
4.1 Slab and slice selection	5
4.2 Readout	6
4.3 Inhomogeneity	8
4.4 Chemical Shift	8
5 Contrast	9
5.1 Relaxation	9
5.2 Inversion	15
5.3 Diffusion	17
5.4 BOLD	18
5.5 Contrast agents	19
6 Fast Imaging	20
6.1 Echo Planar Imaging (EPI)	21
6.2 In-Plane Acceleration (SENSE, GRAPPA, etc.)	22
6.3 Multi-Slice Acceleration (Simultaneous Multi-Slice, Multiband, etc.)	22
7 Quantitative MRI	24

1 Introduction

It is not necessary to understand everything there is to know about MR physics and how an MRI scanner works to successfully run neuroimaging experiments. However, some basics are necessary to be able to design good experiments and interpret results carefully. In this Short Introduction we will outline the basics of MR physics and how MRI scanners work, which builds upon the basic principles that you will also find in the Introduction to Neuroimaging Analysis primer. The key components of the sort of MRI scanner you might meet when doing neuroimaging are shown in Figure 1.1.

The fundamental principle of MRI is that certain atomic nuclei act like tiny bar magnets and interact with magnetic fields in a way that allows us to both measure and manipulate their magnetic state. Since this interaction is with nuclei it is quite reasonable to call MRI a 'nuclear' technique, and the principles will be familiar to any chemist who knows about Nuclear Magnetic Resonance. However, the medical imaging community has dropped the 'nuclear' part of the name to avoid confusion with Nuclear Medicine, where radioactive elements are used. In MRI the interaction with the nuclei does no damage to them or any of the biological processes involving the molecules that they are contained within.

There are a number of different elements that possess the relevant properties that means they can exhibit magnetic resonance. For most MRI applications, it is the hydrogen nuclei within water molecules that are principally targeted, as they are so abundant, though hydrogen in fatty tissues also shows up in some types of scans. A hydrogen nucleus is simply a single proton, thus you may find people referring to proton MRI. Although imaging water may sound uninteresting, it is actually far from it, as there are many different properties of the water molecules that MRI can detect and

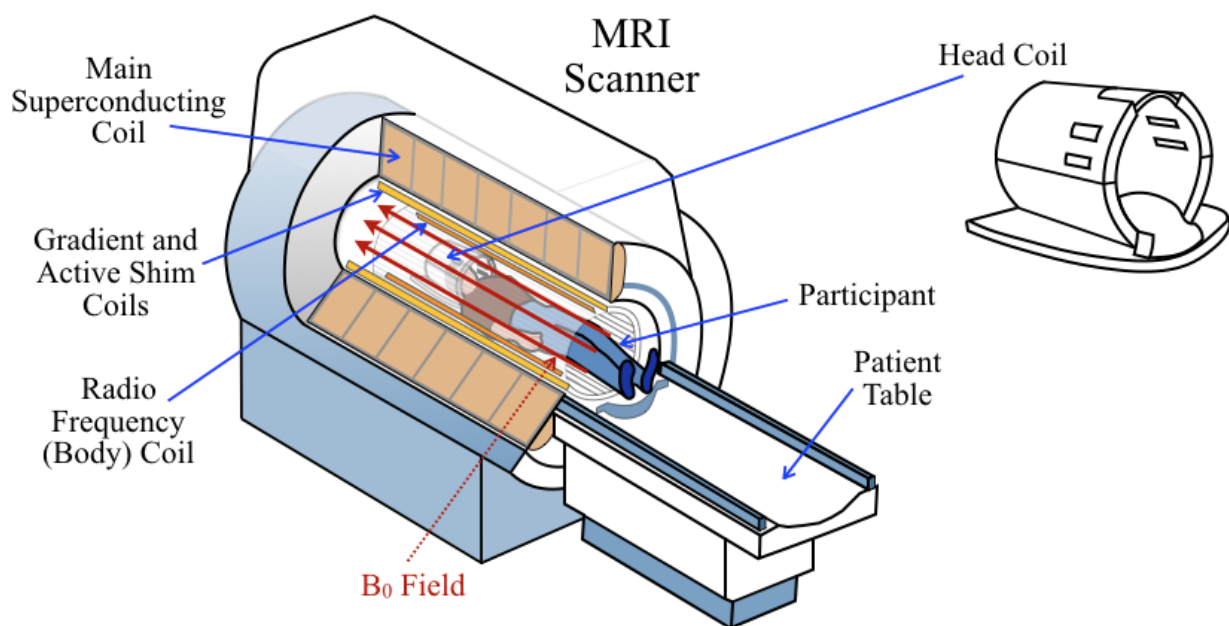


Figure 1.1: Schematic of an MRI scanner, showing main coils and the B_0 field in relation to the subject inside the scanner. The head coil (top right) is placed around the subject's head prior to them being moved into the centre of the scanner (the bore).

this is what gives it the ability to examine the brain in so many different ways. One of the reasons for the success of MRI in neuroimaging is the number of imaginative ways that MR physics principles have been used to exploit the relationship between water and its environment in the brain.

2 M is for Magnetic

To interact with the hydrogen nuclei a very strong static magnetic field is required. We have already noted that the hydrogen nuclei act like tiny bar magnets, thus they have an orientation, just like a bar magnet has both a north and a south pole. In the absence of any external magnetic field all the hydrogen nuclei point in different directions, so there is no overall orientation of the total magnetization from summing all their individual contributions. In this case there would be no signal to detect. It is only in a very strong magnetic field that the tiny bar magnets of the nuclei will tend to point in the same direction (along this main field). This strong magnetic field, known as the B_0 field, is established by a large superconducting coil that is always on and cooled with liquid helium. It is this field that defines the strength of the scanner, as reported in units of *tesla* (or T); for example, a 3T scanner is one where the B_0 field is 3 tesla. To give you a sense of how 'strong' this is, the large electromagnets used in scrapyards to pick up cars, are approximately 1T, whereas the magnetic field you are in all the time, generated by the Earth's magnetic core, is between 25 and 65 *micro*-tesla.

Even in a strong magnetic field the overall, net, magnetization from all the different hydrogen nuclei is still very small. Whilst aligning with the field is more favourable (lower energy state) at normal room temperature many of the hydrogen nuclei have sufficient energy to be pointing in other directions and the bias in favour of aligning with the field is very low (on the order of one in one million). Thus, even to start with, we are trying to detect a small effect with MRI.

3 R is for Resonance

A further important property of the net magnetization of the nuclei in a strong field is that if we can shift it out of alignment with B_0 it will *precess* or rotate around the axis of the B_0 field, making a circular path around it - see Figure 3.1. The frequency of its rotation is proportional to the strength of the magnetic field it experiences (i.e., the B_0 field plus any other fields - see later). The relationship between the frequency of rotation and the field strength is governed by the Larmor

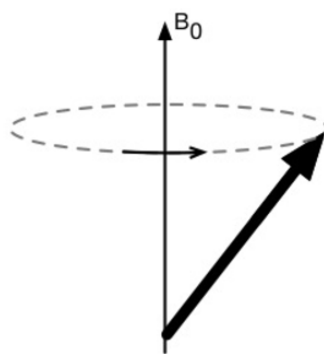


Figure 3.1: In a strong field the net magnetization of all the hydrogen nuclei precess around the direction of the magnetic field at a fixed angle.

equation: the frequency of precession for a hydrogen nucleus is 42.58 MHz times the main B_0 field strength in tesla. This is 63.9 MHz for a 1.5T scanner, 127.7 MHz for a 3T scanner and 298.1 MHz for a 7T scanner. This precession frequency (or Larmor frequency) is thus a very specific and well defined property of the system and thus we can describe it as the *resonant frequency*.

Due to the way that the coils are arranged, we only detect changes in the net magnetization (the sum of magnetization from all the nuclei) in the plane normal (or orthogonal) to the B_0 axis.¹ We call this the *transverse* or axial plane, while the B_0 axis is referred to as the *longitudinal* axis. As we have already noted, if we manipulate the hydrogen nuclei such that the net magnetization is tipped away from the direction of the B_0 field, we find that the net magnetization now precesses at the Larmor frequency - Figure 3.1. This net magnetization, which is varying with time, in turn generates its own time-varying magnetic fields, which can then induce time-varying currents in a nearby coil that we can measure. Thus the head (or body) coil found in the scanner is there to allow us to measure signals from the hydrogen nuclei precessing at the Larmor frequency, and they are specifically built to operate at that frequency. As noted above, these signals are in the tens of MHz range which are normally called radio frequencies (RF), being associated with the transmission of radio waves in communications.

3.1 Excitation

We have established that we can measure a signal from a collection hydrogen nuclei associated with their net (or overall) magnetization, but this is only possible once this net magnetization is no longer aligned (parallel) with the B_0 field. If this happens the net magnetization precesses at the Larmor frequency and we can 'listen' to that using a coil. Thus a vital step for generating an MR image is perturbing the magnetization to move it away from the longitudinal axis - a process called *excitation*. Here we exploit the resonance property: if we take a coil and put into it a current that varies with time at the Larmor frequency, it will create a magnetic field that oscillates at the same frequency. This is the *B_1 field*. We can orient the coil such that this field is always perpendicular to the B_0 field. The hydrogen nuclei will now experience the (static) B_0 field plus a B_1 field (also called the *RF field*) which is tuned to their resonant, Larmor, frequency. The effect of this field is to change the orientation of their magnetization away from B_0 and into the transverse plane, to achieve the effect shown in Figure 3.2. The net magnetization now not only points in a different direction to B_0 , but it will also be precessing; thus if we stop applying B_1 and instead 'listen' to the signal generated in the coil we can now measure the magnetization of the hydrogen nuclei. Note that the excitation process will only be successful if the frequency of the applied B_1 field exactly matches the resonant frequency of the magnetisation, hence the importance of resonance in MRI.

An analogy for this is the classic example of resonance: a child on a swing. To get the swing going you have to push, but to get the swing to go higher and higher you have to keep pushing, but it only makes sense to push when the swing is near you and at the point where it is heading away from you. You time your pushes with the timing of the swing - any attempt to push at a different rate, or frequency, will not reinforce the natural motion of the swing, and result in an unhappy child! If you stop, you can then watch the motion of the swing, which continues with the same timing (i.e.,

¹ We could attempt to measure the magnetization aligned with B_0 , but it is so small compared to the main magnetic field that we have very little chance of measuring it accurately. It is much more convenient to measure the rotating component in the plane normal to B_0 .

