

University of Barcelona
Department of Organic Chemistry

UNDERSTANDING NMR SPECTROSCOPY

James Keeler

University of Cambridge, Department of Chemistry

© James Keeler, 2002 & 2004

1 What this course is about

This course is aimed at those who are already familiar with using NMR on a day-to-day basis, but who wish to deepen their understanding of how NMR experiments work and the theory behind them. It will be assumed that you are familiar with the concepts of chemical shifts and couplings, and are used to interpreting proton and ^{13}C spectra. It will also be assumed that you have at least come across simple two-dimensional spectra such as COSY and HMQC and perhaps may have used such spectra in the course of your work. Similarly, some familiarity with the nuclear Overhauser effect (NOE) will be assumed. That NMR is useful for chemists will be taken as self evident.

This course will always use the same approach. We will first start with something familiar – such as multiplets we commonly see in proton NMR spectra – and then go deeper into the explanation behind this, introducing along the way new ideas and new concepts. In this way the new things that we are learning are always rooted in the familiar, and we should always be able to see *why* we are doing something.

In NMR there is no escape from the plain fact that to understand all but the simplest experiments we need to use *quantum mechanics*. Luckily for us, the quantum mechanics we need for NMR is really rather simple, and if we are prepared to take it on trust, we will find that we can make quantum mechanical calculations simply by applying a set of rules. Also, the quantum mechanical tools we will use are quite intuitive and many of the calculations can be imagined in a very physical way. So, although we will be using quantum mechanical ideas, we will not be using any heavy-duty theory. It is not necessary to have studied quantum mechanics at anything more than the most elementary level.

Inevitably, we will have to use some mathematics in our description of NMR. However, the level of mathematics we need is quite low and should not present any problems for a science graduate. Occasionally we will use a few ideas from calculus, but even then it is not essential to understand this in great detail.

Course structure

The course is accompanied by a detailed set of handouts, which for convenience is divided up into “chapters”. You will notice an inconsistency in the style of these chapters; this comes about because they have been prepared (or at least the early versions of them) over a number of years for a variety of purposes. The notes are sufficiently complete that you should not need to take many extra notes during the lectures.

Each chapter has associated with it some exercises which are intended to

illustrate the course material; unless you do the exercises you will not understand the material. These exercises will give you a feel for what you can do with NMR data and how what you see relates to the theory you have studied.

Chapter 2 considers how we can understand the form of the NMR spectrum in terms of the underlying nuclear spin energy levels. Although this approach is more complex than the familiar “successive splitting” method for constructing multiplets it does help us understand how to think about multiplets in terms of “active” and “passive” spins. This approach also makes it possible to understand the form of multiple quantum spectra, which will be useful to us later on in the course. The chapter closes with a discussion of strongly coupled spectra and how they can be analysed.

Chapter 3 introduces the vector model of NMR. This model has its limitations, but it is very useful for understanding how pulses excite NMR signals. We can also use the vector model to understand the basic, but very important, NMR experiments such as pulse-acquire, inversion recovery and most importantly the spin echo.

Chapter 4 is concerned with data processing. The signal we actually record in an NMR experiment is a function of time, and we have to convert this to the usual representation (intensity as a function of frequency) using Fourier transformation. There are quite a lot of useful manipulations that we can carry out on the data to enhance the sensitivity or resolution, depending on what we require. These manipulations are described and their limitations discussed.

Chapter 5 is concerned with how the spectrometer works. It is not necessary to understand this in great detail, but it does help to have some basic understanding of what is going on when we “shim the magnet” or “tune the probe”. In this chapter we also introduce some important ideas about how the NMR signal is turned into a digital form, and the consequences that this has.

Chapter 6 introduces the product operator formalism for analysing NMR experiments. This approach is quantum mechanical, in contrast to the semi-classical approach taken by the vector model. We will see that the formalism is well adapted to describing pulsed NMR experiments, and that despite its quantum mechanical rigour it retains a relatively intuitive approach. Using product operators we can describe important phenomena such as the evolution of couplings during spin echoes, coherence transfer and the generation of multiple quantum coherences.

Chapter 7 puts the tools from Chapter 6 to immediate use in analysing and understanding two-dimensional spectra. Such spectra have proved to be enormously useful in structure determination, and are responsible for the explosive growth of NMR over the past 20 years or so. We will concentrate on the most important types of spectra, such as COSY and HMQC, analysing these in some detail.

Chapter 8 considers the important topic of relaxation in NMR. We start out by considering the effects of relaxation, concentrating in particular on the very important nuclear Overhauser effect. We then go on to consider the

sources of relaxation and how it is related to molecular properties.

Chapter 9 does not form a part of the course, but is an optional advanced topic. The chapter is concerned with the two methods used in multiple pulse NMR to select a particular outcome in an NMR experiment: phase cycling and field gradient pulses. An understanding of how these work is helpful in getting to grips with the details of how experiments are actually run.

Texts

There are innumerable books written about NMR. Many of these avoid any serious attempt to describe how the experiments work, but rather concentrate on the interpretation of various kinds of spectra. An excellent example of this kind of book is J. K. M. Sanders and B. K. Hunter *Modern NMR Spectroscopy* (OUP).

There are also a number of texts which take a more theory-based approach, at a number of different levels. Probably the best of the more elementary books is P. J. Hore *Nuclear Magnetic Resonance* (OUP).

For a deeper understanding you can do no better than refer to the book recently published by M. H. Levitt *Spin Dynamics* (Wiley).

Acknowledgements

Chapters 2 to 5 were prepared for a graduate course given at the Department of Chemistry, University of California Irvine in the spring of 2002. I wish to express my thanks to Professor AJ Shaka, and to his colleagues at Irvine, for the invitation to give this course and for hosting a period of sabbatical leave.

Chapters 6, 7 and 8 are modified from notes prepared for summer schools held in Mishima and Sapporo (Japan) in 1998 and 1999; thanks are due to Professor F Inagaki for the opportunity to present this material.

Chapter 9 was originally prepared (in a somewhat different form) for an EMBO course held in Turin (Italy) in 1995. It has been modified subsequently for the courses in Japan mentioned above and for another EMBO course held in Lucca in 2000. Once again I am grateful to the organizers and sponsors of these meetings for the opportunity to present this material.

Finally I wish to thank Professor Dr Miquel Pons Vallès and the Department of Organic Chemistry, University of Barcelona, for the opportunity to present this course. Financial support from the International Graduate School of Catalonia is gratefully acknowledged.

James Keeler
University of Cambridge, Department of Chemistry
March 2004
James.Keeler@ch.cam.ac.uk
www-keeler.ch.cam.ac.uk

2 NMR and energy levels

The picture that we use to understand most kinds of spectroscopy is that molecules have a set of *energy levels* and that the lines we see in spectra are due to *transitions* between these energy levels. Such a transition can be caused by a photon of light whose frequency, ν , is related to the energy gap, ΔE , between the two levels according to:

$$\Delta E = h\nu$$

where h is a universal constant known as Planck's constant. For the case shown in Fig. 2.1, $\Delta E = E_2 - E_1$.

In NMR spectroscopy we tend not to use this approach of thinking about energy levels and the transitions between them. Rather, we use different rules for working out the appearance of multiplets and so on. However, it is useful, especially for understanding more complex experiments, to think about how the familiar NMR spectra we see are related to energy levels. To start with we will look at the energy levels of just one spin and then move on quickly to look at two and three coupled spins. In such spin systems, as they are known, we will see that in principle there are other transitions, called *multiple quantum transitions*, which can take place. Such transitions are not observed in simple NMR spectra, but we can detect them indirectly using two-dimensional experiments; there are, as we shall see, important applications of such multiple quantum transitions.

Finally, we will look at *strongly coupled spectra*. These are spectra in which the simple rules used to construct multiplets no longer apply because the shift differences between the spins have become small compared to the couplings. The most familiar effect of strong coupling is the “roofing” or “tilting” of multiplets. We will see how such spectra can be analysed in some simple cases.

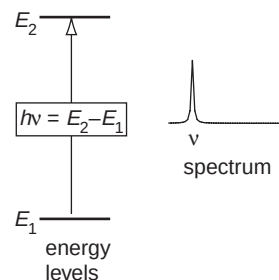


Fig. 2.1 A line in the spectrum is associated with a transition between two energy levels.

2.1 Frequency and energy: sorting out the units

NMR spectroscopists tend to use some rather unusual units, and so we need to know about these and how to convert from one to another if we are not to get into a muddle.

Chemical shifts

It is found to a very good approximation that the frequencies at which NMR absorptions (lines) occur scale linearly with the magnetic field strength. So, if the line from TMS comes out on one spectrometer at 400 MHz, doubling the magnetic field will result in it coming out at 800 MHz. If we wanted to quote the NMR frequency it would be inconvenient to have to specify the

exact magnetic field strength as well. In addition, the numbers we would have to quote would not be very memorable. For example, would you like to quote the shift of the protons in benzene as 400.001234 MHz?

We neatly side-step both of these problems by quoting the *chemical shift* relative to an agreed *reference compound*. For example, in the case of proton NMR the reference compound is TMS. If the frequency of the line we are interested in is ν (in Hz) and the frequency of the line from TMS is ν_{TMS} (also in Hz), the chemical shift of the line is computed as:

$$\delta = \frac{\nu - \nu_{\text{TMS}}}{\nu_{\text{TMS}}}. \quad (2.1)$$

As all the frequencies scale with the magnetic field, this ratio is independent of the magnetic field strength. Typically, the chemical shift is rather small so it is common to multiply the value for δ by 10^6 and then quote its value in *parts per million*, or ppm. With this definition the chemical shift of the reference compound is 0 ppm.

$$\delta_{\text{ppm}} = 10^6 \times \frac{\nu - \nu_{\text{TMS}}}{\nu_{\text{TMS}}}. \quad (2.2)$$

Sometimes we want to convert from shifts in ppm to frequencies. Suppose that there are two peaks in the spectrum at shifts δ_1 and δ_2 in ppm. What is the frequency separation between the two peaks? It is easy enough to work out what it is in ppm, it is just $(\delta_2 - \delta_1)$. Writing this difference out in terms of the definition of chemical shift given in Eq. 2.2 we have:

$$\begin{aligned} (\delta_2 - \delta_1) &= 10^6 \times \frac{\nu_2 - \nu_{\text{TMS}}}{\nu_{\text{TMS}}} - 10^6 \times \frac{\nu_1 - \nu_{\text{TMS}}}{\nu_{\text{TMS}}} \\ &= 10^6 \times \frac{\nu_2 - \nu_1}{\nu_{\text{TMS}}}. \end{aligned}$$

Multiplying both sides by ν_{TMS} now gives us what we want:

$$(\nu_2 - \nu_1) = 10^{-6} \times \nu_{\text{TMS}} \times (\delta_2 - \delta_1).$$

It is often sufficiently accurate to replace ν_{TMS} with the spectrometer *reference frequency*, about which we will explain later.

If we want to change the ppm scale of a spectrum into a frequency scale we need to decide where zero is going to be. One choice for the zero frequency point is the line from the reference compound. However, there are plenty of other possibilities so it is as well to regard the zero point on the frequency scale as arbitrary.

Angular frequency

Frequencies are most commonly quoted in Hz, which is the same as “per second” or s^{-1} . Think about a point on the edge of a disc which is rotating about its centre. If the disc is moving at a constant speed, the point returns to the same position at regular intervals each time it has completed 360° of rotation. The time taken for the point to return to its original position is called the *period*, τ .

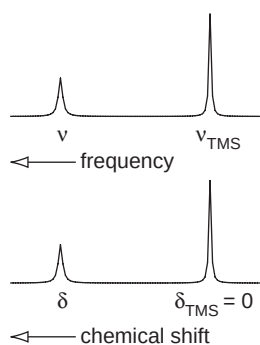


Fig. 2.2 An NMR spectrum can be plotted as a function of frequency, but it is more convenient to use the chemical shift scale in which frequencies are expressed relative to that of an agreed reference compound, such as TMS in the case of proton spectra.

The frequency, ν , is simply the inverse of the period:

$$\nu = \frac{1}{\tau}.$$

For example, if the period is 0.001 s, the frequency is $1/0.001 = 1000$ Hz.

There is another way of expressing the frequency, which is in *angular units*. Recall that 360° is 2π radians. So, if the point completes a rotation in τ seconds, we can say that it has rotated through 2π radians in τ seconds. The angular frequency, ω , is given by

$$\omega = \frac{2\pi}{\tau}.$$

The units of this frequency are “radians per second” or rad s^{-1} . ν and ω are related via

$$\nu = \frac{\omega}{2\pi} \quad \text{or} \quad \omega = 2\pi\nu.$$

We will find that angular frequencies are often the most natural units to use in NMR calculations. Angular frequencies will be denoted by the symbols ω or Ω whereas frequencies in Hz will be denoted ν .

Energies

A photon of frequency ν has energy E given by

$$E = h\nu$$

where h is Planck’s constant. In SI units the frequency is in Hz and h is in J s^{-1} . If we want to express the frequency in angular units then the relationship with the energy is

$$\begin{aligned} E &= h \frac{\omega}{2\pi} \\ &= \hbar \omega \end{aligned}$$

where $\hbar$ (pronounced “h bar” or “h cross”) is Planck’s constant divided by 2π .

The point to notice here is that frequency, in either Hz or rad s^{-1} , is directly proportional to energy. So, there is really nothing wrong with quoting energies in frequency units. All we have to remember is that there is a factor of h or $\hbar$ needed to convert to Joules if we need to. It turns out to be much more convenient to work in frequency units throughout, and so this is what we will do. So, do not be concerned to see an energy expressed in Hz or rad s^{-1} .

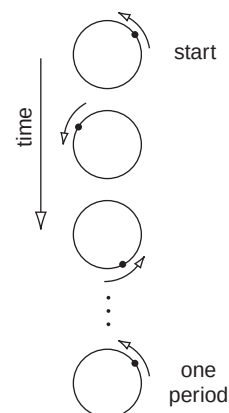


Fig. 2.3 A point at the edge of a circle which is moving at a constant speed returns to its original position after a time called the *period*. During each period the point moves through 2π radians or 360° .

2.2 Nuclear spin and spin states

NMR spectroscopy arises from the fact that nuclei have a property known as *spin*; we will not concern ourselves with where this comes from, but just take it as a fact. Quantum mechanics tells us that this nuclear spin is characterised by a nuclear spin quantum number, I . For all the nuclei that we are going

Strictly, α is the low energy state for nuclei with a positive gyromagnetic ratio, more of which below.

to be concerned with, $I = \frac{1}{2}$, although other values are possible. A spin-half nucleus has an interaction with a magnetic field which gives rise to two energy levels; these are characterised by another quantum number m which quantum mechanics tells us is restricted to the values $-I$ to I in integer steps. So, in the case of a spin-half, there are only two values of m , $-\frac{1}{2}$ and $+\frac{1}{2}$.