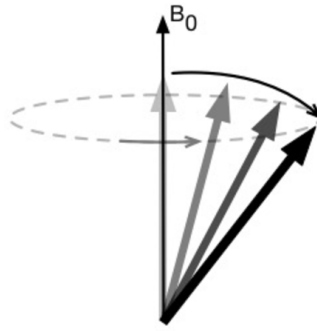


Figure 3.2: Excitation involves changing the orientation of the net magnetization, to move it away from axis of the B_0 field.

frequency) as this timing being dictated by the swing and not you. The B_1 field is very much like the adult, it pushes with just the right timing to force the net magnetization out of alignment with B_0 . In practice, to make the process more efficient the B_1 field is in fact pushing all the time, but in perfect synchronisation with the rotation of the net magnetization.

Since excitation exploits the resonance frequency of the hydrogen nuclei, the process of pushing them out of alignment with B_0 is relatively 'easy', and so a much smaller field is required to do this compared to the B_0 field, and therefore we can use conventional coils and not supercooled ones. These coils are known as **RF coils** and the application of the fields produced by these coils for excitation (or other manipulations) are often called **RF pulses**, as they are only applied for very short periods of time.

4 I is for Imaging

The observed MR signal is made up of contributions from all the hydrogen nuclei inside the bore of the scanner, all summed together. In order to determine where the signals are coming from in space (i.e., the location within the head) we need a way of separating out the individual contributions to this summed signal. The way this is done in MRI is to use the fact that the resonant frequency, and thus the frequency of the signal, depends on the field strength. Since it is possible to separate out the contributions of different frequencies in a combined signal (a bit like listening for different instruments in an orchestra) the location of the signal can be determined if the frequency can be linked to the location.

To relate frequency to location we deliberately add extra, carefully controlled, magnetic fields that vary with location. These fields are added while acquiring the signal measurements so that the signal from different locations have different, and known, frequencies. This then allows us to measure how strong the signal is at each frequency and, from that, work out how much signal came from a given spatial location (in practice this involves some complicated mathematics, but this is the basic principle that is used). These extra fields are called **gradient fields** and are created by three different **gradient coils** in the scanner, see Figure 1.1, and their effect is illustrated in Figure 4.1. Importantly, the main effect of the gradient fields is to vary the strength of the B_0 field and not its direction (which is still along the bore of the scanner). Only fairly subtle variations in the field strength are required in practice, of the order of millitesla, making it possible to create gradient

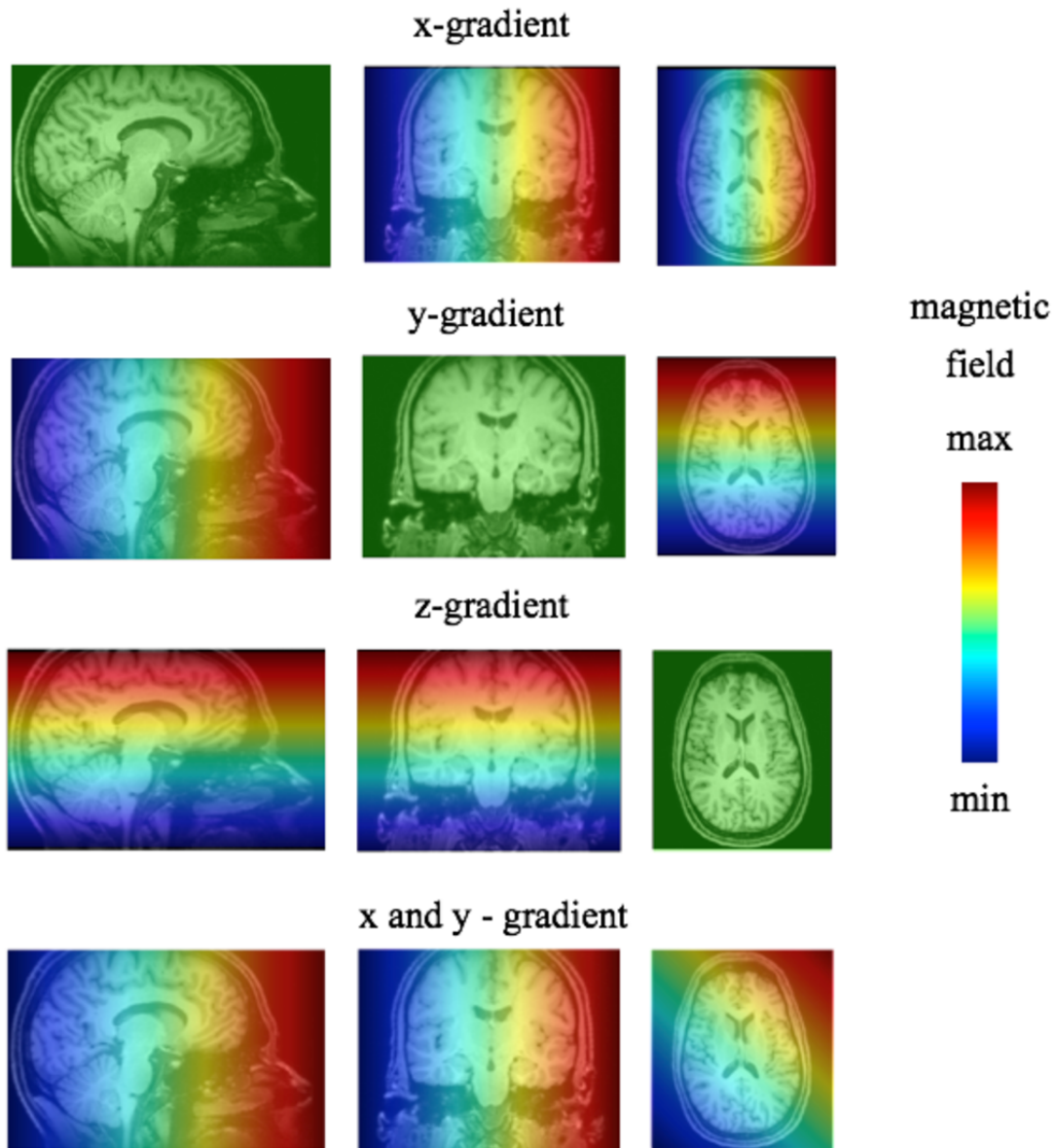


Figure 4.1: Illustration of the three gradient fields; x, y and z. The colors represent the strength of the magnetic field at each location, which is present throughout the head, and makes the signal frequency depend on the spatial location. The first three rows show examples where the gradients are applied independently, although in practice they are also be applied in combinations, as in the example shown in the last row.

fields without the need for superconducting magnets. This also allows them to be switched on and off very rapidly, which is important for the imaging process, as we shall see.

4.1 Slab and slice selection

We can use the fact that excitation is frequency specific, in combination with the gradient fields, to select only a slice or slab of tissue for imaging. The application of a gradient, most commonly for slice-selection in the superior-inferior (z) direction, means that the resonant frequency varies with position according to the strength of the magnetic field, as shown in Figure 4.1. If we apply the excitation using an RF pulse containing a narrow range of frequencies, then only the region of tissue whose resonant frequencies match those within the RF pulse are excited and go on to contribute to the final image. It is the range of frequencies coupled with the variation induced by the applied

gradient (called the slice-selection gradient) that determines the location and thickness of the slab, which could be used to select a volume such as the whole head or a thin slab that might correspond only to a single slice, with multiple slices being collected with separate excitations. There are several schemes used to specify the order that the multiple slices are acquired. The three most common schemes are: (i) sequential - with each slice next to the previous one; (ii) interleaved - starting with odd slices, with gaps in between, and then filling in with even slices later (e.g., slices 1, 3, 5, ..., then 2, 4, 6, ...); (iii) simultaneous multi-slice - where a set of slices (e.g., three slices) are acquired each time and repeated sets used to cover the field of view. See Figure 6.1 for examples of these slice acquisition schemes. Once the excitation has been done we turn the slice-selection gradients back off and move on to the next phase of acquiring spatially resolved information: the readout.

4.2 Readout

Excitation is used to perturb the magnetization, and once this has happened we can then listen to (i.e., measure) the signal created using a coil, which is sometimes the same coil we used to deliver the excitation in the first place. This signal will come from all of the tissue that has been excited, which could be the whole head, or only part of it if we used gradients at the time of excitation to select only a slab or slice of the head. However, unless we do something else we will get an RF signal with a single frequency corresponding to the resonant frequency of the hydrogen nuclei in tissue, which is a measure of the magnetization of the whole slab/slice of tissue. This is the next point where we use the gradients. By applying a gradient to vary the resonant frequency with position in the head, the signal we measure will be a combination of *different frequencies* from *different locations* in the head. Somewhat simplistically, if we made the field stronger toward the right of the head and weaker on the left then the component of the signal we receive at a higher frequency will tell us about the magnetization on the right and the lower frequency component about the magnetization on the left. The signal we receive is a mixture of all of the different frequency components and we will need some mathematics, namely the Fourier Transform, to separate out the different parts and relate them to position. What we have described here is the **Gradient Recalled Echo (GRE)**; also known as just **Gradient Echo** or **GE**). However, the principles are the same for all readout techniques: we send an RF signal in (the excitation), listen to what comes back (the echo), but doing so whilst applying a gradient to encode spatial information.

This simplistic example allows us to distinguish between locations from left to right (or anterior to posterior, or inferior to superior), but no more. What we ultimately want is specific information about every 3D position in the head. This can be achieved by repeating the process above using a different combination of gradients each time, in order to tease apart the contributions from different locations. The main consequence of this is that it will take some time to readout all of the information needed to build the full 3D volume. There are, broadly speaking, two main approaches:

- 2D - select a single slice of tissue with excitation and then use a combination of gradients (in the two within-slice directions) to get all the signals needed to build just the 2D image of that slice. Then repeat the same process for all the other slices.
- 3D - select a full volume with excitation and just use combinations of gradients (in all three directions) to get all the signals needed.

More details on how we use combinations of gradients to get all the information needed to generate a 3D volumetric image can be found in box 4.1.

Box 4.1: Frequency and phase encoding

So far we have considered what happens if you apply a gradient during readout - whilst listening to the signal generated by the net magnetization in a coil. This is called *frequency encoding*, because different locations in physical space are associated with different resonant frequencies as determined by the strength of the gradient field in any given location. The measured signal is a composite of different components with different frequencies, the size of the component telling us how strong the magnetization was at the locations associated with that frequency.

Since the gradient can only vary in strength in a single direction, e.g., left-right, multiple locations that all experience the same gradient field strength have the same resonant frequency and thus we cannot tell them apart. What we need are further measurements where we also impose some variation in the other dimensions of the object: *phase encoding*. We know how to acquire one frequency encoded measurement, what happens if we apply a *different* gradient *before* we acquire the signal? Whenever we apply a gradient we change the resonant frequency. Our aim with phase encoding is to change the initial angle (phase) of the net magnetization in different regions of the imaging volume prior to the frequency encoding. To do this we use the fact that applying a gradient will cause the rotation (precession) to be slightly faster or slower than other parts, leading to differences in the angles of the magnetizations. This means that in some regions the rotation of the net magnetization gets ahead of that in other parts and this manifests in the signal we measure from the receive coil, where the magnetizations in different locations are now out of alignment with each other in time - strictly we have introduced phase differences between them. This gives us extra information to help tease apart different locations. For example, we might firstly apply a gradient in the anterior-posterior direction to make the net magnetization at the front of the head get ahead of that at the back, then we turn this gradient off and apply a new gradient in the left-right direction whilst at the same time listening to the signal (i.e., making several measurements of the signal).

In practice we need to use a combination of phase and frequency encoding across a number of measurements to build up all the information we require to finally arrive at a full 3D image of the head. Thus we would take one measurement in which we apply an AP (anterior-posterior) gradient (the phase encode gradient) first, then subsequently apply a LR (left-right) gradient whilst listening to the signal. Then we take another measurement following the same scheme but with less of the AP gradient applied before the LR gradient and signal recording. By repeating this approach for multiple measurements with different amounts of phase encode gradient we can get enough information to map out the magnitude of the net magnetization in 2D. If we have used slice selection then using this scheme we can get a 2D image through the head and subsequently we can repeat the whole process after selecting other slices until we have built a full 3D volume. There is also nothing to prevent us using phase encode gradients in a different orientation, thus it is possible to acquire a full 3D volume using measurements that combine the application of different durations of AP and IS (inferior-superior) gradients.

Ultimately the choice as to how many measurements are needed with different applications of the phase encode gradients is very important, as this relates to the final resolution of the image. Inevitably you have to have more measurements for a higher resolution final image, although there are various tricks that can be used to reduce the total number of measurements required, something that we consider further in Section 6.

To go any further in understanding how image acquisition works in MRI requires a discussion of 'k-space' and Fourier Transforms, especially to understand the relationship between the choices made during frequency and phase encoding, and the final image quality. If you are interested we recommend you consult one of the texts listed in the further reading section.

4.3 Inhomogeneity

Inhomogeneities in the magnetic fields, or imperfections in the gradient fields, lead to distortions in the final image. The distortions are usually mis-locations of the signal due to the B_0 or gradient fields being imperfect, as the strength of these combined fields is used to determine the location of the signal. However, large inhomogeneities in the B_0 field can also cause loss of signal. On the other hand, inhomogeneities in either the transmitted or received RF (B_1) fields result in changes in the image intensity and manifest as brighter or darker areas in the images.

The main B_0 field that is created by the superconducting coil should ideally be perfectly uniform and homogeneous, but in practice it has very small inhomogeneities, some of which are created by the presence of the object being scanned, due to the geometry and magnetic properties (susceptibility) of the materials. In particular, for the head it is the air-filled sinuses that always cause minor inhomogeneities, whilst metallic materials, such as major dental work, can cause greater inhomogeneities.

Inhomogeneity in the B_0 field means that the hydrogen nuclei in a given location within the brain will have a different resonant frequency to what was expected. This has implications for the readout process - where we use the frequency of the signal we receive to localize the signal to a specific place in the brain when building the final image. Departures in the field, and hence frequency, from the expected values will lead to localisation errors (geometric distortion) because we will associate the signal with the location we think its frequency belongs to. This is often noticeable around the sinuses. B_0 inhomogeneities can also lead, in the more extreme cases, to signal loss associated with T_2^* relaxation (see Section 5.1).

In order to reduce the minor amount of B_0 inhomogeneity in typical scans (normally around one part per million), there are specific coils built into the scanner, called shim coils, that are there to create fields that can cancel the bulk of the inhomogeneities. These coils only produce magnetic fields that are quite smooth in space (i.e., they do not change sharply) and so cannot cancel all the effects created by localised structures such as the sinuses. Nonetheless, the shim fields are very important for reducing the inhomogeneities and getting better, less distorted, images.

It is not only the B_0 field that is important in determining the resonant frequency of the hydrogen nuclei, as noted in section 4.2, we also need to add gradient fields so that different locations are deliberately at different frequencies during the readout phase of the acquisition. It is important for the gradient fields to be as close to linear as possible. Any departures from linearity (known as gradient nonlinearities) also result in signals being mislocated in the image (i.e., geometric distortion).

4.4 Chemical Shift

In the previous sections we have concentrated on hydrogen nuclei in water that have a well defined resonant frequency in a given magnetic field. However, there are other hydrogen nuclei in the head that, because they are in a different chemical environment (e.g., part of a different molecule), will have a subtly different resonant frequency. For many of these hydrogen nuclei this does not matter as they are not sufficiently abundant to be detected. However, there are some that can cause a

problem, with lipid molecules being the most common example. Because hydrogen nuclei in lipids (such as fatty tissue) have a different resonant frequency the signal they produce is at a different frequency to the hydrogen nuclei in water, thus in principle we could tell them apart. However, we use gradients to manipulate the resonant frequency of the hydrogen nuclei in water to provide location information. Without additional measures we therefore cannot separate the two sources of signal, and so the signals from lipid hydrogen nuclei gets put in the wrong location in the final image. To avoid this problem, an acquisition might include fat suppression, which specifically targets and removes the signal from fat prior to, or during, the readout phase (see Section 5.2).

5 Contrast

What we have established so far is how, using MRI, it is possible to obtain signals from hydrogen nuclei in the body and do this in a way that allows us to separate contributions from different locations from each other and thus create images. However, these signals are a measure of the net magnetization of hydrogen nuclei, a quantity that does not have a direct physiological meaning. The magnitude of the net magnetization and thus the signal we measure does depend upon how many hydrogen nuclei there are, thus at the very least we can create an image where intensity is related to the density of hydrogen nuclei (protons), which is known as a *proton-density (PD) weighted image*: an image that tells us about the relative density of water everywhere. Because water density varies between different brain tissues, different regions of the image will be lighter or darker, and these differences in density provide *contrast* in the image between tissues.

Since most methods for producing contrast only aim to exaggerate signal differences based on some effect, rather than directly measuring that effect in isolation, you will find many MRI methods described as producing *weighted* images. So PD-weighted images do not tell you what the proton (water) density is quantitatively, but regions of higher proton density have a different intensity to regions of lower proton density, although other factors may also have a minor influence on the image intensity. Commonly, MRI methods (and medical imaging in general) rely on this idea of contrast between tissues or regions with different inherent properties, hence you will often hear people refer to methods as creating or generating a particular contrast - what this means is that the method has been developed to generate signals that vary specifically so that contrast is seen in the final image between different tissues. For MRI the process of contrast 'generation' often involves some extra choices, prior to excitation.

5.1 Relaxation

The most common way to produce brain images is to exploit the relaxation properties of hydrogen nuclei to generate contrast between different tissues. Relaxation describes the process(es) of the net magnetization returning to alignment with B_0 after excitation has ended. Returning to our child on a swing analogy, once you stop pushing the swing it still keeps swinging but over time the height reached on each swing gets smaller and eventually the swing comes to rest. For the net magnetization in MRI there are two independent relaxation processes that happen simultaneously called T_1 and T_2 relaxation for simplicity (or perhaps lack of originality).

T₂ relaxation

It is easiest to consider T₂ relaxation first as this is closest conceptually to the swing analogy. When we use a coil to detect the net magnetization it is strictly the component perpendicular to B₀ that we are measuring. This signal dies away relatively rapidly, typically over the course of tens of milliseconds in brain tissues, described by the T₂ time constant. Thus a tissue with a long T₂ time constant will continue to produce a signal for longer than one with a short T₂. One way to generate a contrast between different tissues is to choose how long we wait after excitation before we measure the signal: the *echo time (TE)*. A very short echo time, relative to the T₂ time constants of the tissues being imaged, will produce a large signal, but all the tissues will end up with approximately the same magnetization, and hence intensity. If we wait longer we get a smaller signal, but tissues with short T₂ will be even smaller relative to than those with long T₂, and thus we have a better contrast. This is illustrated in Figure 5.1. Note that there is an important trade-off here - better contrast is