By tradition in NMR the energy level (or *state*, as it is sometimes called) with $m = \frac{1}{2}$ is denoted α and is sometimes described as “spin up”. The state with $m = -\frac{1}{2}$ is denoted β and is sometimes described as “spin down”. For the nuclei we are interested in, the α state is the one with the lowest energy.

If we have two spins in our molecule, then each spin can be in the α or β state, and so there are four possibilities: $\alpha_1\alpha_2$, $\alpha_1\beta_2$, $\beta_1\alpha_2$ and $\beta_1\beta_2$. These four possibilities correspond to four energy levels. Note that we have added a subscript 1 or 2 to differentiate the two nuclei, although often we will dispense with these and simply take it that the first spin state is for spin 1 and the second for spin 2. So $\alpha\beta$ implies $\alpha_1\beta_2$ etc.

We can continue the same process for three or more spins, and as each spin is added the number of possible combinations increases. So, for example, for three spins there are 8 combinations leading to 8 energy levels. It should be noted here that there is only a one-to-one correspondence between these spin state combinations and the energy levels in the case of weak coupling, which we will assume from now on. Further details are to be found in section 2.6.

2.3 One spin

There are just two energy levels for a single spin, and these can be labelled with either the m value of the labels α and β . From quantum mechanics we can show that the energies of these two levels, E_α and E_β , are:

$$E_\alpha = +\frac{1}{2}\nu_{0,1} \quad \text{and} \quad E_\beta = -\frac{1}{2}\nu_{0,1}$$

where $\nu_{0,1}$ is the *Larmor frequency* of spin 1 (we will need the 1 later on, but it is a bit superfluous here as we only have one spin). In fact it is easy to see that the energies just depend on m :

$$E_m = m\nu_{0,1}.$$

You will note here that, as explained in section 2.1 we have written the energies in frequency units. The Larmor frequency depends on a quantity known as the *gyromagnetic ratio*, γ , the chemical shift δ , and the strength of the applied magnetic field, B_0 :

$$\nu_{0,1} = -\frac{1}{2\pi}\gamma_1(1 + \delta_1)B_0 \quad (2.3)$$

where again we have used the subscript 1 to distinguish the nucleus. The magnetic field is normally given in units of Tesla (symbol T). The gyromagnetic ratio is characteristic of a particular kind of nucleus, such as proton or

This negative Larmor frequency for nuclei with a positive γ sometimes seems a bit unnatural and awkward, but to be consistent we need to stick with this convention. We will see in a later chapter that all this negative frequency really means is that the spin precesses in a particular sense.

carbon-13; its units are normally $\text{rad s}^{-1} \text{ T}^{-1}$. In fact, γ can be positive or negative; for the commonest nuclei (protons and carbon-13) it is positive. For such nuclei we therefore see that the Larmor frequency is negative.

To take a specific example, for protons $\gamma = +2.67 \times 10^8 \text{ rad s}^{-1} \text{ T}^{-1}$, so in a magnetic field of 4.7 T the Larmor frequency of a spin with chemical shift zero is

$$\begin{aligned}\nu_0 &= -\frac{1}{2\pi}\gamma(1 + \delta)B_0 \\ &= -\frac{1}{2\pi} \times 2.67 \times 10^8 \times 4.7 = -200 \times 10^6 \text{ Hz}.\end{aligned}$$

In other words, the Larmor frequency is -200 MHz .

We can also calculate the Larmor frequency in angular units, ω_0 , in which case the factor of $1/2\pi$ is not needed:

$$\omega_0 = -\gamma(1 + \delta)B_0$$

which gives a value of $1.255 \times 10^9 \text{ rad s}^{-1}$.

Spectrum

As you may know from other kinds of spectroscopy you have met, only certain transitions are *allowed* i.e. only certain ones actually take place. There are usually rules – called *selection rules* – about which transitions can take place; these rules normally relate to the quantum numbers which are characteristic of each state or energy level.

In the case of NMR, the selection rule refers to the quantum number m : only transitions in which m changes by one (up or down) are allowed. This is sometimes expressed as

$$\begin{aligned}\Delta m &= m(\text{initial state}) - m(\text{final state}) \\ &= \pm 1.\end{aligned}$$

Another way of saying this is that one spin can flip between “up” and “down” or vice versa.

In the case of a single spin-half, the change in m between the two states is $(+\frac{1}{2} - (-\frac{1}{2})) = 1$ so the transition is allowed. We can now simply work out the frequency of the allowed transition:

$$\begin{aligned}\nu_{\alpha\beta} &= E_\beta - E_\alpha \\ &= -\frac{1}{2}\nu_{0,1} - (+\frac{1}{2}\nu_{0,1}) \\ &= -\nu_{0,1}.\end{aligned}$$

Note that we have taken the energy of the upper state minus that of the lower state. In words, therefore, we see one transition at the minus the Larmor frequency, $-\nu_{0,1}$.

You would be forgiven for thinking that this is all an enormous amount of effort to come up with something very simple! However, the techniques and ideas developed in this section will enable us to make faster progress with the case of two and three coupled spins, which we consider next.

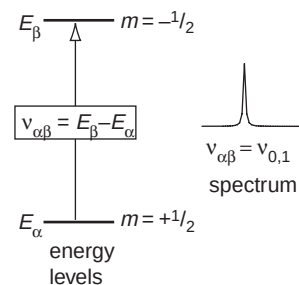


Fig. 2.4 The transition between the two energy levels of a spin-half is allowed, and results in a single line at the Larmor frequency of the spin.

2.4 Two spins

We know that the spectrum of two coupled spins consists of two doublets, each split by the same amount, one centred at the chemical shift of the first spin and one at the shift of the second. The splitting of the doublets is the scalar coupling, J_{12} , quoted in Hz; the subscripts indicate which spins are involved. We will write the shifts of the two spins as δ_1 and δ_2 , and give the corresponding Larmor frequencies, $\nu_{0,1}$ and $\nu_{0,2}$ as:

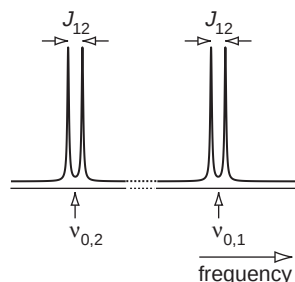


Fig. 2.5 Schematic spectrum of two coupled spins showing two doublets with equal splittings. As indicated by the dashed lines, the separation of the Larmor frequencies is much larger than the coupling between the spins.

$$\nu_{0,1} = -\frac{1}{2\pi}\gamma_1(1 + \delta_1)B_0$$

$$\nu_{0,2} = -\frac{1}{2\pi}\gamma_2(1 + \delta_2)B_0.$$

If the two nuclei are of the same type, such a proton, then the two gyromagnetic ratios are equal; such a two spin system would be described as *homonuclear*. The opposite case is where the two nuclei are of different types, such as proton and carbon-13; such a spin system is described as *heteronuclear*.

Energy levels

As was already described in section 2.2, there are four possible combinations of the spin states of two spins and these combinations correspond to four energy levels. Their energies are given in the following table:

number	spin states	energy
1	$\alpha\alpha$	$+\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} + \frac{1}{4}J_{12}$
2	$\alpha\beta$	$+\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} - \frac{1}{4}J_{12}$
3	$\beta\alpha$	$-\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} - \frac{1}{4}J_{12}$
4	$\beta\beta$	$-\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} + \frac{1}{4}J_{12}$

The second column gives the spin states of spins 1 and 2, in that order. It is easy to see that these energies have the general form:

$$E_{m_1 m_2} = m_1 \nu_{0,1} + m_2 \nu_{0,2} + m_1 m_2 J_{12}$$

where m_1 and m_2 are the m values for spins 1 and 2, respectively.

For a homonuclear system $\nu_{0,1} \approx \nu_{0,2}$; also both Larmor frequencies are much greater in magnitude than the coupling (the Larmor frequencies are of the order of hundreds of MHz, while couplings are at most a few tens of Hz). Therefore, under these circumstances, the energies of the $\alpha\beta$ and $\beta\alpha$ states are rather similar, but very different from the other two states. For a heteronuclear system, in which the Larmor frequencies differ significantly, the four levels are all at markedly different energies. These points are illustrated in Fig. 2.6.

Spectrum

The selection rule is the same as before, but this time it applies to the quantum number M which is found by adding up the m values for each of the spins. In

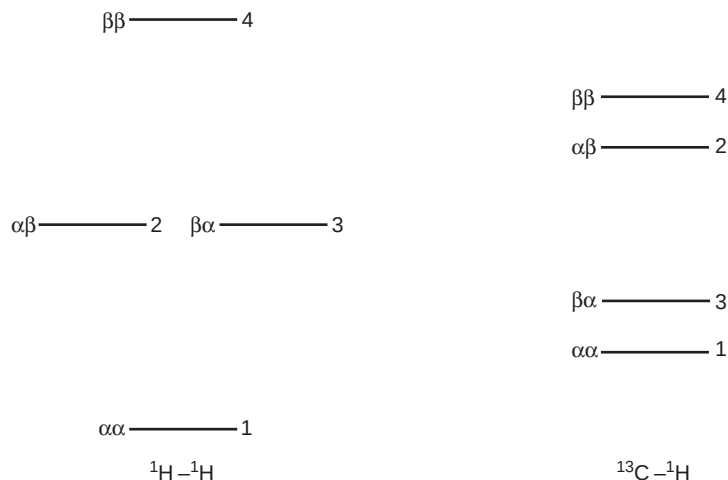


Fig. 2.6 Energy levels, drawn approximately to scale, for two spin systems. On the left is shown a homonuclear system (two protons); on this scale the $\alpha\beta$ and $\beta\alpha$ states have the same energy. On the right is the case for a carbon-13 – proton pair. The Larmor frequency of proton is about four times that of carbon-13, and this is clear reflected in the diagram. The $\alpha\beta$ and $\beta\alpha$ states now have substantially different energies.

this case:

$$M = m_1 + m_2.$$

The resulting M values for the four levels are:

number	spin states	M
1	$\alpha\alpha$	1
2	$\alpha\beta$	0
3	$\beta\alpha$	0
4	$\beta\beta$	-1

The selection rule is that $\Delta M = \pm 1$, i.e. the value of M can change up or down by one unit. This means that the allowed transitions are between levels 1 & 2, 3 & 4, 1 & 3 and 2 & 4. The resulting frequencies are easily worked out; for example, the 1-2 transition:

$$\begin{aligned}
 \nu_{12} &= E_2 - E_1 \\
 &= +\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} - \frac{1}{4}J_{12} - \left(\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} + \frac{1}{4}J_{12}\right) \\
 &= -\nu_{0,2} - \frac{1}{2}J_{12}.
 \end{aligned}$$

Throughout we will use the convention that when computing the transition frequency we will take the energy of the upper state minus the energy of the lower:
 $\Delta E = E_{\text{upper}} - E_{\text{lower}}$.

The complete set of transitions are:

transition	spin states	frequency
1 $\rightarrow$ 2	$\alpha\alpha \rightarrow \alpha\beta$	$-\nu_{0,2} - \frac{1}{2}J_{12}$
3 $\rightarrow$ 4	$\beta\alpha \rightarrow \beta\beta$	$-\nu_{0,2} + \frac{1}{2}J_{12}$
1 $\rightarrow$ 3	$\alpha\alpha \rightarrow \beta\alpha$	$-\nu_{0,1} - \frac{1}{2}J_{12}$
2 $\rightarrow$ 4	$\alpha\beta \rightarrow \beta\beta$	$-\nu_{0,1} + \frac{1}{2}J_{12}$

The energy levels and corresponding schematic spectrum are shown in Fig. 2.7. There is a lot we can say about this spectrum. Firstly, each allowed

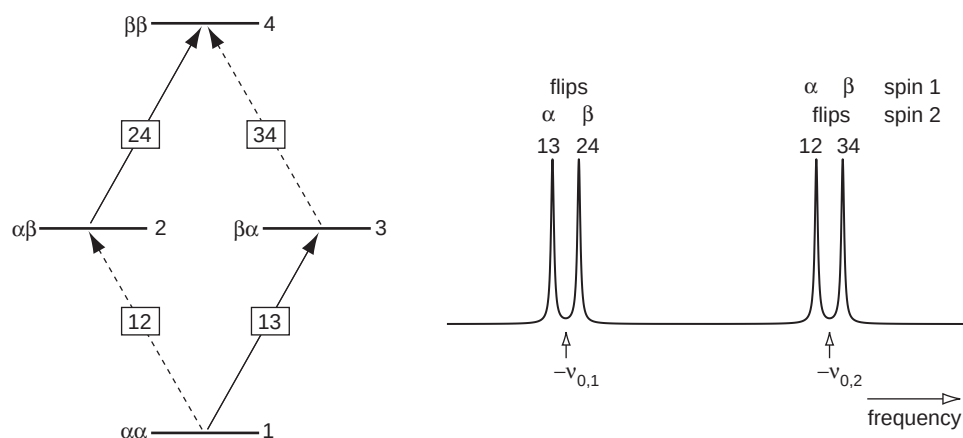


Fig. 2.7 On the left, the energy levels of a two-spin system; the arrows show the allowed transitions: solid lines for transitions in which spin 1 flips and dotted for those in which spin 2 flips. On the right, the corresponding spectrum; it is assumed that the Larmor frequency of spin 2 is greater in magnitude than that of spin 1 and that the coupling J_{12} is positive.

transition corresponds to one of the spins flipping from one spin state to the other, while the spin state of the other spin remains fixed. For example, transition 1–2 involves a spin 2 going from α to β whilst spin 1 remains in the α state. In this transition we say that spin 2 is *active* and spin 1 is *passive*. As spin 2 flips in this transition, it is not surprising that the transition forms one part of the doublet for spin 2.

Transition 3–4 is similar to 1–2 except that the passive spin (spin 1) is in the β state; this transition forms the second line of the doublet for spin 2. This discussion illustrates a very important point, which is that the lines of a multiplet can be associated with different spin states of the coupled (passive) spins. We will use this kind of interpretation very often, especially when considering two-dimensional spectra.

The two transitions in which spin 1 flips are 1–3 and 2–4, and these are associated with spin 2 being in the α and β spin states, respectively. Which spin flips and the spins states of the passive spins are shown in Fig. 2.7.

What happens if the coupling is negative? If you work through the table you will see that there are still four lines at the same frequencies as before. All that changes is the labels of the lines. So, for example, transition 1–2 is now the right line of the doublet, rather than the left line. From the point of view of the spectrum, what swaps over is the spin state of the passive spin associated with each line of the multiplet. The overall appearance of the spectrum is therefore independent of the sign of the coupling constant.

Multiple quantum transitions

There are two more transitions in our two-spin system which are not allowed by the usual selection rule. The first is between states 1 and 4 ($\alpha\alpha \rightarrow \beta\beta$) in

which both spins flip. The ΔM value is 2, so this is called a *double-quantum* transition. Using the same terminology, all of the allowed transitions described above, which have $\Delta M = 1$, are *single-quantum* transitions. From the table of energy levels it is easy to work out that its frequency is $(-\nu_{0,1} - \nu_{0,2})$ i.e. the sum of the Larmor frequencies. Note that the coupling has no effect on the frequency of this line.

The second transition is between states 2 and 3 ($\alpha\beta \rightarrow \beta\alpha$); again, both spins flip. The ΔM value is 0, so this is called a *zero-quantum* transition, and its frequency is $(-\nu_{0,1} + \nu_{0,2})$ i.e. the difference of the Larmor frequencies. As with the double-quantum transition, the coupling has no effect on the frequency of this line.

In a two spin system the double- and zero-quantum spectra are not especially interesting, but we will see in a three-spin system that the situation is rather different. We will also see later on that in two-dimensional spectra we can exploit to our advantage the special properties of these multiple-quantum spectra.

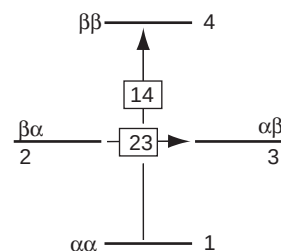


Fig. 2.8 In a two-spin system there is one double quantum transition (1–4) and one zero-quantum transition (2–3); the frequency of neither of these transitions are affected by the size of the coupling between the two spins.

2.5 Three spins

If we have three spins, each of which is coupled to the other two, then the spectrum consists of three doublets of doublets, one centred at the shift of each of the three spins; the *spin topology* is shown in Fig. 2.9. The appearance of these multiplets will depend on the relative sizes of the couplings. For example, if $J_{12} = J_{13}$ the doublet of doublets from spin 1 will collapse to a 1:2:1 triplet. On the other hand, if $J_{12} = 0$, only doublets will be seen for spin 1 and spin 2, but spin 3 will still show a doublet of doublets.

Energy levels

Each of the three spins can be in the α or β spin state, so there are a total of 8 possible combinations corresponding to 8 energy levels. The energies are given by:

$$E_{m_1 m_2 m_3} = m_1 \nu_{0,1} + m_2 \nu_{0,2} + m_3 \nu_{0,3} + m_1 m_2 J_{12} + m_1 m_3 J_{13} + m_2 m_3 J_{23}$$

where m_i is the value of the quantum number m for the i th spin. The energies and corresponding M values ($= m_1 + m_2 + m_3$) are shown in the table:

number	spin states	M	energy
1	$\alpha\alpha\alpha$	$\frac{3}{2}$	$+\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} + \frac{1}{2}\nu_{0,3} + \frac{1}{4}J_{12} + \frac{1}{4}J_{13} + \frac{1}{4}J_{23}$
2	$\alpha\beta\alpha$	$\frac{1}{2}$	$+\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} + \frac{1}{2}\nu_{0,3} - \frac{1}{4}J_{12} + \frac{1}{4}J_{13} - \frac{1}{4}J_{23}$
3	$\beta\alpha\alpha$	$\frac{1}{2}$	$-\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} + \frac{1}{2}\nu_{0,3} - \frac{1}{4}J_{12} - \frac{1}{4}J_{13} + \frac{1}{4}J_{23}$
4	$\beta\beta\alpha$	$-\frac{1}{2}$	$-\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} + \frac{1}{2}\nu_{0,3} + \frac{1}{4}J_{12} - \frac{1}{4}J_{13} - \frac{1}{4}J_{23}$
5	$\alpha\alpha\beta$	$\frac{1}{2}$	$+\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} - \frac{1}{2}\nu_{0,3} + \frac{1}{4}J_{12} - \frac{1}{4}J_{13} - \frac{1}{4}J_{23}$
6	$\alpha\beta\beta$	$-\frac{1}{2}$	$+\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} - \frac{1}{2}\nu_{0,3} - \frac{1}{4}J_{12} - \frac{1}{4}J_{13} + \frac{1}{4}J_{23}$
7	$\beta\alpha\beta$	$-\frac{1}{2}$	$-\frac{1}{2}\nu_{0,1} + \frac{1}{2}\nu_{0,2} - \frac{1}{2}\nu_{0,3} - \frac{1}{4}J_{12} + \frac{1}{4}J_{13} - \frac{1}{4}J_{23}$
8	$\beta\beta\beta$	$-\frac{3}{2}$	$-\frac{1}{2}\nu_{0,1} - \frac{1}{2}\nu_{0,2} - \frac{1}{2}\nu_{0,3} + \frac{1}{4}J_{12} + \frac{1}{4}J_{13} + \frac{1}{4}J_{23}$

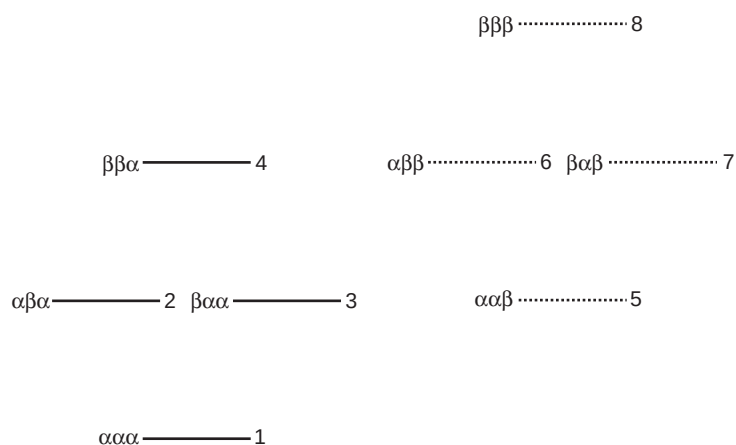


Fig. 2.10 Energy levels for a homonuclear three-spin system. The levels can be grouped into two sets of four: those with spin 3 in the α state (shown on the left with solid lines) and those with spin 3 in the β state, shown on the right (dashed lines).

We have grouped the energy levels into two groups of four; the first group all have spin 3 in the α state and the second have spin 3 in the β state. The energy levels (for a homonuclear system) are shown schematically in Fig. 2.10.

Spectrum

The selection rule is as before, that is M can only change by 1. However, in the case of more than two spins, there is the additional constraint that *only one* spin can flip. Applying these rules we see that there are four allowed transitions in which spin 1 flips: 1–3, 2–4, 5–7 and 6–8. The frequencies of these lines can easily be worked out from the table of energy levels on page 2-10. The results are shown in the table, along with the spin states of the passive spins (2 and 3 in this case).

transition	state of spin 2	state of spin 3	frequency
1-3	α	α	$-\nu_{0,1} - \frac{1}{2}J_{12} - \frac{1}{2}J_{13}$
2-4	β	α	$-\nu_{0,1} + \frac{1}{2}J_{12} - \frac{1}{2}J_{13}$
5-7	α	β	$-\nu_{0,1} - \frac{1}{2}J_{12} + \frac{1}{2}J_{13}$
6-8	β	β	$-\nu_{0,1} + \frac{1}{2}J_{12} + \frac{1}{2}J_{13}$

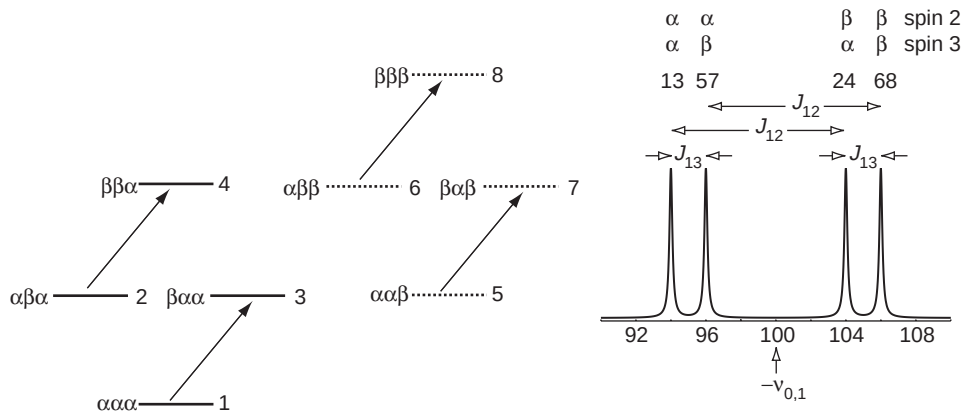


Fig. 2.11 Energy levels for a three-spin system showing by the arrows the four allowed transitions which result in the doublet of doublets at the shift of spin 1. The schematic multiplet is shown on the right, where it has been assuming that $\nu_{0,1} = -100$ Hz, $J_{12} = 10$ Hz and $J_{13} = 2$ Hz. The multiplet is labelled with the spin states of the passive spins.

These four transitions form the four lines of the multiplet (a doublet of doublets) at the shift of spin 1. The schematic spectrum is illustrated in Fig. 2.11. As in the case of a two-spin system, we can label each line of the multiplet with the spin states of the passive spins – in the case of the multiplet from spin 1, this means the spin states of spins 2 and 3. In the same way, we can identify the four transitions which contribute to the multiplet from spin 2 (1-2, 3-4, 5-6 and 7-8) and the four which contribute to that from spin 3 (1-5, 3-7, 2-6 and 4-8).

Subspectra

One way of thinking about the spectrum from the three-spin system is to divide up the lines in the multiplets for spins 1 and 2 into two groups or *subspectra*. The first group consists of the lines which have spin 3 in the α state and the second group consists of the lines which have spin 3 in the β state. This separation is illustrated in Fig. 2.12.

There are four lines which have spin-3 in the α state, and as can be seen from the spectrum these form two doublets with a common separation of J_{12} . However, the two doublets are not centred at $-\nu_{0,1}$ and $-\nu_{0,2}$, but at $(-\nu_{0,1} - \frac{1}{2}J_{13})$ and $(-\nu_{0,2} - \frac{1}{2}J_{23})$. We can define an *effective* Larmor frequency for spin 1 with spin 3 in the α spin state, $\nu_{0,1}^{\alpha 3}$, as

$$\nu_{0,1}^{\alpha 3} = \nu_{0,1} + \frac{1}{2}J_{13}$$

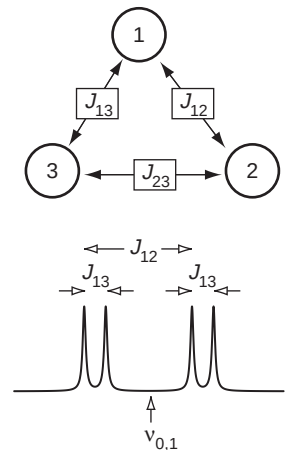


Fig. 2.9 The topology – that is the number of spins and the couplings between them – for a three-spin system in which each spin is coupled to both of the others. The resulting spectrum consists of three doublets of doublets, a typical example of which is shown for spin 1 with the assumption that J_{12} is greater than J_{13} .

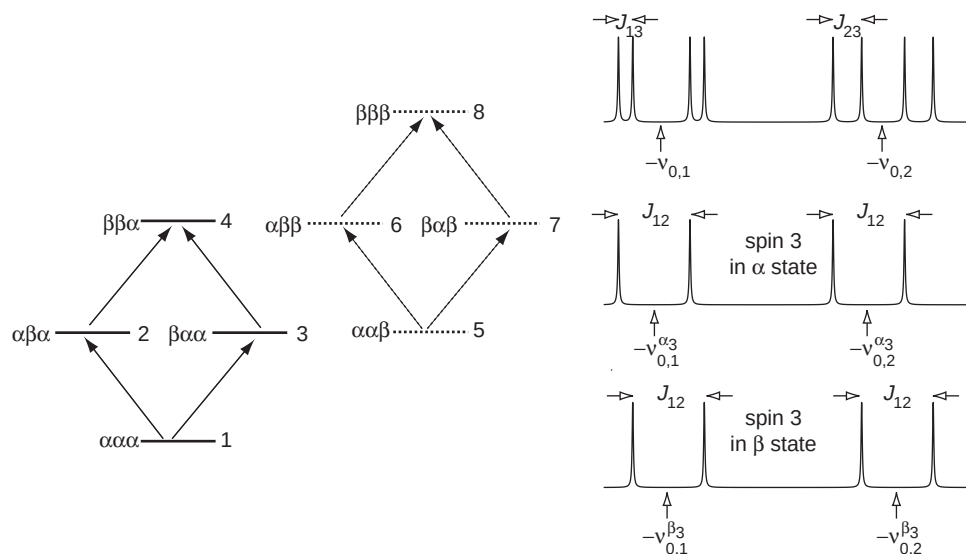


Fig. 2.12 Illustration of the division of the two multiplets from spins 1 and 2 into subspectra according to the spin state of spin 3. The transitions associated with spin 3 in the α state (indicated by the full lines on the energy level diagram) give rise to a pair of doublets, but with their centres shifted from the Larmor frequencies by half the coupling to spin 3. The same is true of those transitions associated with spin 3 being in the β state (dashed lines), except that the shift is in the opposite direction.

and likewise for spin 2:

$$\nu_{0,2}^{\alpha_3} = \nu_{0,2} + \frac{1}{2}J_{23}.$$

The two doublets in the sub-spectrum corresponding to spin 3 being in the α state are thus centred at $-\nu_{0,1}^{\alpha_3}$ and $-\nu_{0,2}^{\alpha_3}$. Similarly, we can define effective Larmor frequencies for spin 3 being in the β state:

$$\nu_{0,1}^{\beta_3} = \nu_{0,1} - \frac{1}{2}J_{13} \quad \nu_{0,2}^{\beta_3} = \nu_{0,2} - \frac{1}{2}J_{23}.$$

The two doublets in the β sub-spectrum are centred at $-\nu_{0,1}^{\beta_3}$ and $-\nu_{0,2}^{\beta_3}$.

We can think of the spectrum of spin 1 and 2 as being composed of two subspectra, each from a two spin system but in which the Larmor frequencies are effectively shifted one way or the other by half the coupling to the third spin. This kind of approach is particularly useful when it comes to dealing with strongly coupled spin systems, section 2.6

Note that the separation of the spectra according to the spin state of spin 3 is arbitrary. We could just as well separate the two multiplets from spins 1 and 3 according to the spin state of spin 2.