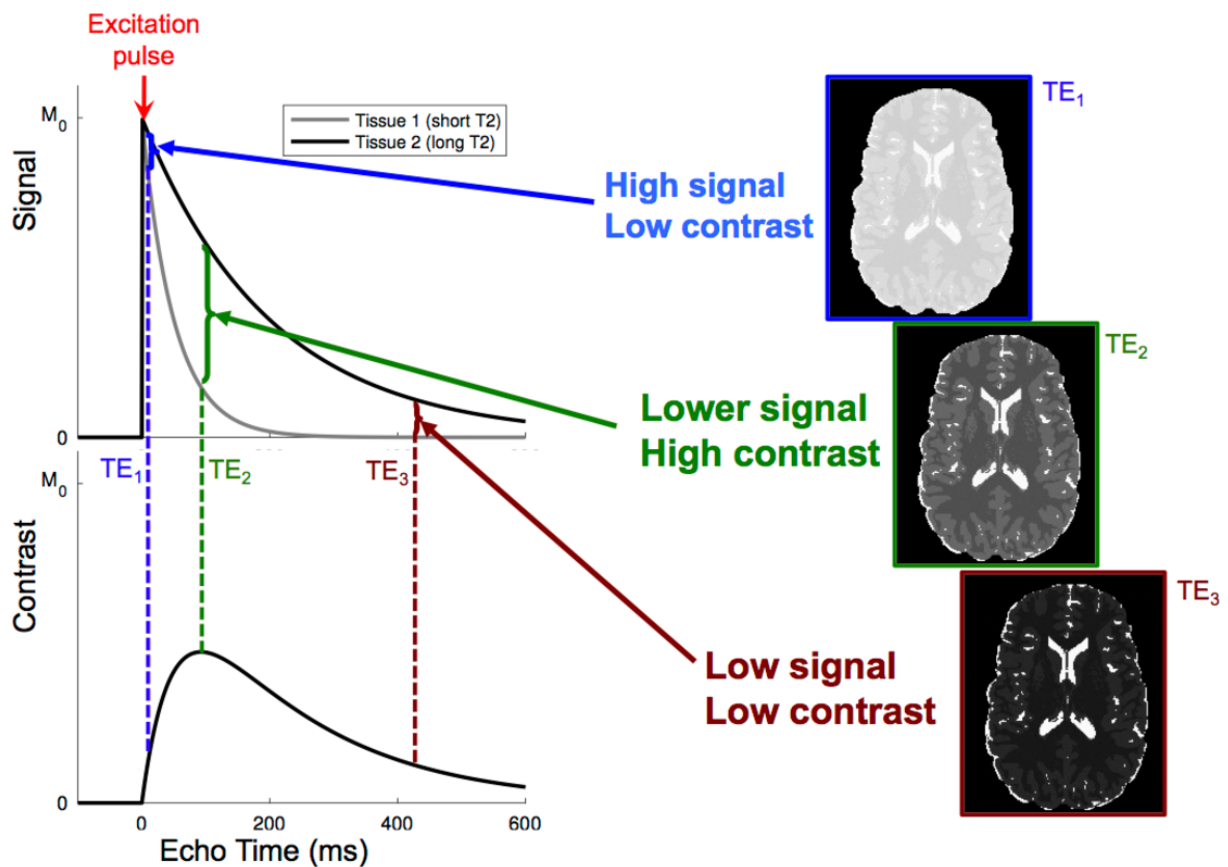


Figure 5.1: The process of generating contrast from T₂ decay. After excitation the measurable net magnetization that is perpendicular to the B₀ field, the signal, reduces over time following an exponential curve with time constant T₂. Different brain tissues have different inherent T₂ values, if we choose the time at which we measure, the echo time (TE), appropriately we can maximise the difference in the signal received from specific tissues. If we wait too long we get a very small signal and thus poor SNR. However, a short TE might provide poor contrast even if the resulting SNR is good. This example illustrates the signals from gray matter and white matter, highlighting the contrast between these two tissues in the plots on the left and in the images on the right (note that we are ignoring the CSF here, which has a large signal at all these echo times).

achieved by waiting for some time (although not too long), but this comes at the expense of smaller signal and thus poorer signal to noise ratio.

Box 5.1: Dephasing and T_2 decay

The T_2 effect arises from the fact that whilst the resonant frequency is a property of the nuclei and the strength of the main magnetic field, it is also influenced by the local environment. We have already seen how gradient fields can be used to subtly alter the resonant frequency to give us spatial localisation, and how B_0 inhomogeneities also affect this. In the body there are many hydrogen nuclei, all in slightly different chemical environments, for example based on proximity to other molecules. These influence the local magnetic field experienced by each nucleus, and thus all the hydrogen nuclei have subtly different resonant frequencies which can vary over time as the molecules move around. These tiny differences mean that, on average, some are rotating slightly faster than others and thus, over time, they get out of synchronization with each other and the net magnetization perpendicular to B_0 starts to disappear as the faster ones start to cancel out the contribution from the slower ones, this is illustrated in 2da800. This is often called dephasing, as the quantity that relates the relative rotation of each of the contributions to the net magnetization from hydrogen nuclei is the phase: this would be represented by the angles of each of the arrows in Figure 5.2.

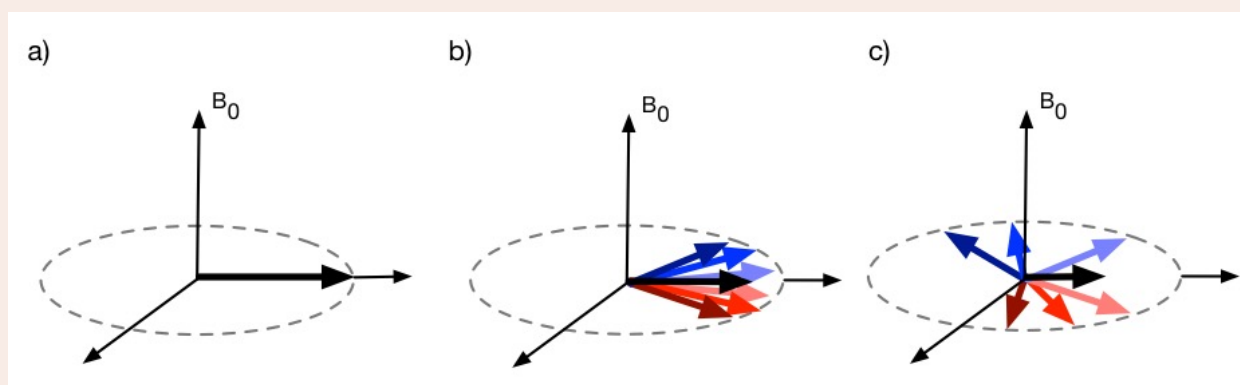


Figure 5.2: The process of T_2 decay: a) After excitation the net magnetization has been tipped out of alignment with the B_0 field, b) the net magnetization is comprised of contributions from all the different hydrogen nuclei, subtle differences in their resonant frequencies due to local environment mean that some precess slightly faster than others and c) over time they get further out of synchronisation with each other and the net value (summation of them all - shown in black) tends to zero.

T_2^* and T_2'

There are in fact two forms of T_2 relaxation: what we might call the 'true' T_2 and the time constant we observe in practice (at least when using gradient echo), T_2^* . The difference between T_2 and T_2^* is the T_2' constant (such that $1/T_2^* = 1/T_2 + 1/T_2'$).² It is useful to make the distinction between the two processes because the T_2' relaxation effect can be compensated for (refocused or simply 'undone') using a *spin echo (SE)*, allowing us to obtain a 'pure' T_2 weighted image, as explained in box 5.2.

² It would be neater if we were talking about rates rather than time constants: i.e. $R_2 = 1/T_2$ etc, since then $R_2^* = R_2 + R_2'$.

Box 5.2: Spin echo refocusing

The essential difference between T_2 and T_2' is that T_2' is a static effect, at least over the time scale of our measurements. This means that the desynchronization we considered in Box 5.1 develops at the same rate all the way through the measurement. By contrast the T_2 component is random, e.g., due to changes in resonant frequency brought about by molecules moving around in the environment.

We can recover the fixed, and therefore predictable, T_2' component using a spin echo. For a spin echo acquisition, halfway through the TE we apply another RF pulse (a “refocussing” pulse) to flip the net magnetization that is perpendicular to B_0 . This means that the ‘faster’ hydrogen nuclei that were previously ahead of the ‘slower’ ones are now behind and at the echo time (TE) they have caught back up again, thus removing the contribution from T_2' . This is illustrated in Figure 5.3. Ultimately we can choose to generate either T_2 or T_2^* weighted images depending upon what will generate the best contrast for an application, by choosing either a GRE or SE type of acquisition.

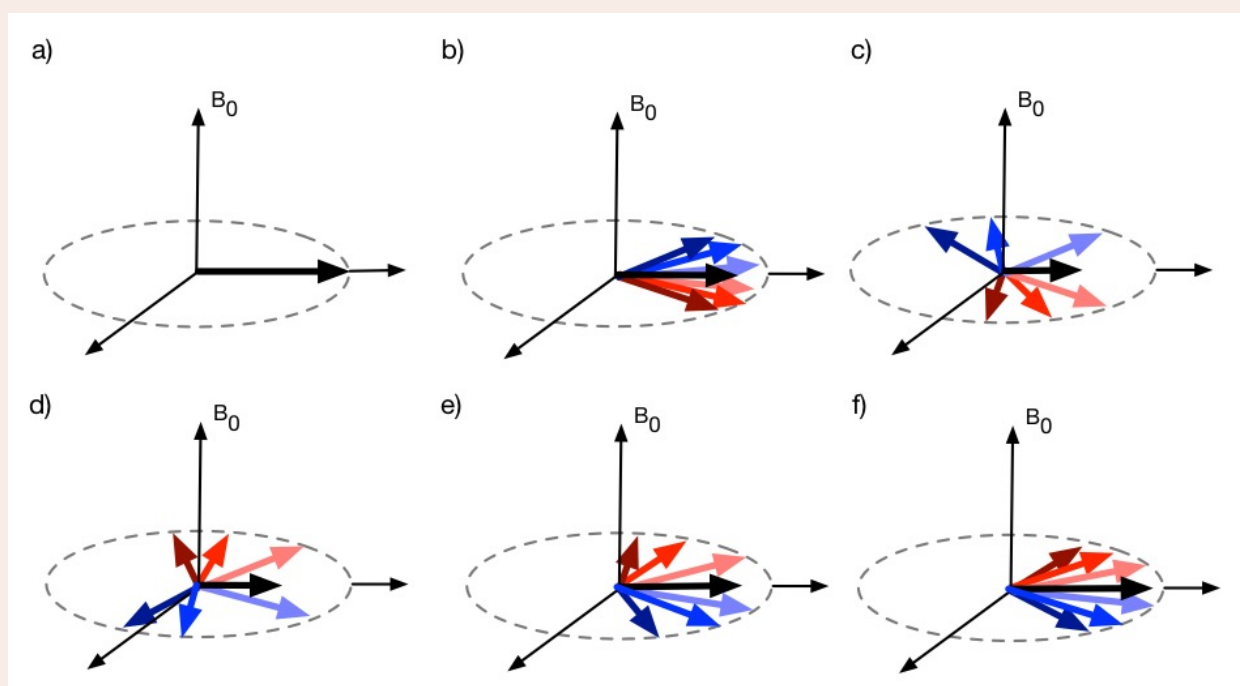


Figure 5.3: The process of generating a spin echo. a) excitation followed by, b) and c), the different hydrogen nuclei getting out of synchronisation due to magnetic field variations, d) halfway through the TE the B_1 field is applied again to reverse the magnetization of the hydrogen nuclei, e) the ones that were ahead (in blue), because they had a faster resonant frequency, are now behind and catch up the slower ones (in red), f) up to the point where the net magnetization is partially recovered.