Multiple quantum transitions

There are six transitions in which M changes by 2. Their frequencies are given in the table.

transition	initial state	final state	frequency
1-4	$\alpha\alpha\alpha$	$\beta\beta\alpha$	$-\nu_{0,1} - \nu_{0,2} - \frac{1}{2}J_{13} - \frac{1}{2}J_{23}$
5-8	$\alpha\alpha\beta$	$\beta\beta\beta$	$-\nu_{0,1} - \nu_{0,2} + \frac{1}{2}J_{13} + \frac{1}{2}J_{23}$
1-7	$\alpha\alpha\alpha$	$\beta\alpha\beta$	$-\nu_{0,1} - \nu_{0,3} - \frac{1}{2}J_{12} - \frac{1}{2}J_{23}$
2-8	$\alpha\beta\alpha$	$\beta\beta\beta$	$-\nu_{0,1} - \nu_{0,3} + \frac{1}{2}J_{12} + \frac{1}{2}J_{23}$
1-6	$\alpha\alpha\alpha$	$\alpha\beta\beta$	$-\nu_{0,2} - \nu_{0,3} - \frac{1}{2}J_{12} - \frac{1}{2}J_{13}$
3-8	$\beta\alpha\alpha$	$\beta\beta\beta$	$-\nu_{0,2} - \nu_{0,3} + \frac{1}{2}J_{12} + \frac{1}{2}J_{13}$

These transitions come in three pairs. Transitions 1-4 and 5-8 are centred at the sum of the Larmor frequencies of spins 1 and 2; this is not surprising as we note that in these transitions it is the spin states of both spins 1 and 2 which flip. The two transitions are separated by the *sum* of the couplings to spin 3 ($J_{13} + J_{23}$), but they are unaffected by the coupling J_{12} which is between the two spins which flip.

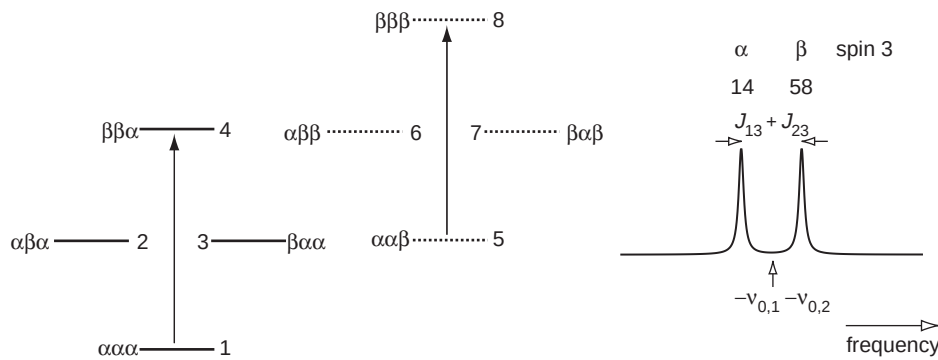


Fig. 2.13 There are two double quantum transitions in which spins 1 and 2 both flip (1-4 and 5-8). The two resulting lines form a doublet which is centred at the sum of the Larmor frequencies of spins 1 and 2 and which is split by the sum of the couplings to spin 3. As with the single-quantum spectra, we can associate the two lines of the doublet with different spin states of the third spin. It has been assumed that both couplings are positive.

We can describe these transitions as a kind of double quantum doublet. Spins 1 and 2 are both active in these transitions, and spin 3 is passive. Just as we did before, we can associate one line with spin 3 being in the α state (1-4) and one with it being in the β state (5-8). A schematic representation of the spectrum is shown in Fig. 2.13.

There are also six zero-quantum transitions in which M does not change. Like the double quantum transitions these group in three pairs, but this time centred around the *difference* in the Larmor frequencies of two of the spins. These zero-quantum doublets are split by the *difference* of the couplings to the spin which does not flip in the transitions. There are thus many similarities between the double- and zero-quantum spectra.

In a three spin system there is one triple-quantum transition, in which M changes by 3, between levels 1 ($\alpha\alpha\alpha$) and 8 ($\beta\beta\beta$). In this transition all of

the spins flip, and from the table of energies we can easily work out that its frequency is $-\nu_{0,1} - \nu_{0,2} - \nu_{0,3}$, i.e. the sum of the Larmor frequencies.

We see that the single-quantum spectrum consists of three doublets of doublets, the double-quantum spectrum of three doublets and the triple-quantum spectrum of a single line. This illustrates the idea that as we move to higher orders of multiple quantum, the corresponding spectra become simpler. This feature has been used in the analysis of some complex spin systems.

Combination lines

There are three more transitions which we have not yet described. For these, M changes by 1 but all three spins flip; they are called *combination lines*. Such lines are not seen in normal spectra but, like multiple quantum transitions, they can be detected indirectly using two-dimensional spectra. We will also see in section 2.6 that these lines may be observable in strongly coupled spectra. The table gives the frequencies of these three lines:

transition	initial state	final state	frequency
2-7	$\alpha\beta\alpha$	$\beta\alpha\beta$	$-\nu_{0,1} + \nu_{0,2} - \nu_{0,3}$
3-6	$\beta\alpha\alpha$	$\alpha\beta\beta$	$+\nu_{0,1} - \nu_{0,2} - \nu_{0,3}$
4-5	$\beta\beta\alpha$	$\alpha\alpha\beta$	$+\nu_{0,1} + \nu_{0,2} - \nu_{0,3}$

Notice that the frequencies of these lines are not affected by any of the couplings.

2.6 Strong coupling

So far all we have said about energy levels and spectra applies to what are called *weakly coupled* spin systems. These are spin systems in which the differences between the Larmor frequencies (in Hz) of the spins are much greater in magnitude than the magnitude of the couplings between the spins.

Under these circumstances the rules from predicting spectra are very simple – they are the ones you are already familiar with which you use for constructing multiplets. In addition, in the weak coupling limit it is possible to work out the energies of the levels present simply by making all possible combinations of spin states, just as we have done above. Finally, in this limit all of the lines in a multiplet have the same intensity.

If the separation of the Larmor frequencies is not sufficient to satisfy the weak coupling criterion, the system is said to be *strongly coupled*. In this limit *none* of the rules outlined in the previous paragraph apply. This makes predicting or analysing the spectra much more difficult than in the case of weak coupling, and really the only practical approach is to use computer simulation. However, it is useful to look at the spectrum from two strongly coupled spins as the spectrum is simple enough to analyse by hand and will reveal most of the of the crucial features of strong coupling.

We need to be careful here as Larmor frequencies, the differences between Larmor frequencies and the values of couplings can be positive or negative! In deciding whether or not a spectrum will be strongly coupled we need to compare the *magnitude* of the difference in the Larmor frequencies with the *magnitude* of the coupling.

You will often find that people talk of two spins being *strongly coupled* when what they really mean is the coupling between the two spins is *large*. This is sloppy usage; we will always use the term strong coupling in the sense described in this section.

It is a relatively simple exercise in quantum mechanics to work out the energy levels and hence frequencies and intensities of the lines from a strongly coupled two-spin system¹. These are given in the following table.

transition	frequency	intensity
1-2	$\frac{1}{2}D - \frac{1}{2}\Sigma - \frac{1}{2}J_{12}$	$(1 + \sin 2\theta)$
3-4	$\frac{1}{2}D - \frac{1}{2}\Sigma + \frac{1}{2}J_{12}$	$(1 - \sin 2\theta)$
1-3	$-\frac{1}{2}D - \frac{1}{2}\Sigma - \frac{1}{2}J_{12}$	$(1 - \sin 2\theta)$
2-4	$-\frac{1}{2}D - \frac{1}{2}\Sigma + \frac{1}{2}J_{12}$	$(1 + \sin 2\theta)$

In this table Σ is the sum of the Larmor frequencies:

$$\Sigma = \nu_{0,1} + \nu_{0,2}$$

and D is the *positive* quantity defined as

$$D^2 = (\nu_{0,1} - \nu_{0,2})^2 + J_{12}^2. \quad (2.4)$$

The angle θ is called the *strong coupling parameter* and is defined via

$$\sin 2\theta = \frac{J_{12}}{D}.$$

The first thing to do is to verify that these formulae give us the expected result when we impose the condition that the separation of the Larmor frequencies is large compared to the coupling. In this limit it is clear that

$$\begin{aligned} D^2 &= (\nu_{0,1} - \nu_{0,2})^2 + J_{12}^2 \\ &\approx (\nu_{0,1} - \nu_{0,2})^2 \end{aligned}$$

and so $D = (\nu_{0,1} - \nu_{0,2})$. Putting this value into the table above gives us exactly the frequencies we had before on page 2-7.

When D is very much larger than J_{12} the fraction J_{12}/D becomes small, and so $\sin 2\theta \approx 0$ ($\sin \phi$ goes to zero as ϕ goes to zero). Under these circumstances all of the lines have unit intensity. So, the weak coupling limit is regained.

Figure 2.14 shows a series of spectra computed using the above formulae in which the Larmor frequency of spin 1 is held constant while the Larmor frequency of spin 2 is progressively moved towards that of spin 1. This makes the spectrum more and more strongly coupled. The spectrum at the bottom is almost weakly coupled; the peaks are just about all the same intensity and where we expect them to be.

¹See, for example, Chapter 10 of *NMR: The Toolkit*, by P J Hore, J A Jones and S Wimperis (Oxford University Press, 2000)

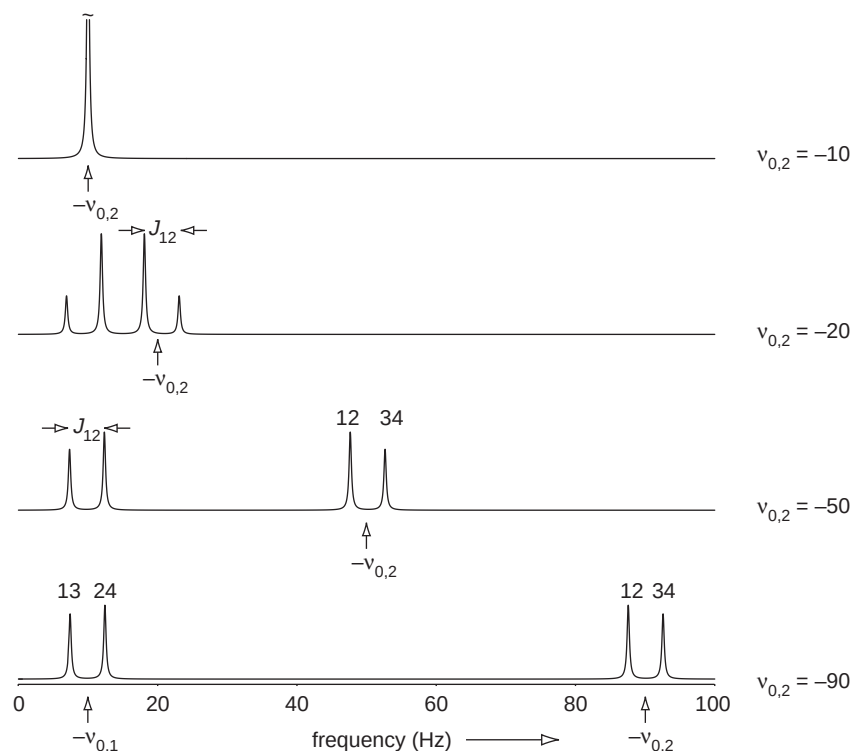


Fig. 2.14 A series of spectra of a two spin system in which the Larmor frequency of spin 1 is held constant and that of spin 2 is moved in closer to spin 1. The spectra become more and more strongly coupled showing a pronounced roof effect until in the limit that the two Larmor frequencies are equal only one line is observed. Note that as the “outer” lines get weaker the “inner” lines get proportionately stronger. The parameters used for these spectra were $\nu_{0,1} = -10$ Hz and $J_{12} = 5$ Hz; the peak in the top most spectrum has been truncated.

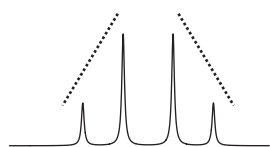


Fig. 2.15 The intensity distributions in multiplets from strongly-coupled spectra are such that the multiplets “tilt” towards one another; this is called the “roof” effect.

However, as the Larmor frequencies of the two spins get closer and closer together we notice two things: (1) the “outer” two lines get weaker and the “inner” two lines get stronger; (2) the two lines which originally formed the doublet are no longer symmetrically spaced about the Larmor frequency; in fact the stronger of the two lines moves progressively closer to the Larmor frequency. There is one more thing to notice which is not so clear from the spectra but is clear if one looks at the frequencies in the table. This is that the two lines that originally formed the spin 1 doublet are always separated by J_{12} ; the same is true for the other doublet.

These spectra illustrate the so-called *roof* effect in which the intensities of the lines in a strongly coupled multiplet tilt upwards towards the multiplet from the coupled spin, making a kind of roof; Fig. 2.15 illustrates the idea. The spectra in Fig 2.14 also illustrate the point that when the two Larmor frequencies are identical there is only one line seen in the spectrum and this is at this Larmor frequency. In this limit lines 1–2 and 2–4 both appear at the Larmor frequency and with intensity 2; lines 1–3 and 3–4 appear elsewhere but have intensity zero.

The “take home message” is that from such strongly coupled spectra we can easily measure the coupling, but the Larmor frequencies (the shifts) are no longer mid-way between the two lines of the doublet. In fact it is easy

enough to work out the Larmor frequencies using the following method; the idea is illustrated in Fig. 2.16.