 T_1 relaxation

Whereas T_2 relaxation was about the component of the net magnetization that is perpendicular to the B_0 field, T_1 relaxation is the recovery of the component that is aligned with the field (parallel to

it). For the brain, this happens over a few seconds, i.e., a lot slower than T_2 decay, which happens in tens-hundreds of milliseconds. This process means that after an excitation, if we wait long enough, the net magnetization will ‘recover’ to its initial magnitude and direction (parallel with the B_0 field) and then we can excite again and start another, equivalent measurement process, and repeat this over and over. The time taken between excitations of the same region (e.g., the same slice if we are using 2D slice-wise acquisitions) is called the *repetition time* (TR) and is a very important parameter for all MRI sequences, largely governing timing and contrast. If we choose a repetition time that is sufficiently short, so that the net magnetization aligned with the field does not fully recover, then the next excitation will start with a smaller magnetization and so the signal we measure (at our chosen echo time) will be reduced. Since different brain tissues have different T_1 values we can use this partial recovery of the magnetization to give different signal magnitudes for different tissues, as illustrated in Figure 5.5. Specifying the contrast in a T_1 -weighted image is thus largely a matter of choosing the TR. Normally it is a good idea to wait a couple of TRs before starting to generate any images, in order for the magnetization to reach a steady state, hence you often find that the MRI scanner will normally take a couple of ‘dummy’ scans before collecting your data.

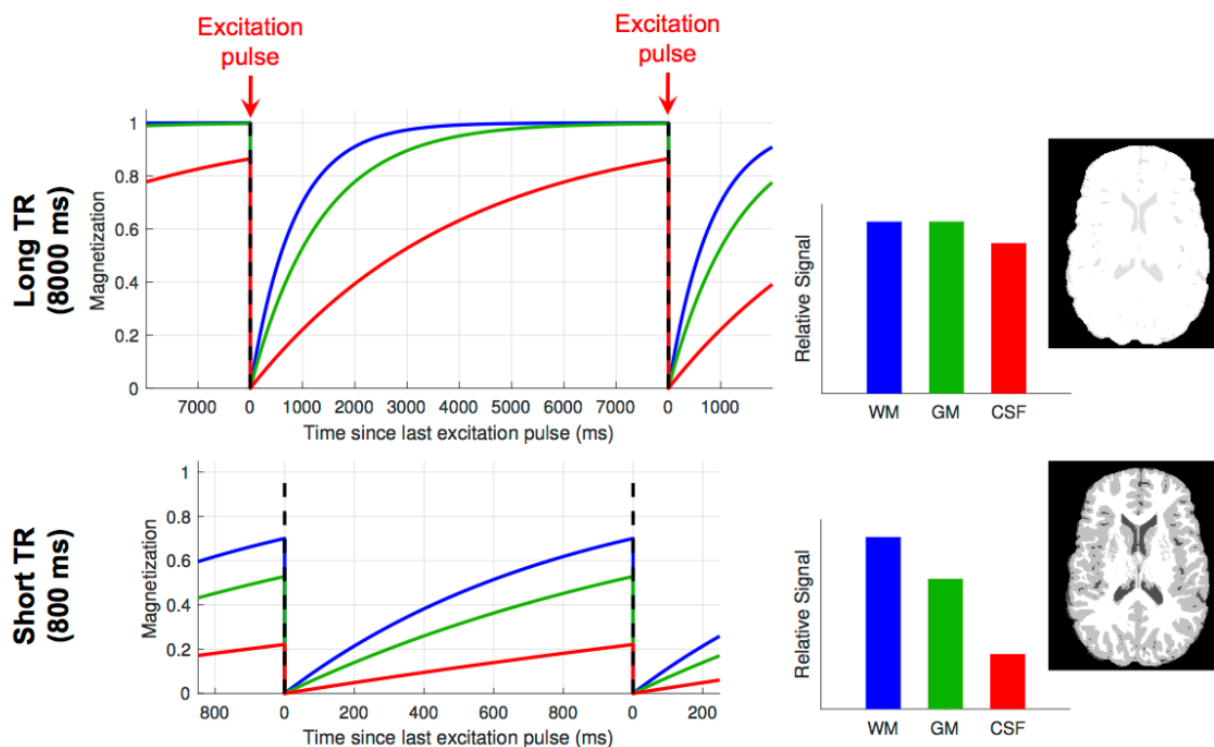


Figure 5.5: Creating images with T_1 contrast: The net magnetization parallel to the B_0 field for tissues with different T_1 will recover to a different degree during the TR. Thus the choice of TR will affect how big a signal is received from the different tissues, introducing a T_1 -weighting. For a long TR (top row), all the tissues have mostly recovered back to equilibrium and there is very little T_1 -weighting. For a short TR (bottom row), the partial recovery that is achieved is quite different for each tissue, meaning that the signal has a strong T_1 -weighting.

Box 5.3: Combined T_1 and T_2 recovery

It is tempting to think of the recovery of the net magnetization as being the exact reverse of the excitation process where we tipped the magnetization toward the transverse plane, perpendicular to B_0 . However, it is not correct to think of the net magnetization simply tipping back to align with B_0 again, as the processes that govern the size of magnetization in the transverse plane, T_2 , and the longitudinal direction, T_1 are independent and happen over different time scales. It is better to regard the magnetization after excitation as having two components: the transverse component and the longitudinal component. After a 90° excitation we have zero longitudinal component and a maximal transverse component. The transverse component now reduces according to T_2 and the longitudinal component increases according to T_1 . Since T_1 is longer than T_2 the growth of the longitudinal component is slower than the decay of the transverse component, leading to an overall recovery trajectory of the net magnetization that looks something like that shown in Figure 5.4.

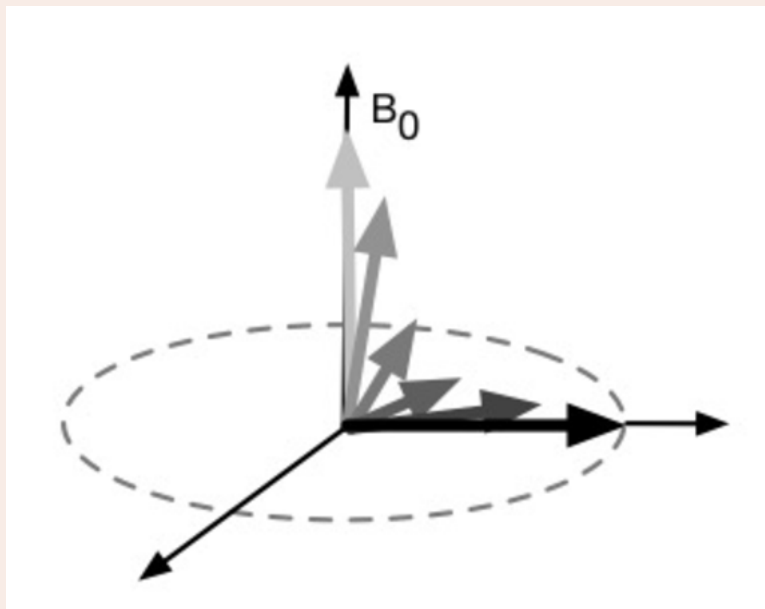


Figure 5.4: The recovery of the net magnetization does not look like a simple reversal of the excitation that tipped it into the transverse plane. The transverse component of the net magnetization decays more quickly (with T_2) than the ‘regrowth’ of the longitudinal component (with T_1).

 T_1 , T_2 and PD weighted images

In practice T_1 and T_2 weighting both occur, but we can choose to favour one or the other because one involves the choice of the TR and the other the choice of the TE, which are independent parameters. Figure 5.6 shows how the choice of TR and TE can give rise to different types of image, here ‘short’ and ‘long’ are judged relative to the T_1 and T_2 values of the tissues. Note that if we choose short TE and long TR we remove almost all of the T_1 and T_2 weightings and get back to the proton density weighted image where we are just measuring the maximum net magnetization of the tissue, which is proportional to the density of water.

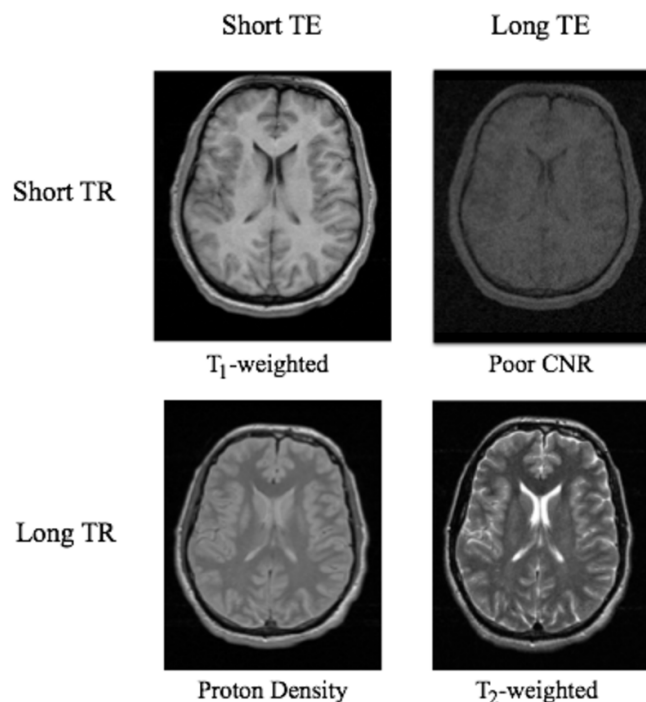


Figure 5.6: The choice of TE and TR determines the type of weighting (contrast) seen in the final image. Short and long in this context is relative to the T_1 and T_2 values of the tissues involved. Note that the combination of short TR and long TE is not very useful, since this involves very little measurable signal and thus achieves poor contrast compared to the noise (or poor contrast-to-noise ratio, CNR).

5.2 Inversion

Thus far we have considered the case of excitation perturbing the magnetization out of (parallel) alignment with the B_0 field, after which the perpendicular part decays according to the T_2 time constant, and the parallel part recovers according to the T_1 time constant. It is also possible to do more than simply push the magnetization out of alignment with the B_0 field, but to flip it so that it is aligned opposite to the B_0 field: *inversion*. If we do this and then wait, the magnetization will then recover, according to T_1 , back toward being aligned with B_0 again. Note that we will not detect any signal nor see any T_2 decay during this process, as none of the magnetization is perpendicular to B_0 . Figure 5.7 illustrates this process of inversion-recovery for a number of tissues. At some point during this recovery process we could perform a normal excitation and readout and thus get a measure of magnetization of the tissues that would, like the T_1 -weighted images from before, distinguish between them on the basis of T_1 but with even greater contrast. Inversion-recovery is thus a useful building block for various types of MRI sequences and is also used extensively in tissue suppression.

Suppression and Nulling

The combination of T_1 and T_2 relaxation also provides a means to suppress or null specific tissues. In the previous section we met the concept of inversion-recovery, and if you examine Figure 5.7 again you will see that at a certain point on each tissue's recovery curve the magnetization goes through zero. If we performed further excitation at this moment no signal would be generated by that tissue in the final image - it would have been suppressed. Thus inversion-recovery can be used to suppress a tissue with a known T_1 value, and this is used in the FLAIR (FLuid-Attenuated Inversion Recovery) sequence to suppress CSF contributions (Figure 5.7, T_{14}), which is widely used clinically.

Things can get a little more complicated if there are multiple tissues to be suppressed with different T_1 values or, as is commonly the case, we want to suppress tissues with a range of T_1 values, for example, because we do not precisely know the relevant T_1 *a priori*. This can still be achieved using inversion-recovery, but we have to repeat the inversion process multiple times allowing for different recovery times, such that each specific value of T_1 is effectively suppressed; for example, Double Inversion Recovery (DIR) uses two inversions prior to readout to suppress both CSF and white matter.

A tissue that we often want to suppress in brain imaging is fat. Although there is very little lipid present inside the skull, the presence of a layer of fat around the outside of the skull can introduce artifacts inside the brain during imaging due to chemical shift, as discussed in section 4.4. Inversion-recovery can be one means to suppress the lipid signal. Another, that exploits the resonant frequency difference of hydrogen nuclei in fat, is to apply a separate fat saturation RF pulse prior to the main excitation. This is, in effect, an extra excitation pulse but tuned specifically to only influence the hydrogen nuclei in fatty tissue. Once excited, their signal contribution will decay with their T_2 , reducing their influence on any subsequent imaging. In practice normal T_2 decay is not fast enough to make this sufficiently effective, but through the use of “crushing” gradients the T_2 decay (dephasing) process can be accelerated and good fat suppression achieved. There are various other more sophisticated fat saturation schemes, often specifically tailored to suit particular applications.

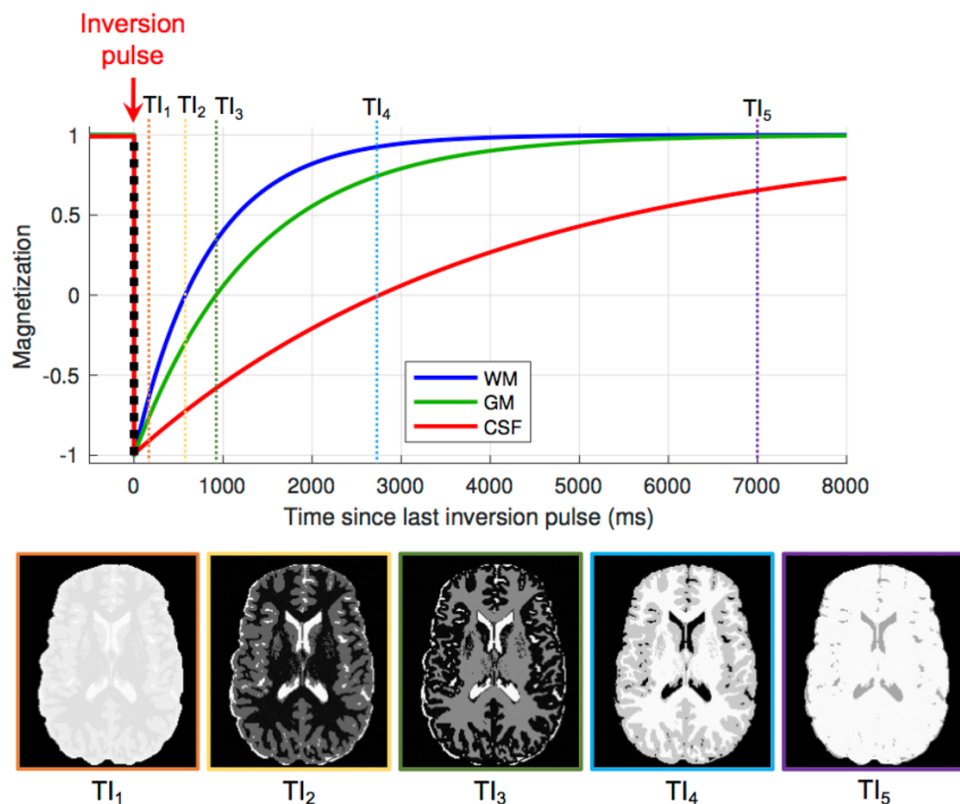


Figure 5.7: Inversion recovery contrast. The inversion pulse flips the magnetization so that it points in the opposite direction to B_0 . After this, the magnetization recovers back towards equilibrium according to the T_1 time constant of each tissue. Introducing a delay between the inversion pulse and the imaging time (TI) introduces tissue contrast based on the T_1 differences. In addition, if the TI is chosen such that a particular tissue is at its “null” point (i.e., has zero magnetization), then the signal from that tissue is totally removed, as is demonstrated in TI_2 , TI_3 and TI_4 for WM, GM and CSF respectively.

Spin labelling

The process of inversion can also be used to create a 'label'. Arterial Spin Labeling (ASL) uses inversion of the hydrogen nuclei that are in water within blood, whilst the blood is in the arteries in the neck. This 'labeled' blood then travels into the brain where it replaces non-labeled blood. Two images, one with and one without labeling, can be subtracted to reveal the quantity of labeled blood-water that has reached the brain. Since water is freely exchanged between the blood and tissue, most of the labeled blood-water ends up in the tissue where it stays, and this allows the delivery of blood-water and hence perfusion (often referred to as Cerebral Blood Flow) to be measured. One challenge for this technique is that after inversion, the hydrogen nuclei in the blood-water undergo T_1 recovery back toward their unlabeled state. Furthermore, the rate of blood-water delivery to the brain means that only a relatively small amount of water is replaced within the time it takes for T_1 recovery to have destroyed the label. Inherently, therefore, ASL has relatively poor signal-to-noise ratio. Details of effective ASL acquisition and analysis, including perfusion quantification, can be found in the primer *Introduction to Perfusion Quantification using Arterial Spin Labeling*.

5.3 Diffusion

As the name implies, diffusion imaging exploits diffusion, namely that of water, to generate contrast. To make the signal sensitive to diffusion we use a preparation process that has the effect of reducing the measured signal based on how far, on average, water molecules have moved during the preparation part of the imaging, before we do the main excitation. We start by flipping the magnetization to be perpendicular to B_0 . Once flipped, the processes behind T_2 relaxation kick in. As we saw in section 5.1, if we use a spin echo we can, in principle, recover the T_2' component of the relaxation and only be sensitive to T_2 relaxation. For diffusion imaging we use both gradient fields and the spin echo process, before the readout, to create a signal that depends on the diffusion of the water molecules. During the first part of the spin echo we apply a gradient field that causes hydrogen nuclei in different locations to have a higher or lower frequency. When water molecules move around in this field they experience different fields as they move, and their magnetization more easily gets out of synchronisation with that of other molecules, as each one is moving differently. This is effectively like an additional T_2 decay, and so the spin echo is able to recover less signal from hydrogen nuclei that have been moving around than from ones that stay in, or near, the same location. In fact, the greater the movement, the more decay there is and so the measurable signal is decreased, for more details see Box 5.4. Thus diffusion MRI provides a measure of the average distance travelled by the water and hence the diffusion. For more information on diffusion imaging, and the information that can be extracted from it, see the primer *Introduction to Neuroimaging Analysis*.

Box 5.4: Gradients, resonance and diffusion

As explained in the main text, diffusion measurements depend on making the signal sensitive to moving water molecules using gradient fields. One way to make an MRI measurement that is sensitive to diffusion would be to insert a preparation period before the excitation, in which we apply a gradient for the first half of the preparation time, then subsequently apply the exact opposite gradient for the rest. This modification has no impact on the signal unless the hydrogen nuclei move.

Applying a gradient will cause hydrogen nuclei to have different resonant frequencies from each other, and this enhances the effect of desynchronization that we saw in Figure 5.2. If the hydrogen nuclei have not moved, during the second half of the preparation, those that had the lower resonant frequency before now have a higher value and vice versa. Thus those that were 'behind' catch up with those in 'front' reversing the effect of desynchronization introduced by the gradient. Thus without any motion of the hydrogen nuclei we only get the T_2 component of decay that we had previously. However, any hydrogen nuclei that have moved will have spent time at different locations in the gradient field during both halves of the preparation and thus will have had different resonant frequencies in the first half to the second half. The consequence is that they will not end up back in synchronization with the other hydrogen nuclei and the final signal will be smaller: having both a T_2 decay and an extra movement (i.e., diffusion) component.

In practice, we can achieve the same effect by modifying the spin echo, Figure 5.3. If we apply the gradient during the first half then we enhance the dephasing, but the RF pulse we apply halfway through the spin echo flips the magnetization, so that if we apply exactly the same gradient during the second half we recover all of the phase dispersion that was due to the first gradient for any hydrogen nuclei that have not moved. Where there has been motion the final spin echo signal is still decreased due to dephasing, which is what we want.

We can vary how strong the influence of the diffusion weighting is by adjusting the timing of the spin echo and/or increasing the strength of the gradient field that is applied. You will often see people quote the 'b-value' for a diffusion acquisition, this combines these two factors into a single value. As we saw in section 4.2, gradients are applied in a given direction in 3D in the scanner. For diffusion this means that the effect depends upon the direction we choose and it is only sensitive to movement along that direction. Thus we use a number of different directions to summarise the movement in simple diffusion scans, as at least 3 directions are needed if we want to account for the main directions of water movement. We can also use a combination of many different directions, and even different b-values, to get more detailed information about the movement of water within the voxels in the image and this is the basis of diffusion tensor imaging and tractography.