If we denote the frequency of transition 1–2 as ν_{12} and so on, it is clear from the table that the frequency separation of the left-hand lines of the two multiplets (3–4 and 2–4) is D

$$\begin{aligned}\nu_{34} - \nu_{24} &= \left(\frac{1}{2}D - \frac{1}{2}\Sigma + \frac{1}{2}J_{12}\right) - \left(-\frac{1}{2}D - \frac{1}{2}\Sigma + \frac{1}{2}J_{12}\right) \\ &= D.\end{aligned}$$

The separation of the other two lines is also D . Remember we can easily measure J_{12} directly from the splitting, and so once we know D it is easy to compute $(\nu_{0,1} - \nu_{0,2})$ from its definition, Eqn. 2.4.

$$\begin{aligned}D^2 &= (\nu_{0,1} - \nu_{0,2})^2 + J_{12}^2 \\ \text{therefore } (\nu_{0,1} - \nu_{0,2}) &= \sqrt{D^2 - J_{12}^2}.\end{aligned}$$

Now we notice from the table on 2–15 that the sum of the frequencies of the two stronger lines (1–2 and 2–4) or the two weaker lines (3–4 and 2–4) gives us $-\Sigma$:

$$\begin{aligned}\nu_{12} + \nu_{24} &= \left(\frac{1}{2}D - \frac{1}{2}\Sigma - \frac{1}{2}J_{12}\right) + \left(-\frac{1}{2}D - \frac{1}{2}\Sigma + \frac{1}{2}J_{12}\right) \\ &= -\Sigma.\end{aligned}$$

Now we have a values for $\Sigma = (\nu_{0,1} + \nu_{0,2})$ and a value for $(\nu_{0,1} - \nu_{0,2})$ we can find $\nu_{0,1}$ and $\nu_{0,2}$ separately:

$$\nu_{0,1} = \frac{1}{2}(\Sigma + (\nu_{0,1} - \nu_{0,2})) \quad \nu_{0,2} = \frac{1}{2}(\Sigma - (\nu_{0,1} - \nu_{0,2})).$$

In this way we can extract the Larmor frequencies of the two spins (the shifts) and the coupling from the strongly coupled spectrum.

Notation for spin systems

There is a traditional notation for spin systems which it is sometimes useful to use. Each spin is given a letter (rather than a number), with a different letter for spins which have different Larmor frequencies (chemical shifts). If the Larmor frequencies of two spins are widely separated they are given letters which are widely separated in the alphabet. So, two weakly coupled spins are usually denoted as A and X; whereas three weakly coupled spins would be denoted AMX.

If spins are strongly coupled, then the convention is to used letters which are close in the alphabet. So, a strongly coupled two-spin system would be denoted AB and a strongly coupled three-spin system ABC. The notation ABX implies that two of the spins (A and B) are strongly coupled but the Larmor frequency of the third spin is widely separated from that of A and B.

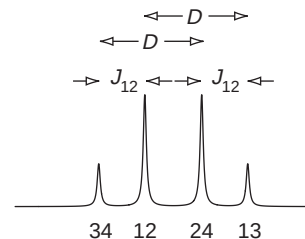


Fig. 2.16 The quantities J_{12} and D are readily measurable from the spectrum of two strongly coupled spins.

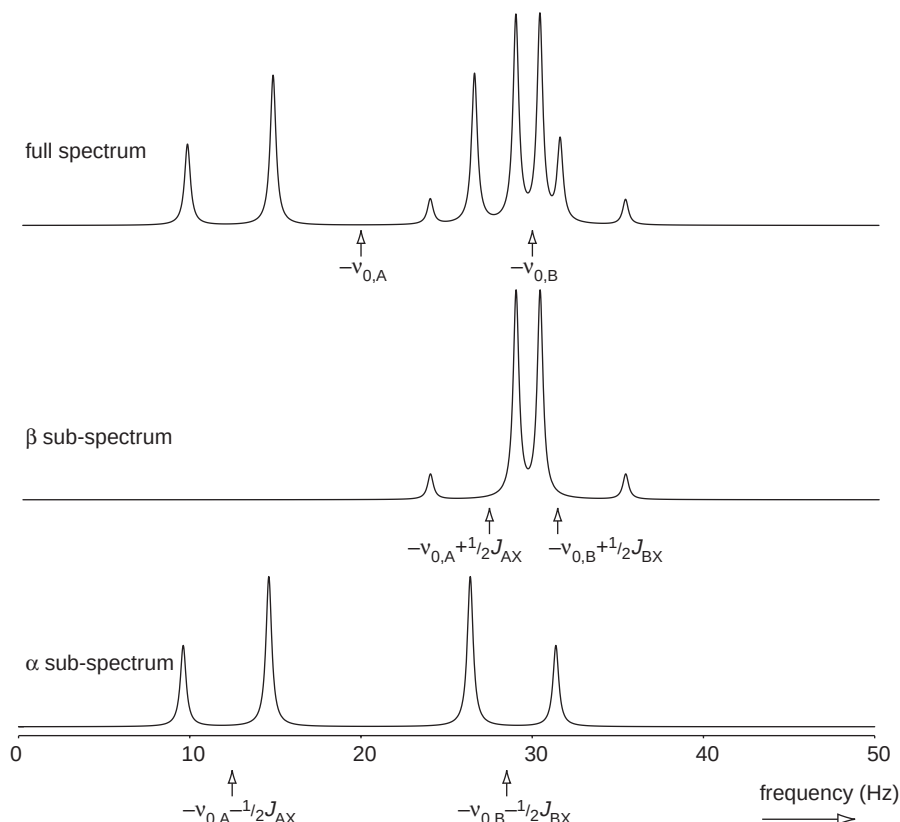


Fig. 2.17 AB parts of an ABX spectrum illustrating the decomposition into two sub-spectra with different effective Larmor frequencies (indicated by the arrows). The parameters used in the simulation were $\nu_{0,A} = -20$ Hz, $\nu_{0,B} = -30$ Hz, $J_{AB} = 5$ Hz, $J_{AX} = 15$ Hz and $J_{BX} = 3$ Hz.

The ABX spin system

We noted in section 2.5 that we could think about the spectrum of three coupled spins in terms of sub-spectra in which the Larmor frequencies were replaced by effective Larmor frequencies. This kind of approach is very useful for understanding the AB part of the ABX spectrum.

As spin X is weakly coupled to the others we can think of the AB part of the spectrum as two superimposed AB sub-spectra; one is associated with the X spin being in the α state and the other with the spin being in the β state. If spins A and B have Larmor frequencies $\nu_{0,A}$ and $\nu_{0,B}$, respectively, then one sub-spectrum has effective Larmor frequencies $\nu_{0,A} + \frac{1}{2}J_{AX}$ and $\nu_{0,B} + \frac{1}{2}J_{BX}$. The other has effective Larmor frequencies $\nu_{0,A} - \frac{1}{2}J_{AX}$ and $\nu_{0,B} - \frac{1}{2}J_{BX}$.

The separation between the two effective Larmor frequencies in the two subspectra can easily be different, and so the degree of strong coupling (and hence the intensity patterns) in the two subspectra will be different. All we can measure is the complete spectrum (the sum of the two sub-spectra) but once we know that it is in fact the sum of two AB-type spectra it is usually possible to disentangle these two contributions. Once we have done this, the two sub-spectra can be analysed in exactly the way described above for an AB system. Figure 2.17 illustrates this decomposition.

The form of the X part of the ABX spectrum cannot be deduced from this simple analysis. In general it contains 6 lines, rather than the four which would be expected in the weak coupling limit. The two extra lines are combination lines which become observable when strong coupling is present.

2.7 Exercises

E 2-1

In a proton spectrum the peak from TMS is found to be at 400.135705 MHz. What is the shift, in ppm, of a peak which has a frequency of 400.136305 MHz? Recalculate the shift using the spectrometer frequency, ν_{spec} quoted by the manufacturer as 400.13 MHz rather than ν_{TMS} in the denominator of Eq. 2.2:

$$\delta_{\text{ppm}} = 10^6 \times \frac{\nu - \nu_{\text{TMS}}}{\nu_{\text{spec}}}.$$

Does this make a significant difference to the value of the shift?

E 2-2

Two peaks in a proton spectrum are found at 1.54 and 5.34 ppm. The spectrometer frequency is quoted as 400.13 MHz. What is the separation of these two lines in Hz and in rad s^{-1} ?

E 2-3

Calculate the Larmor frequency (in Hz and in rad s^{-1}) of a carbon-13 resonance with chemical shift 48 ppm when recorded in a spectrometer with a magnetic field strength of 9.4 T. The gyromagnetic ratio of carbon-13 is $+6.7283 \times 10^7 \text{ rad s}^{-1} \text{ T}^{-1}$.

E 2-4

Consider a system of two weakly coupled spins. Let the Larmor frequency of the first spin be -100 Hz and that of the second spin be -200 Hz , and let the coupling between the two spins be -5 Hz . Compute the frequencies of the lines in the normal (single quantum) spectrum.

Make a sketch of the spectrum, roughly to scale, and label each line with the energy levels involved (i.e. 1-2 etc.). Also indicate for each line which spin flips and the spin state of the passive spin. Compare your sketch with Fig. 2.7 and comment on any differences.

E 2-5

For a three spin system, draw up a table similar to that on page 2-11 showing the frequencies of the four lines of the multiplet from spin 2. Then, taking $\nu_{0,2} = -200 \text{ Hz}$, $J_{23} = 4 \text{ Hz}$ and the rest of the parameters as in Fig. 2.11, compute the frequencies of the lines which comprise the spin 2 multiplet. Make a sketch of the multiplet (roughly to scale) and label the lines in the same way as is done in Fig. 2.11. How would these labels change if $J_{23} = -4 \text{ Hz}$?

On an energy level diagram, indicate the four transitions which comprise the spin 2 multiplet, and which four comprise the spin 3 multiplet.

Of course in reality the Larmor frequencies are to be tens or hundreds of MHz, not 100 Hz! However, it makes the numbers easier to handle if we use these unrealistic small values; the principles remain the same, however.

E 2-6

For a three spin system, compute the frequencies of the six zero-quantum transitions and also mark these on an energy level diagram. Do these six transitions fall into natural groups? How would you describe the spectrum?

E 2-7

Calculate the line frequencies and intensities of the spectrum for a system of two spins with the following parameters: $\nu_{0,1} = -10$ Hz, $\nu_{0,2} = -20$ Hz, $J_{12} = 5$ Hz. Make a sketch of the spectrum (roughly to scale) indicating which transition is which and the position of the Larmor frequencies.

E 2-8

The spectrum from a strongly-coupled two spin system showed lines at the following frequencies, in Hz, (intensities are given in brackets): 32.0 (1.3), 39.0 (0.7), 6.0 (0.7), 13.0 (1.3). Determine the values of the coupling constant and the two Larmor frequencies. Show that the values you find are consistent with the observed intensities.

Make sure that you have your calculator set to "radians" when you compute $\sin 2\theta$.

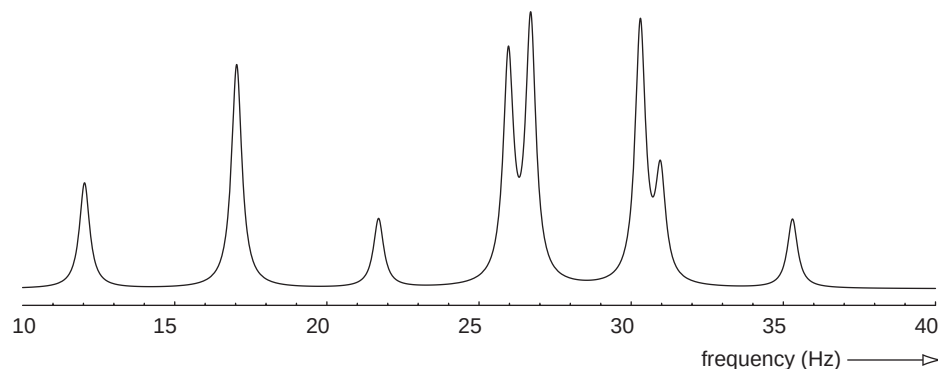
E 2-9

Fig. 2.18 The AB part of an ABX spectrum

Figure 2.18 shows the AB part of an ABX spectrum. Disentangle the two subspectra, mark in the rough positions of the effective Larmor frequencies and hence estimate the size of the AX and BX couplings. Also, give the value of the AB coupling.

3 The vector model

For most kinds of spectroscopy it is sufficient to think about energy levels and selection rules; this is not true for NMR. For example, using this energy level approach we cannot even describe how the most basic pulsed NMR experiment works, let alone the large number of subtle two-dimensional experiments which have been developed. To make any progress in understanding NMR experiments we need some more tools, and the first of these we are going to explore is the *vector model*.

This model has been around as long as NMR itself, and not surprisingly the language and ideas which flow from the model have become the language of NMR to a large extent. In fact, in the strictest sense, the vector model can only be applied to a surprisingly small number of situations. However, the ideas that flow from even this rather restricted area in which the model can be applied are carried over into more sophisticated treatments. It is therefore essential to have a good grasp of the vector model and how to apply it.

3.1 Bulk magnetization

We commented before that the nuclear spin has an interaction with an applied magnetic field, and that it is this which gives rise to the energy levels and ultimately an NMR spectrum. In many ways, it is permissible to think of the nucleus as behaving like a small bar magnet or, to be more precise, a *magnetic moment*. We will not go into the details here, but note that the quantum mechanics tells us that the magnetic moment can be aligned in any direction¹.

In an NMR experiment, we do not observe just one nucleus but a very large number of them (say 10^{20}), so what we need to be concerned with is the net effect of all these nuclei, as this is what we will observe.

If the magnetic moments were all to point in random directions, then the small magnetic field that each generates will cancel one another out and there will be no net effect. However, it turns out that at equilibrium the magnetic moments are not aligned randomly but in such a way that when their contributions are all added up there is a net magnetic field along the direction of the applied field (B_0). This is called the *bulk magnetization* of the sample.

The magnetization can be represented by a vector – called the *magnetization vector* – pointing along the direction of the applied field (z), as shown in Fig. 3.1. From now on we will only be concerned with what happens to this vector.

This would be a good point to comment on the axis system we are going to

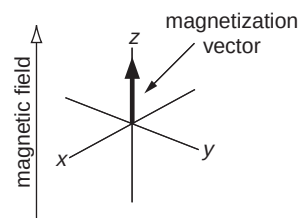


Fig. 3.1 At equilibrium, a sample has a net magnetization along the magnetic field direction (the z axis) which can be represented by a magnetization vector. The axis set in this diagram is a right-handed one, which is what we will use throughout these lectures.

¹It is a common misconception to state that the magnetic moment must either be aligned with or against the magnetic field. In fact, quantum mechanics says no such thing (see Levitt Chapter 9 for a very lucid discussion of this point).

use – it is called a *right-handed set*, and such a set of axes is shown in Fig. 3.1. The name right-handed comes about from the fact that if you imagine grasping the z axis with your right hand, your fingers curl from the x to the y axes.

3.2 Larmor precession

Suppose that we have managed, some how, to tip the magnetization vector away from the z axis, such that it makes an angle β to that axis. We will see later on that such a tilt can be brought about by a radiofrequency pulse. Once tilted away from the z axis we find that the magnetization vector rotates about the direction of the magnetic field sweeping out a cone with a constant angle; see Fig. 3.2. The vector is said to *precess* about the field and this particular motion is called *Larmor precession*.