5.4 BOLD

The Blood Oxygen Level Dependent (BOLD) effect is the basis for most functional brain imaging experiments, including both task and resting-state fMRI. The BOLD signal is indirectly sensitive to

neuronal activity through changes in metabolic demand, which are then reflected in changes in localised blood volume, blood flow and oxygen extraction, which finally affect the local B_0 field due to the magnetic properties of deoxygenated blood. BOLD is essentially a T_2^* dependent contrast which relies on the different magnetic properties of oxygenated and deoxygenated hemoglobin within the blood. By rapidly acquiring images that have T_2^* weighting, it is possible to detect changes in T_2^* associated with changes in blood oxygenation and thus map neuronal activity - for more details see the primer *Introduction to Neuroimaging Analysis*.

As we saw in section 5.1, localized magnetic field inhomogeneities will affect how the MR signal decays, as the hydrogen nuclei end up with a range of resonant frequencies. This means that they get out of synchronisation with each other more quickly, which reduces the net magnetization and hence reduces the signal. This effect is quantified by the T_2^* relaxation time, which incorporates all the T_2 relaxation process based on many microscopic properties as well as the extra effect that arises from localized B_0 field inhomogeneities. An increased concentration of deoxygenated hemoglobin nearby will induce stronger variations in the B_0 field and result in a quicker relaxation process, with the signal decaying away faster. That is, the signal decreases when there is more deoxygenated hemoglobin in the vicinity. Counterintuitively, this is the case when there is *less* neuronal activity. When neuronal activity increases the result is an increase in the locally measured MRI signal since the body more than compensates for the extra metabolic demand by providing plenty of oxygenated blood (higher blood volume and flow) leading to an overall reduction in deoxyhemoglobin - this is illustrated in Figure 5.8.

BOLD requires that the sequence that is used is sensitive to T_2^* relaxation. This means that *spin-echo* sequences are generally not used (unlike diffusion MRI) as the spin-echo removes most of the T_2' component of the change, leaving the image only sensitive to the T_2 relaxation, which is largely unaffected by the deoxyhemoglobin concentration. Hence most functional imaging is based on fast GRE acquisitions that are tuned to be as sensitive to T_2^* relaxation effects as possible.

5.5 Contrast agents

Contrast agents are not as widely used in neuroimaging research as they are in clinical practice, in no small part because they are less flexible and cannot readily be used repeatedly in the same individual compared to alternative endogenous mechanisms that we can use with MRI. Generally MRI contrast agents rely on materials that change either the T_1 or T_2^* properties, or both, of the tissues in which they are present. The most common contrast agent is based on gadolinium, introduced in the form of a chelate, in which the body is protected from the gadolinium itself (which is highly poisonous), but water molecules can still interact with it. A gadolinium agent shortens both T_1 and T_2^* ; the T_1 effect is mainly used to visualise arterial structures or to observe when the agent leaves the bloodstream and enters the tissue (into the extracellular spaces). It is not that common for the agent to leave the bloodstream when in the brain, unless the blood-brain-barrier is broken; for example, in a tumour. What is more often exploited are changes in T_2^* that are seen when the agent is passing through the tissue in the capillary network: the presence of the gadolinium creates local gradients in the magnetic field (similar to deoxyhemoglobin and BOLD) which accelerates the T_2^* decay process. This causes a signal decrease in T_2^* weighted images as the agent passes and can be used to image perfusion and related measures, such as blood volume. Rapid T_2^* imaging with a gadolinium agent is the basis of dynamic susceptibility contrast (DSC) perfusion weighted imaging. For more information see the primer *Introduction to Neuroimaging Analysis*.

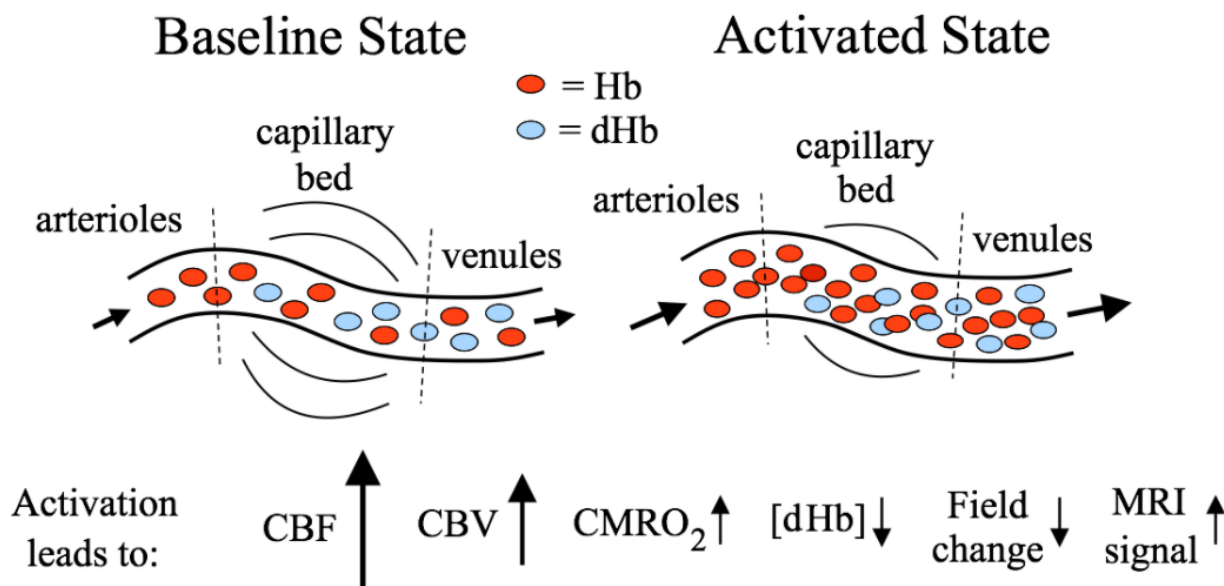


Figure 5.8: Illustration of the BOLD effect: an increase in local neuronal activation (moving from the baseline state, left, to the activated state, right) is accompanied by large increases in cerebral blood flow (CBF), cerebral blood volume (CBV), more modest increases in oxygen extraction (CMRO₂), and thus an overall increase in the amount of oxygenated blood (Hb) present compared to deoxygenated blood (dHb). As a consequence the concentration of deoxyhemoglobin decreases, as do the magnetic field inhomogeneities that it causes, leading to an increase in the T₂*-weighted MRI signal.

6 Fast Imaging

For many applications, such as diffusion and functional imaging, we need to acquire multiple images quickly - in a matter of a few seconds per image rather than the several minutes often required for conventional MRI scans. To achieve this requires specialised *fast imaging* techniques, which is what this section will discuss.

The volume acquisition time (i.e., time to acquire one image of the whole brain) of a sequence is dependent on a number of factors, with the main ones being:

- Time taken for each readout and any time spent waiting for the desired amount of relaxation to occur; for example, using a relatively long TR to achieve a specified T₁ weighting.
- Time to prepare magnetization.
- Number of separate readouts needed to collect the information for the whole image (related to the resolution).
- Sampling scheme and number of samples per readout.

One way of reducing the acquisition time is to lower the spatial resolution, and thus reduce the number of voxels in the image. This can have a substantial impact on the acquisition time, although it is not typically enough to reduce the acquisition time of a scan from a few minutes to a few seconds. For example, changing the resolution of a structural scan from 1x1x1mm to 3x3x3mm reduces the number of readouts required by roughly a factor of 10, but that will still only turn a 5 minute scan into a 30 second scan.

The greatest speed ups are usually achieved by manipulating the sampling scheme. Two common techniques for doing this are:

- Acquiring a full slice's worth of data after each excitation pulse, as is done with EPI (see section 6.1).
- Acquiring samples in parallel: reducing the number of samples required by using the additional information provided by multiple receive coils. This can be done with in-plane acceleration techniques such as SENSE or GRAPPA (see section 6.2), or with simultaneous multi-slice or multiband methods (see section 6.3).

6.1 Echo Planar Imaging (EPI)

The main principle on which EPI is based is capturing all the data associated with a slice using a single readout - that is, making all the necessary measurements to form an image of a slice following a single RF excitation pulse. This is different from most other sequences where a single readout only captures enough data for one "line" of the image (though note that there is no simple correspondence between the data acquired and the lines in the image, as the "lines" refer to an alternative version of the image - the k-space version - the details of which are not important except that many "lines" still need to be acquired to make a 2D image of a slice).

In order to get the necessary frequency and phase encodings, at least two sets of gradients need to be changed during the measurement (readout) so that enough information about the 2D distribution of the magnetization is present in the signal. This differs from slower sequences where only one gradient is typically present during the readout. In the case of EPI it is necessary to measure many "lines" worth of data after each RF excitation pulse, as we want to get all the data for a whole slice in one go. By applying a series of different gradients in rapid succession, spanning the two directions of the slice, we can acquire all the data we need for a single slice before the next RF excitation pulse, rather than spending time waiting for an appropriate amount of relaxation to occur (see section 5.1). This is why EPI is faster than conventional sequences.

There are disadvantages or tradeoffs involved with using EPI, which are typically related to the speed of the measurements and the length of the readout following each RF excitation pulse. A longer readout means there is more time for T_2 (or T_2^*) decay to occur during the time that the imaging data are being recorded. The consequences of this are that there is some blurring associated with the image and, more significantly, that there is a limit to how many "lines" of data can be acquired before the signal decays too much. Both of these effectively limit the spatial resolution that can be obtained from EPI, although there are other techniques (e.g., acceleration methods - see the next sections) that can help to offset this and achieve higher resolutions.

In addition to the limited resolution, a major tradeoff involved when using the EPI sequence for fast imaging is an increase in the magnitude of geometric distortions. This is due to the longer total readout time, which allows more time for B_0 inhomogeneities (see section 4.3) to affect the signal. A

key parameter that is related to the total readout time is the *echo spacing*, which is the amount of time required to acquire each “line” or, approximately, the total readout time divided by the number of voxels in the phase encode direction.

6.2 In-Plane Acceleration (SENSE, GRAPPA, etc.)

Another way of speeding up acquisitions is to use the multiple, individual coils within a head coil array (as most modern head coils consist of an array of small coils) to make simultaneous measurements and hence increase the amount of data collected in a given amount of time. This relies on the fact that each individual coil is more sensitive to signals emitted by molecules that are close to it, and hence each coil then acquires somewhat independent data that includes partial information about the spatial location of the signal. In practice the data are not entirely independent, as there is overlap between areas that nearby coils are sensitive to, and this is one limitation onto how much acceleration can be achieved.

What makes the acceleration possible is the fact that fewer “lines” need to be acquired, leading to less readouts or shorter readouts. The “missing lines” (and hence a reduced number of measured samples) are then effectively “filled in” by sophisticated reconstruction algorithms, which use the extra data from the multiple coils as a way of regaining the information about location that otherwise would have been provided by the lines that were skipped. In essence these reconstruction methods take various combinations of data from the different coils to synthesize data for the samples in the missing lines, often using some calibration data to factor in aspects of the coil geometry and couplings. Two well known examples of such reconstruction algorithms (which require slightly different data too) are GRAPPA and SENSE. Both of these can be applied to a wide range of sequences and can work with 2D or 3D acquisitions.

As with most alternatives in MRI there is a tradeoff that is made when using acceleration to speed up the image acquisition. The main penalties are a reduction in SNR and increased levels of artifacts. Both the changes in SNR and the artifacts are not uniform across the image and tend to be different in the periphery, nearer to the coils, compared to the centre of the brain, which is the furthest away from the coils. For small acceleration factors (e.g., 2 or 3 times acceleration) the changes in SNR and artifact levels are often acceptable and use of this degree of acceleration is very common. However, with higher acceleration factors the changes in SNR and artifact levels can be quite dramatic and can lead to poor or even unusable image quality. Hence it is important to carefully determine how much acceleration is acceptable, which will depend on the nature of the head coil hardware (especially the number of coils in the coil array) and the imaging sequence being used. In general, coil arrays that contain more individual coils are capable of higher accelerations before the increased artifacts and reduced SNR prevent the images being useful.

6.3 Multi-Slice Acceleration (Simultaneous Multi-Slice, Multiband, etc.)

Another way of utilizing the independent data from different coils is to take a multi-slice (2D) acquisition and modify it so that more than one slice’s data is acquired at any one time. This is done by exciting multiple slices simultaneously and then separating their contributions later on, based on

the different local sensitivities of the coils - that is, coils nearer to a given slice record a stronger signal from that slice. In the same way as in-plane acceleration methods, this requires multiple coils in the head coil array, and typically works better when there are more coils.

In these methods a group of slices are excited at once, with the number of slices in each group equal to the acceleration factor (also called a multiband factor). That is, for an acceleration factor of three, each group would contain three slices, separated over the FOV - see figure 6.1. The whole volume is then acquired a group at a time, with one group of slices for each excitation, until the FOV is fully covered. In each case the measured signal is a combination of the signal across the slices within one group, but is recorded separately for each coil and then the reconstruction method separates the signal into different slices using the fact that each coil is more sensitive to the signals from nearby molecules. This is similar to the idea used in the in-plane acceleration methods, but can utilise different tricks in the sequence that can assist in separating the signals more cleanly.

Both multi-slice and in-plane acceleration methods can be combined to achieve the advantages of both approaches. The multi-slice acceleration methods share similar limitations to the in-plane acceleration methods in that they can reduce SNR and generate artifacts, neither of which is uniform across the image. These penalties are reduced when there are more coils in the head coil array, but vary with different sequences and parameter settings. In general, the use of both of these

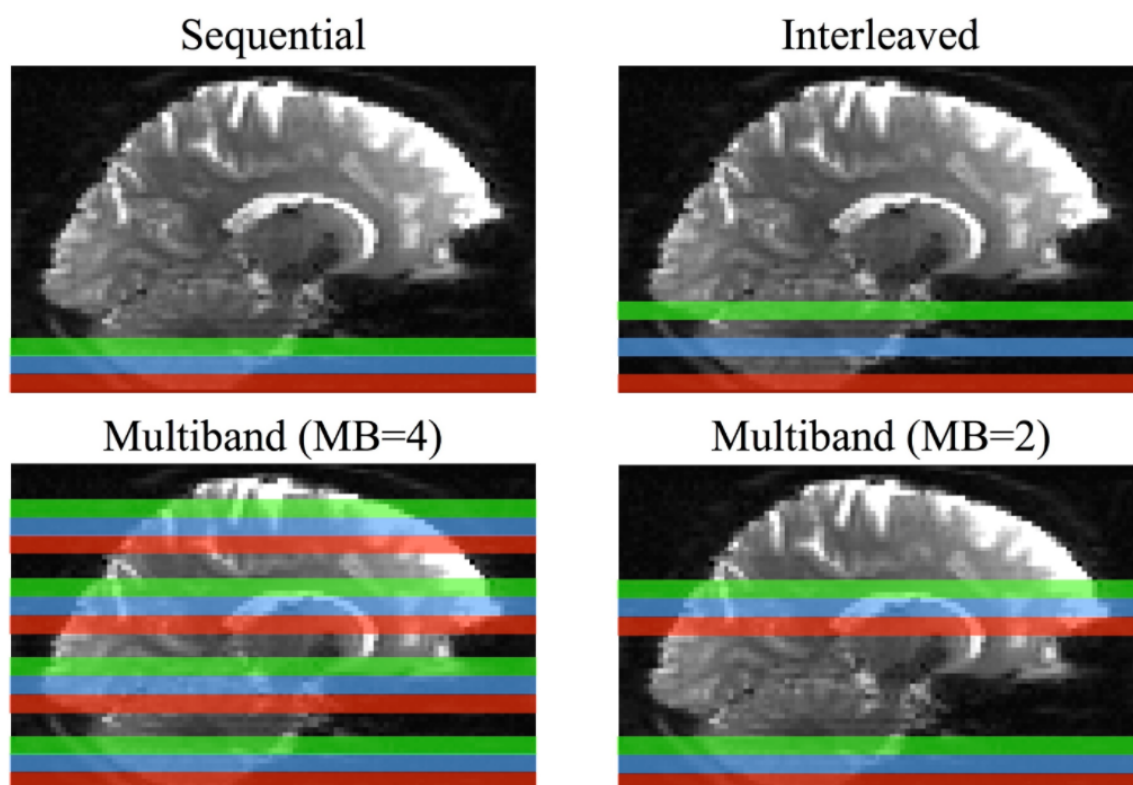


Figure 6.1: Illustration of different slice acquisition options, for both conventional multi-slice methods (top row) and simultaneous multi-slice methods (bottom row). Only the first three slices (or slice groups) acquired are shown here, in order to simplify the illustration: first slice in red, second slice in blue, third slice in green. The different options shown here are: (a) sequential (bottom-up); (b) interleaved (bottom-up); (c) multiband (MB=4); (d) multiband (MB=2). Multiband (i.e., simultaneous multi-slice acceleration) acquires more than one slice at a time, the number of slices being given by the multi-band factor, MB (e.g., the four red slices in MB=4 are acquired at the same time). Only bottom-up orders are shown here but equivalent top-down versions of all of these options also exist.

methods (either separately or together) is becoming more common, especially for reducing acquisition time and for improving spatial resolution.

7 Quantitative MRI

Much of the application of MRI in neuroimaging is about the generation of contrast, for example, structural imaging of the brain with T_1 - or T_2 -weighted imaging, or the detection of change, for example, changes in BOLD contrast under different stimulation conditions. Many of these methods are either: non-quantitative, with images being generated for visual interpretation or the extraction of numerical measures from features (e.g. size, shape) in the image; or 'semi-quantitative', with changes being measured and converted into statistical measures of whether the change observed can be considered significant.

There are also examples of more quantitative uses of MRI. As we saw in section 5.3, diffusion can use a range of images with different diffusion weighting to derive quantitative measures of water diffusion. It is also possible, again through the use of multiple different images, to estimate the numerical values of T_1 and T_2 in tissue. For example, combining images with different TE values to sample the T_2 decay, or at different times during inversion recovery to sample T_1 recovery, with mathematical models of the recovery process being fit to the data to extract quantitative T_1 and T_2 values. The advantage of quantitative measurements is that, at least for an ideal acquisition, they will be consistent from day-to-day, site-to-site, and scanner-to-scanner. However, quantitative values are often not exploited as absolute measures in their own right because it can be tricky to link the quantities measured to physiological processes happening in the brain. For example, there are a range of biological and physiological changes that affect the diffusion of water and could give rise to a change in the quantitative diffusion MRI measures or T_1 , and without extra information it can be almost impossible to link them to a specific change in physiology that might be expected in a disease or due to stimulation. Some examples of specific quantitative and physiological imaging do exist, of which perfusion imaging is one, where the images can be used to directly measure a physiological process of interest.

FURTHER READING

- Huettel, S. A., Song, A. W., & McCarthy, G. (2014). *Functional Magnetic Resonance Imaging* (3rd ed.). Sinauer Associates.
 - This is an accessible, introductory-level textbook, with more of a focus on the physics of acquisition as well as on functional physiology and task fMRI studies.
- Buxton, R. (2009). *Introduction to Functional Magnetic Resonance Imaging: Principles and Techniques* (2nd ed.). Cambridge University Press.
 - A more detailed introduction to both MR physics and neural physiology, which also offers an overview of the principles of ASL and BOLD fMRI.
- Jenkinson M., & Chappell, M. (2017). *Introduction to Neuroimaging Analysis* (Oxford Neuroimaging Primers). Oxford University Press.
 - This primer provides a general overview of the acquisition and analysis of structural, diffusion, functional and perfusion MRI.
- Chappell, M., MacIntosh, B., & Okell, T. (2017). *Introduction to Perfusion Quantification Using Arterial Spin Labelling* (Oxford Neuroimaging Primers). Oxford University Press.
 - This primer provides an introduction to the acquisition and analysis of perfusion MRI.

FURTHER READING

- Brown, R. W., Cheng, Y.-C. N., Haacke, E. M., Thompson, M. R., & Venkatesan, R. (2014). *Magnetic Resonance Imaging: Physical Properties and Sequence Design* (2nd ed.). Wiley-Blackwell.
 - This is a comprehensive and advanced textbook covering many areas of MRI Physics - it is most suitable for physicists and engineers.
- Bernstein M. A., King K. F., & Zhou X. J. (2004). *Handbook of MRI Pulse Sequences*. Academic Press.
 - This is an advanced textbook covering the principles and details of MRI pulse sequences - it is most suitable for physicists and engineers.