If the magnetic field strength is B_0 , then the frequency of the Larmor precession is ω_0 (in rad s^{-1})

$$\omega_0 = -\gamma B_0$$

or if we want the frequency in Hz, it is given by

$$\nu_0 = -\frac{1}{2\pi}\gamma B_0$$

where γ is the gyromagnetic ratio. These are of course exactly the same frequencies that we encountered in section 2.3. In words, the frequency at which the magnetization precesses around the B_0 field is exactly the same as the frequency of the line we see from the spectrum on one spin; this is no accident.

As was discussed in section 2.3, the Larmor frequency is a signed quantity and is negative for nuclei with a positive gyromagnetic ratio. This means that for such spins the precession frequency is negative, which is precisely what is shown in Fig. 3.2.

We can sort out positive and negative frequencies in the following way. Imagine grasping the z axis with your right hand, with the thumb pointing along the $+z$ direction. The fingers then curl in the sense of a positive precession. Inspection of Fig. 3.2 will show that the magnetization vector is rotating in the opposite sense to your fingers, and this corresponds to a negative Larmor frequency.

3.3 Detection

The precession of the magnetization vector is what we actually detect in an NMR experiment. All we have to do is to mount a small coil of wire round the sample, with the axis of the coil aligned in the xy -plane; this is illustrated in Fig. 3.3. As the magnetization vector “cuts” the coil a current is induced which we can amplify and then record – this is the so-called *free induction signal* which is detected in a pulse NMR experiment. The whole process is analogous to the way in which electric current can be generated by a magnet rotating inside a coil.

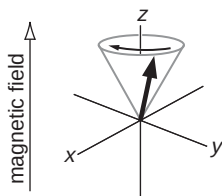


Fig. 3.2 If the magnetization vector is tilted away from the z axis it executes a precessional motion in which the vector sweeps out a cone of constant angle to the magnetic field direction. The direction of precession shown is for a nucleus with a positive gyromagnetic ratio and hence a negative Larmor frequency.

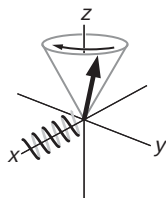


Fig. 3.3 The precessing magnetization will cut a coil wound round the x axis, thereby inducing a current in the coil. This current can be amplified and detected; it is this that forms the free induction signal. For clarity, the coil has only been shown on one side of the x axis.

Essentially, the coil detects the x -component of the magnetization. We can easily work out what this will be. Suppose that the equilibrium magnetization vector is of size M_0 ; if this has been tilted through an angle β towards the x axis, the x -component is $M_0 \sin \beta$; Fig. 3.4 illustrates the geometry.

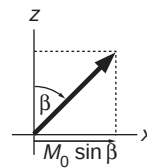


Fig. 3.4 Tilting the magnetization through an angle θ gives an x -component of size $M_0 \sin \beta$.

Although the magnetization vector precesses on a cone, we can visualize what happens to the x - and y -components much more simply by just thinking about the projection onto the xy -plane. This is shown in Fig. 3.5.

At time zero, we will assume that there is only an x -component. After a time τ_1 the vector has rotated through a certain angle, which we will call ϵ_1 . As the vector is rotating at ω_0 radians per second, in time τ_1 the vector has moved through $(\omega_0 \times \tau_1)$ radians; so $\epsilon_1 = \omega_0 \tau_1$. At a later time, say τ_2 , the vector has had longer to precess and the angle ϵ_2 will be $(\omega_0 \tau_2)$. In general, we can see that after time t the angle is $\epsilon = \omega_0 t$.

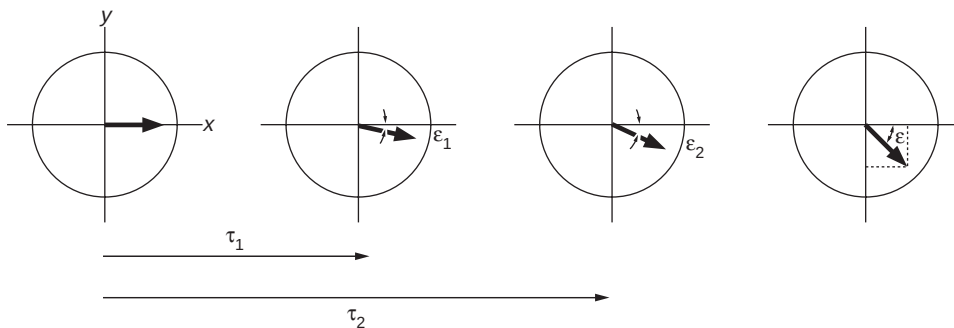


Fig. 3.5 Illustration of the precession of the magnetization vector in the xy -plane. The angle through which the vector has precessed is given by $\omega_0 t$. On the right-hand diagram we see the geometry for working out the x and y components of the vector.

We can now easily work out the x - and y -components of the magnetization using simple geometry; this is illustrated in Fig. 3.5. The x -component is proportional to $\cos \epsilon$ and the y -component is negative (along $-y$) and proportional to $\sin \epsilon$. Recalling that the initial size of the vector is $M_0 \sin \beta$, we can deduce that the x - and y -components, M_x and M_y respectively, are:

$$M_x = M_0 \sin \beta \cos(\omega_0 t)$$

$$M_y = -M_0 \sin \beta \sin(\omega_0 t).$$

Plots of these signals are shown in Fig. 3.6. We see that they are both simple oscillations at the Larmor frequency. Fourier transformation of these signals gives us the familiar spectrum – in this case a single line at ω_0 ; the details of how this works will be covered in a later chapter. We will also see in a later section that in practice we can easily detect *both* the x - and y -components of the magnetization.

3.4 Pulses

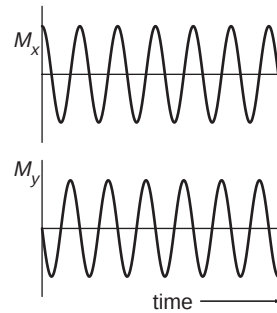


Fig. 3.6 Plots of the x - and y -components of the magnetization predicted using the approach of Fig. 3.5. Fourier transformation of these signals will give rise to the usual spectrum.

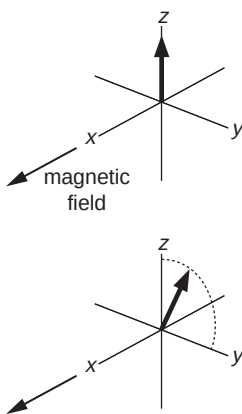


Fig. 3.7 If the magnetic field along the z axis is replaced quickly by one along x , the magnetization will then precess about the x axis and so move towards the transverse plane.

We now turn to the important question as to how we can rotate the magnetization away from its equilibrium position along the z axis. Conceptually it is easy to see what we have to do. All that is required is to (suddenly) replace the magnetic field along the z axis with one in the xy -plane (say along the x axis). The magnetization would then precess about the new magnetic field which would bring the vector down away from the z axis, as illustrated in Fig. 3.7.

Unfortunately it is all but impossible to switch the magnetic field suddenly in this way. Remember that the main magnetic field is supplied by a powerful superconducting magnet, and there is no way that this can be switched off; we will need to find another approach, and it turns out that the key is to use the idea of *resonance*.

The idea is to apply a very small magnetic field along the x axis but one which is oscillating at or near to the Larmor frequency – that is *resonant* with the Larmor frequency. We will show that this small magnetic field is able to rotate the magnetization away from the z axis, even in the presence of the very strong applied field, B_0 .

Conveniently, we can use the same coil to generate this oscillating magnetic field as the one we used to detect the magnetization (Fig. 3.3). All we do is feed some radiofrequency (RF) power to the coil and the resulting oscillating current creates an oscillating magnetic field along the x -direction. The resulting field is called the *radiofrequency* or *RF field*. To understand how this weak RF field can rotate the magnetization we need to introduce the idea of the *rotating frame*.

Rotating frame

When RF power is applied to the coil wound along the x axis the result is a magnetic field which oscillates along the x axis. The magnetic field moves back and forth from $+x$ to $-x$ passing through zero along the way. We will take the frequency of this oscillation to be ω_{RF} (in rad s^{-1}) and the size of the magnetic field to be $2B_1$ (in T); the reason for the 2 will become apparent later. This frequency is also called the *transmitter frequency* for the reason that a radiofrequency transmitter is used to produce the power.

It turns out to be a lot easier to work out what is going on if we replace, in our minds, this *linearly* oscillating field with two counter-rotating fields;

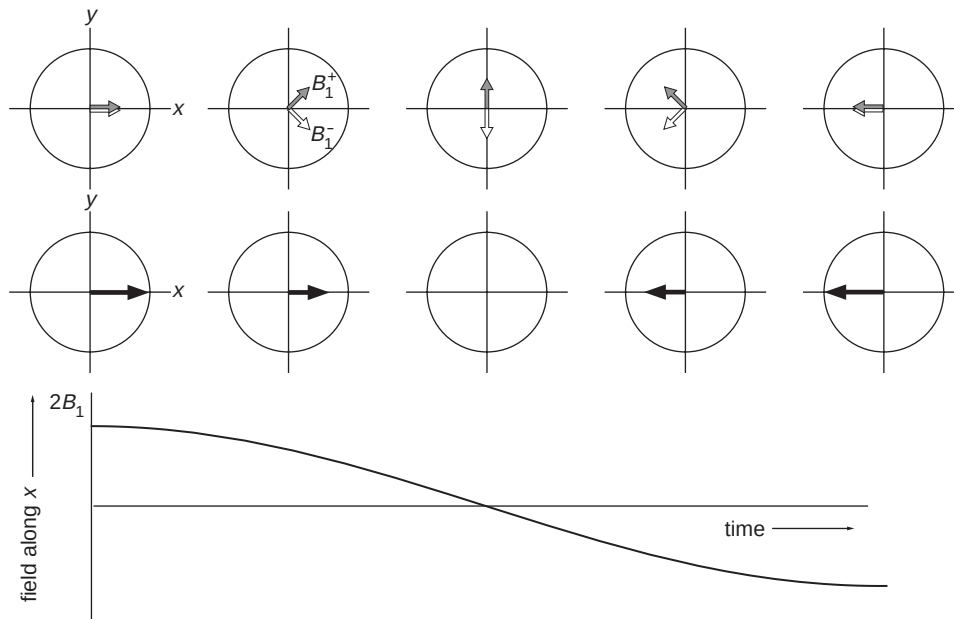


Fig. 3.8 Illustration of how two counter-rotating fields (shown in the upper part of the diagram and marked B_1^+ and B_1^-) add together to give a field which is oscillating along the x axis (shown in the lower part). The graph at the bottom shows how the field along x varies with time.

Fig. 3.8 illustrates the idea. The two counter rotating fields have the same magnitude B_1 . One, denoted B_1^+ , rotates in the positive sense (from x to y) and the other, denoted B_1^- , rotates in the negative sense; both are rotating at the transmitter frequency ω_{RF} .

At time zero, they are both aligned along the x axis and so add up to give a total field of $2B_1$ along the x axis. As time proceeds, the vectors move away from x, in opposite directions. As the two vectors have the same magnitude and are rotating at the same frequency the y-components *always cancel* one another out. However, the x-components shrink towards zero as the angle through which the vectors have rotated approaches $\frac{1}{2}\pi$ radians or 90° . As the angle increases beyond this point the x-component grows once more, but this time along the $-x$ axis, reaching a maximum when the angle of rotation is π . The fields continue to rotate, causing the x-component to drop back to zero and rise again to a value $2B_1$ along the $+x$ axis. Thus we see that the two counter-rotating fields add up to the linearly oscillating one.

Suppose now that we think about a nucleus with a positive gyromagnetic ratio; recall that this means the Larmor frequency is negative so that the sense of precession is from x towards $-y$. This is the same direction as the rotation of B_1^- . It turns out that the other field, which is rotating in the opposite sense to the Larmor precession, has no significant interaction with the magnetization and so from now on we will ignore it.

We now employ a mathematical trick which is to move to a co-ordinate system which, rather than being static (called the *laboratory frame*) is rotating about the z axis in the same direction and at the same rate as B_1^- (i.e. at $-\omega_{\text{RF}}$).

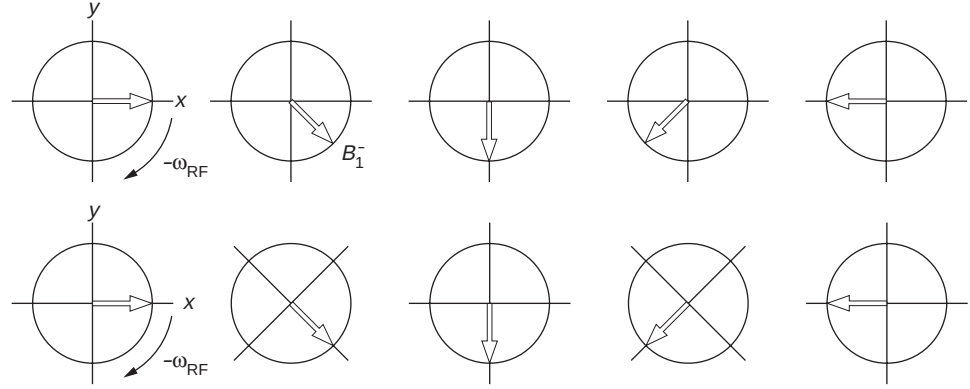


Fig. 3.9 The top row shows a field rotating at $-\omega_{RF}$ when viewed in a fixed axis system. The same field viewed in a set of axes rotating at $-\omega_{RF}$ appears to be static.

In this rotating set of axes, or *rotating frame*, B_1^- appears to be static and directed along the x axis of the rotating frame, as is shown in Fig. 3.9. This is a very nice result as the time dependence has been removed from the problem.

Larmor precession in the rotating frame

We need to consider what happens to the Larmor precession of the magnetization when this is viewed in this rotating frame. In the fixed frame the precession is at ω_0 , but suppose that we choose the rotating frame to be at the same frequency. In the rotating frame the magnetization will appear *not to move* i.e. the apparent Larmor frequency will be zero! It is clear that moving to a rotating frame has an effect on the *apparent* Larmor frequency.

The general case is when the rotating frame is at frequency $\omega_{\text{rot. fram.}}$; in such a frame the Larmor precession will appear to be at $(\omega_0 - \omega_{\text{rot. fram.}})$. This difference frequency is called the offset and is given the symbol Ω :

$$\Omega = \omega_0 - \omega_{\text{rot. fram.}} \quad (3.1)$$

We have used several times the relationship between the magnetic field and the precession frequency:

$$\omega = -\gamma B. \quad (3.2)$$

From this it follows that if the apparent Larmor frequency in the rotating frame is different from that in fixed frame it must also be the case that the *apparent* magnetic field in the rotating frame must be different from the actual applied magnetic field. We can use Eq. 3.2 to compute the apparent magnetic field, given the symbol ΔB , from the apparent Larmor frequency, Ω :

$$\begin{aligned} \Omega &= -\gamma \Delta B \\ \text{hence } \Delta B &= -\frac{\Omega}{\gamma}. \end{aligned}$$

This apparent magnetic field in the rotating frame is usually called the *reduced field*, ΔB .

If we choose the rotating frame to be at the Larmor frequency, the offset Ω will be zero and so too will the reduced field. This is the key to how the very weak RF field can affect the magnetization in the presence of the much stronger B_0 field. In the rotating frame this field along the z axis appears to shrink, and under the right conditions can become small enough that the RF field is dominant.

The effective field

From the discussion so far we can see that when an RF field is being applied there are two magnetic fields in the rotating frame. First, there is the RF field (or B_1 field) of magnitude B_1 ; we will make this field static by choosing the rotating frame frequency to be equal to $-\omega_{\text{RF}}$. Second, there is the reduced field, ΔB , given by $(-\Omega/\gamma)$. Since $\Omega = (\omega_0 - \omega_{\text{rot. fram.}})$ and $\omega_{\text{rot. fram.}} = -\omega_{\text{RF}}$ it follows that the offset is

$$\begin{aligned}\Omega &= \omega_0 - (-\omega_{\text{RF}}) \\ &= \omega_0 + \omega_{\text{RF}}.\end{aligned}$$

This looks rather strange, but recall that ω_0 is negative, so if the transmitter frequency and the Larmor frequency are comparable the offset will be small.

In the rotating frame, the reduced field (which is along z) and the RF or B_1 field (which is along x) add vectorially to give an effective field, B_{eff} as illustrated in Fig. 3.10. The size of this effective field is given by:

$$B_{\text{eff}} = \sqrt{B_1^2 + \Delta B^2}. \quad (3.3)$$

The magnetization precesses about this effective field at frequency ω_{eff} given by

$$\omega_{\text{eff}} = \gamma B_{\text{eff}}$$

just in the same way that the Larmor precession frequency is related to B_0 .

By making the offset small, or zero, the effective field lies close to the xy-plane, and so the magnetization will be rotated from z down to the plane, which is exactly what we want to achieve. The trick is that although B_0 is much larger than B_1 we can affect the magnetization with B_1 by making it oscillate close to the Larmor frequency. This is the phenomena of resonance.

The angle between ΔB and B_{eff} is called the *tilt angle* and is usually given the symbol θ . From Fig. 3.10 we can see that:

$$\sin \theta = \frac{B_1}{B_{\text{eff}}} \quad \cos \theta = \frac{\Delta B}{B_{\text{eff}}} \quad \tan \theta = \frac{B_1}{\Delta B}.$$

All three definitions are equivalent.

The effective field in frequency units

For practical purposes the thing that is important is the precession frequency about the effective field, ω_{eff} . It is therefore convenient to think about the construction of the effective field not in terms of magnetic fields but in terms of the precession frequencies that they cause.

In this discussion we will assume that the gyromagnetic ratio is positive so that the Larmor frequency is negative.

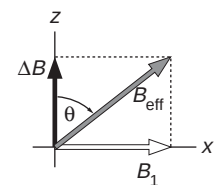


Fig. 3.10 In the rotating frame the effective field B_{eff} is the vector sum of the reduced field ΔB and the B_1 field. The tilt angle, θ , is defined as the angle between ΔB and B_{eff} .

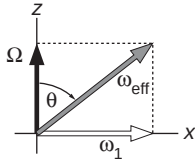


Fig. 3.11 The effective field can be thought of in terms of frequencies instead of the fields used in Fig 3.10.

For each field the precession frequency is proportional to the magnetic field; the constant of proportion is γ , the gyromagnetic ratio. For example, we have already seen that in the rotating frame the apparent Larmor precession frequency, Ω , depends on the reduced field:

$$\Omega = -\gamma \Delta B.$$

We define ω_1 as the precession frequency about the B_1 field (the positive sign is intentional):

$$\omega_1 = \gamma B_1$$

and we already have

$$\omega_{\text{eff}} = \gamma B_{\text{eff}}.$$

Using these definitions in Eq. 3.3 ω_{eff} can be written as

$$\omega_{\text{eff}} = \sqrt{\omega_1^2 + \Omega^2}.$$

Figure 3.10 can be redrawn in terms of frequencies, as shown in Fig. 3.11. Similarly, the tilt angle can be expressed in terms of these frequencies:

$$\sin \theta = \frac{\omega_1}{\omega_{\text{eff}}} \quad \cos \theta = \frac{\Omega}{\omega_{\text{eff}}} \quad \tan \theta = \frac{\omega_1}{\Omega}.$$

On-resonance pulses

The simplest case to deal with is where the transmitter frequency is exactly the same as the Larmor frequency – it is said that the pulse is exactly *on resonance*. Under these circumstances the offset, Ω , is zero and so the reduced field, ΔB , is also zero. Referring to Fig. 3.10 we see that the effective field is therefore the same as the B_1 field and lies along the x axis. For completeness we also note that the tilt angle, θ , of the effective field is $\pi/2$ or 90° .

In this situation the motion of the magnetization vector is very simple. Just as in Fig. 3.7 the magnetization precesses about the field, thereby rotating in the zy -plane. As we have seen above the precession frequency is ω_1 . If the RF field is applied for a time t_p , the angle, β , through which the magnetization has been rotated will be given by

$$\beta = \omega_1 t_p.$$

β is called the *flip angle of the pulse*. By altering the time for which the pulse has been applied we can alter the angle through which the magnetization is rotated.

In many experiments the commonly used flip angles are $\pi/2$ (90°) and π (180°). The motion of the magnetization vector during on-resonance 90° and 180° pulses are shown in Fig. 3.12. The 90° pulse rotates the magnetization from the equilibrium position to the $-y$ axis; this is because the rotation is in the positive sense. Imagine grasping the axis about which the rotation is taking place (the x axis) with your right hand; your fingers then curl in the sense of a positive rotation.

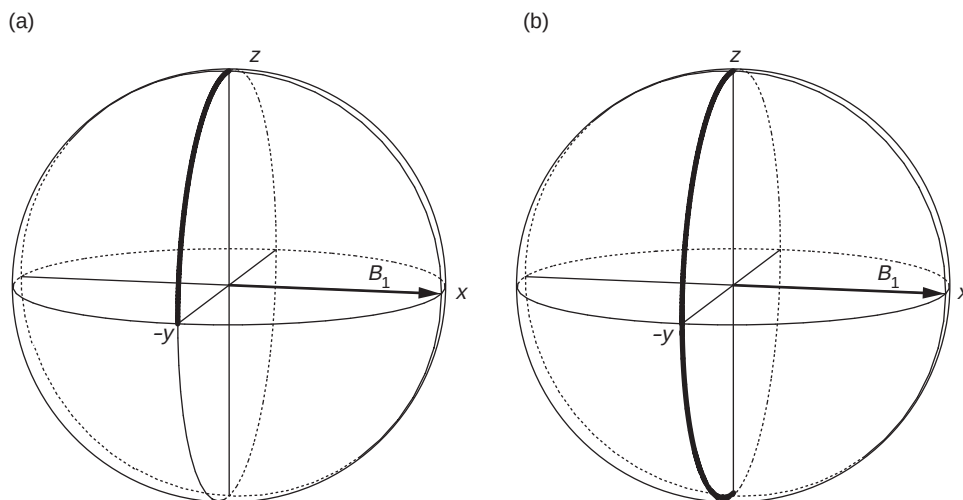


Fig. 3.12 A “grapefruit” diagram in which the thick line shows the motion of a magnetization vector during an on-resonance pulse. The magnetization is assumed to start out along $+z$. In (a) the pulse flip angle is 90° . The effective field lies along the x axis and so the magnetization precesses in the yz -plane. The rotation is in the positive sense about x so the magnetization moves toward the $-y$ axis. In (b) the pulse flip angle is 180° and so the magnetization ends up along $-z$.

If the pulse flip angle is set to 180° the magnetization is taken all the way from $+z$ to $-z$; this is called an *inversion pulse*. In general, for a flip angle β simple geometry tells us that the z - and y -components are

$$M_z = M_0 \cos \beta \quad M_y = -M_0 \sin \beta;$$

this is illustrated in Fig. 3.13.

Hard pulses

In practical NMR spectroscopy we usually have several resonances in the spectrum, each of which has a different Larmor frequency; we cannot therefore be on resonance with all of the lines in the spectrum. However, if we make the RF field strong enough we can effectively achieve this condition.

By “hard” enough we mean that the B_1 field has to be large enough that it is much greater than the size of the reduced field, ΔB . If this condition holds, the effective field lies along B_1 and so the situation is identical to the case of an on-resonance pulse. Thinking in terms of frequencies this condition translates to ω_1 being greater in magnitude than the offset, Ω .

It is often relatively easy to achieve this condition. For example, consider a proton spectrum covering about 10 ppm; if we put the transmitter frequency at about 5 ppm, the maximum offset is 5 ppm, either positive or negative; this is illustrated in Fig. 3.14. If the spectrometer frequency is 500 MHz the maximum offset is $5 \times 500 = 2500$ Hz. A typical spectrometer might have a 90° pulse lasting $12 \mu\text{s}$. From this we can work out the value of ω_1 . We start from

$$\beta = \omega_1 t_p \quad \text{hence} \quad \omega_1 = \frac{\beta}{t_p}.$$

We know that for a 90° pulse $\beta = \pi/2$ and the duration, t_p is 12×10^{-6} s;

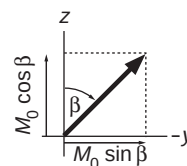


Fig. 3.13 If the pulse flip angle is β we can use simple geometry to work out the y - and z -components.

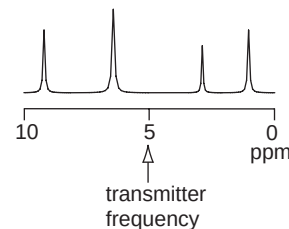


Fig. 3.14 Illustration of the range of offsets one might see in a typical proton spectrum. If the transmitter frequency is placed as shown, the maximum offset of a resonance will be 5 ppm.

therefore

$$\begin{aligned}\omega_1 &= \frac{\pi/2}{12 \times 10^{-6}} \\ &= 1.3 \times 10^5 \text{ rad s}^{-1}.\end{aligned}$$

The maximum offset is 2500 Hz, which is $2\pi \times 2500 = 1.6 \times 10^4 \text{ rad s}^{-1}$. We see that the RF field is about eight times the offset, and so the pulse can be regarded as strong over the whole width of the spectrum.

3.5 Detection in the rotating frame

To work out what is happening during an RF pulse we need to work in the rotating frame, and we have seen that to get this simplification the frequency of the rotating frame must match the transmitter frequency, ω_{RF} . Larmor precession can be viewed just as easily in the laboratory and rotating frames; in the rotating frame the precession is at the offset frequency, Ω .

It turns out that because of the way the spectrometer works the signal that we detect appears to be that in the rotating frame. So, rather than detecting an oscillation at the Larmor frequency, we see an oscillation at the offset, Ω .²

It also turns out that we can detect both the x - and y -components of the magnetization in the rotating frame. From now on we will work exclusively in the rotating frame.

3.6 The basic pulse–acquire experiment

At last we are in a position to describe how the simplest NMR experiment works – the one we use every day to record spectra. The timing diagram – or *pulse sequence* as it is usually known – is shown in Fig. 3.15.

The experiment comes in three periods:

1. The sample is allowed to come to equilibrium.
2. RF power is switched on for long enough to rotate the magnetization through 90° i.e. a 90° pulse is applied.
3. After the RF power is switched off we start to detect the signal which arises from the magnetization as it rotates in the transverse plane.

During period 1 equilibrium magnetization builds up along the z axis. As was described above, the 90° pulse rotates this magnetization onto the $-y$ axis; this takes us to the end of period 2. During period 3 the magnetization precesses in the transverse plane at the offset Ω ; this is illustrated in Fig. 3.16. Some simple geometry, shown in Fig. 3.17, enables us to deduce how the x - and y -magnetizations vary with time. The offset is Ω so after time t the

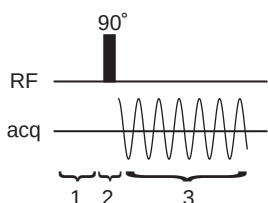


Fig. 3.15 Timing diagram or pulse sequence for the simple pulse–acquire experiment. The line marked “RF” shows the location of the pulses, and the line marked “acq” shows when the signal is recorded or *acquired*.

²Strictly this is only true if we set the receiver reference frequency to be equal to the transmitter frequency; this is almost always the case. More details will be given in Chapter 5.

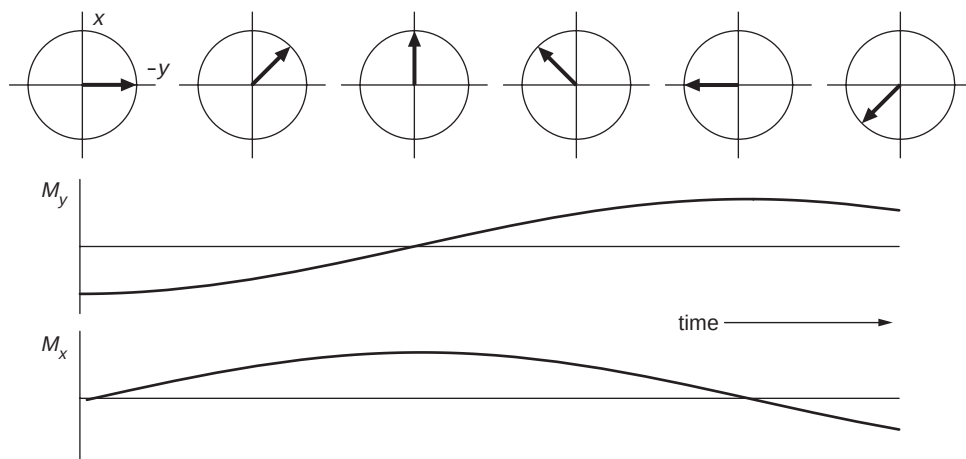


Fig. 3.16 Evolution during the acquisition time (period 3) of the pulse–acquire experiment. The magnetization starts out along $-y$ and evolves at the offset frequency, Ω (here assumed to be positive). The resulting x - and y -magnetizations are shown below.

vector has precessed through an angle $(\Omega \times t)$. The y -component is therefore proportional to $\cos \Omega t$ and the x -component to $\sin \Omega t$. In full the signals are:

$$M_y = -M_0 \cos(\Omega t)$$

$$M_x = M_0 \sin(\Omega t).$$

As we commented on before, Fourier transformation of these signals will give the usual spectrum, with a peak appearing at frequency Ω .

Spectrum with several lines

If the spectrum has more than one line, then to a good approximation we can associate a magnetization vector with each. Usually we wish to observe all the lines at once, so we choose the B_1 field to be strong enough that for the range of offsets that these lines cover all the associated magnetization vectors will be rotated onto the $-y$ axis. During the acquisition time each precesses at its own offset, so the detected signal will be:

$$M_y = -M_{0,1} \cos \Omega_1 t - M_{0,2} \cos \Omega_2 t - M_{0,3} \cos \Omega_3 t \dots$$

where $M_{0,1}$ is the equilibrium magnetization of spin 1, Ω_1 is its offset and so on for the other spins. Fourier transformation of the free induction signal will produce a spectrum with lines at Ω_1 , Ω_2 etc.

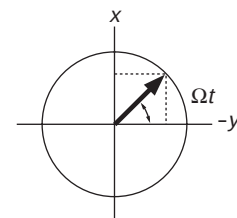


Fig. 3.17 The magnetization starts out along the $-y$ axis and rotates through an angle Ωt during time t .

3.7 Pulse calibration

It is crucial that the pulses we use in NMR experiments have the correct flip angles. For example, to obtain the maximum intensity in the pulse–acquire experiment we must use a 90° pulse, and if we wish to invert magnetization we must use a 180° pulse. Pulse calibration is therefore an important preliminary to any experiment.

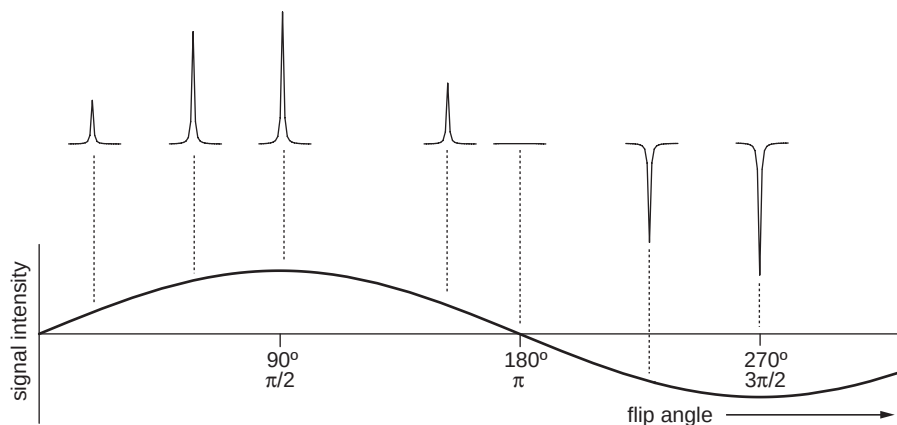


Fig. 3.18 Illustration of how pulse calibration is achieved. The signal intensity varies as $(\sin \beta)$ as shown by the curve at the bottom of the picture. Along the top are the spectra which would be expected for various different flip angles (indicated by the dashed lines). The signal is a maximum for a flip angle of 90° and goes through a null at 180° ; after that, the signal goes negative.

If we imagine an on-resonance or hard pulse we have already determined from Fig. 3.13 that the y-component of magnetization after a pulse of flip angle β is proportional to $\sin \beta$. If we therefore do a pulse-acquire experiment (section 3.6) and vary the flip angle of the pulse, we should see that the intensity of the signal varies as $\sin \beta$. A typical outcome of such an experiment is shown in Fig. 3.18.

The normal practice is to increase the flip angle until a null is found; the flip angle is then 180° . The reason for doing this is that the null is sharper than the maximum. Once the length of a 180° pulse is found, simply halving the time gives a 90° pulse.

Suppose that the 180° pulse was found to be of duration t_{180} . Since the flip angle is given by $\beta = \omega_1 t_p$ we can see that for a 180° pulse in which the flip angle is π

$$\pi = \omega_1 t_{180}$$

$$\text{hence } \omega_1 = \frac{\pi}{t_{180}}.$$

In this way we can determine ω_1 , usually called the RF *field strength* or the B_1 *field strength*.

It is usual to quote the field strength not in rad s^{-1} but in Hz, in which case we need to divide by 2π :

$$(\omega_1/2\pi) = \frac{1}{2t_{180}} \text{ Hz.}$$

For example, let us suppose that we found the null condition at $15.5 \mu\text{s}$; thus

$$\omega_1 = \frac{\pi}{t_{180}} = \frac{\pi}{15.5 \times 10^{-6}} = 2.03 \times 10^5 \text{ rad s}^{-1}.$$

In frequency units the calculation is

$$(\omega_1/2\pi) = \frac{1}{2t_{180}} = \frac{1}{2 \times 15.5 \times 10^{-6}} = 32.3 \text{ kHz.}$$

In normal NMR parlance we would say “the B_1 field is 32.3 kHz”. This is mixing the units up rather strangely, but the meaning is clear once you know what is going on!

3.8 The spin echo

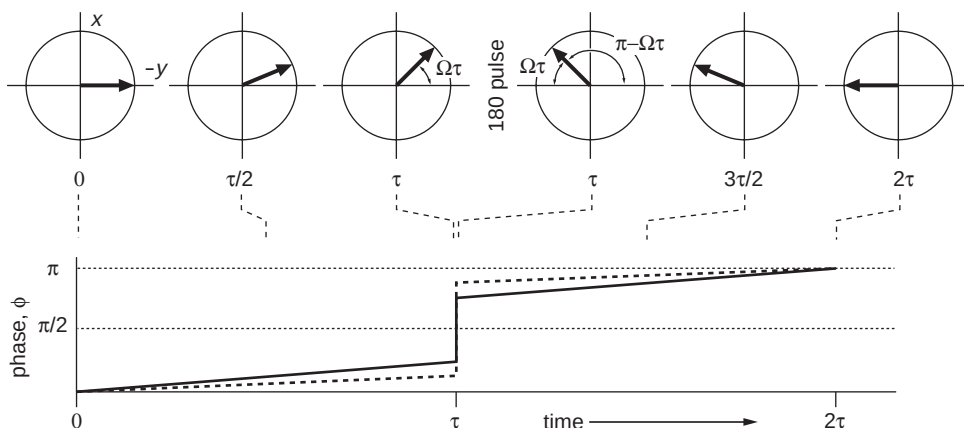


Fig. 3.19 Vector diagrams showing how a spin echo refocuses the evolution of the offset; see text for details. Also shown is a phase evolution diagram for two different offsets (the solid and the dashed line).

We are now able to analyse the most famous pulsed NMR experiment, the *spin echo*, which is a component of a very large number of more complex experiments. The pulse sequence is quite simple, and is shown in Fig. 3.20. The special thing about the spin echo sequence is that at the end of the second τ delay the magnetization ends up along the *same* axis, *regardless* of the values of τ and the offset, Ω .

We describe this outcome by saying that “the offset has been refocused”, meaning that at the end of the sequence it is just as if the offset had been zero and hence there had been no evolution of the magnetization. Figure 3.19 illustrates how the sequence works after the initial 90° pulse has placed the magnetization along the $-y$ axis.

During the first delay τ the vector precesses from $-y$ towards the x axis. The angle through which the vector rotates is simply $(\Omega\tau)$, which we can describe as a *phase*, ϕ . The effect of the 180° pulse is to move the vector to a mirror image position, with the mirror in question being in the xz -plane. So, the vector is now at an angle $(\Omega\tau)$ to the y axis rather than being at $(\Omega\tau)$ to the $-y$ axis.

During the second delay τ the vector continues to evolve; during this time it will rotate through a further angle of $(\Omega\tau)$ and therefore at the end of the second delay the vector will be aligned along the y axis. A few moments thought will reveal that as the angle through which the vector rotates during the first τ delay must be *equal* to that through which it rotates during the second τ delay; the vector will therefore always end up along the y axis *regardless of the offset*, Ω .

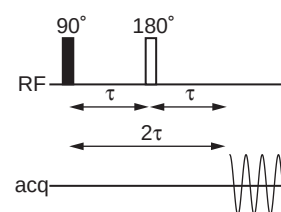


Fig. 3.20 Pulse sequence for the *spin echo* experiment. The 180° pulse (indicated by an open rectangle as opposed to the closed one for a 90° pulse) is in the centre of a delay of duration 2τ , thus separating the sequence into two equal periods, τ . The signal is acquired after the second delay τ , or put another way, when the time from the beginning of the sequence is 2τ . The durations of the pulses are in practice very much shorter than the delays τ but for clarity the length of the pulses has been exaggerated.

The 180° pulse is called a *refocusing pulse* because of the property that the evolution due to the offset during the first delay τ is refocused during the second delay. It is interesting to note that the spin echo sequence gives exactly the same result as the sequence $90^\circ - 180^\circ$ with the delays omitted.

Another way of thinking about the spin echo is to plot a phase evolution diagram; this is done at the bottom of Fig. 3.19. Here we plot the phase, ϕ , as a function of time. During the first τ delay the phase increases linearly with time. The effect of the 180° pulse is to change the phase from $(\Omega\tau)$ to $(\pi - \Omega\tau)$; this is the jump on the diagram at time τ . Further evolution for time τ causes the phase to increase by $(\Omega\tau)$ leading to a final phase at the end of the second τ delay of π . This conclusion is independent of the value of the offset Ω ; the diagram illustrates this by the dashed line which represents the evolution of vector with a smaller offset.

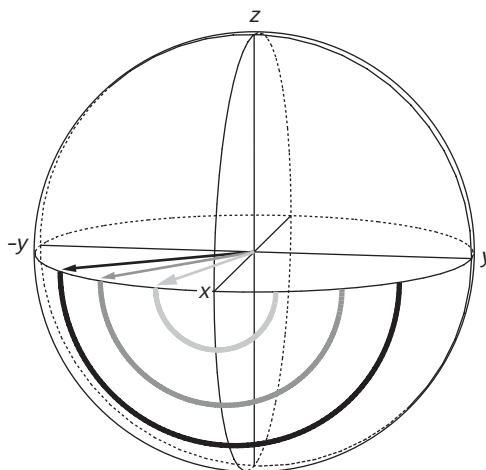


Fig. 3.21 Illustration of the effect of a 180° pulse on three vectors which start out at different angles from the $-y$ axis (coloured in black, grey and light grey). All three are rotated by 180° about the x axis on the trajectories indicated by the thick lines which dip into the southern hemisphere. As a result, the vectors end up in mirror image positions with respect to the xz -plane.

As has already been mentioned, the effect of the 180° pulse is to reflect the vectors in the xz -plane. The way this works is illustrated in Fig. 3.21. The arc through which the vectors are moved is different for each, but all the vectors end up in mirror image positions.

3.9 Pulses of different phases

So far we have assumed that the B_1 field is applied along the x axis; this does not have to be so, and we can just as easily apply it along the y axis, for example. A pulse about y is said to be “phase shifted by 90° ” (we take an x -pulse to have a phase shift of zero); likewise a pulse about $-x$ would be said to be phase shifted by 180° . On modern spectrometers it is possible to produce pulses with arbitrary phase shifts.

If we apply a 90° pulse about the y axis to equilibrium magnetization we find that the vector rotates in the xz -plane such that the magnetization ends up

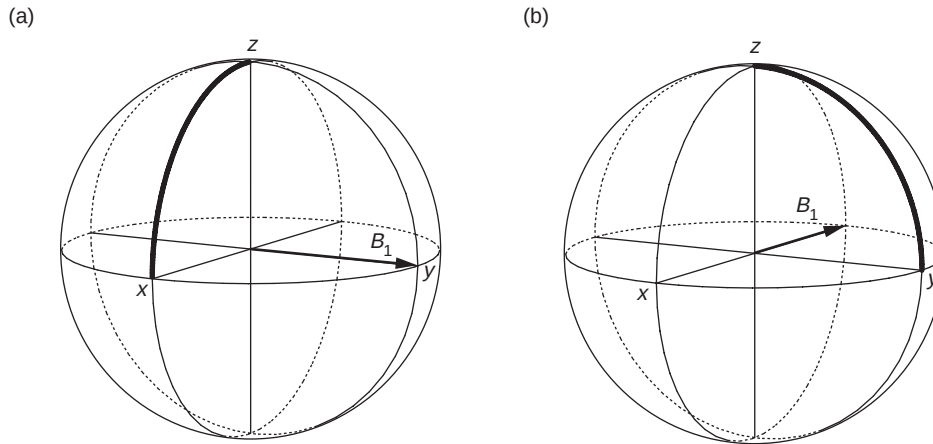


Fig. 3.22 Grapefruit plots showing the effect on equilibrium magnetization of (a) a 90° pulse about the y axis and (b) a 90° pulse about the $-x$ axis. Note the position of the B_1 field in each case.

along x ; this is illustrated in Fig. 3.22. As before, we can determine the effect of such a pulse by thinking of it as causing a positive rotation about the y axis. A 90° pulse about $-x$ causes the magnetization to appear along y , as is also shown in Fig. 3.22.

We have seen that a 180° pulse about the x axis causes the vectors to move to mirror image positions with respect to the xz -plane. In a similar way, a 180° pulse about the y axis causes the vectors to be reflected in the yz -plane.

3.10 Relaxation

We will have a lot more to say about relaxation later on, but at this point we will just note that the magnetization has a tendency to return to its equilibrium position (and size) – a process known as *relaxation*. Recall that the equilibrium situation has magnetization of size M_0 along z and no transverse (x or y) magnetization.

So, if we have created some transverse magnetization (for example by applying a 90° pulse) over time relaxation will cause this magnetization to decay away to zero. The free induction signal, which results from the magnetization precessing in the xy plane will therefore decay away in amplitude. This loss of x - and y -magnetization is called *transverse relaxation*.

Once perturbed, the z -magnetization will try to return to its equilibrium position, and this process is called *longitudinal relaxation*. We can measure the rate of this process using the *inversion recovery* experiment whose pulse sequence is shown in Fig. 3.23. The 180° pulse rotates the equilibrium magnetization to the $-z$ axis. Suppose that the delay τ is very short so that at the end of this delay the magnetization has not changed. Now the 90° pulse will rotate the magnetization onto the $+y$ axis; note that this is in contrast to the case where the magnetization starts out along $+z$ and it is rotated onto $-y$. If

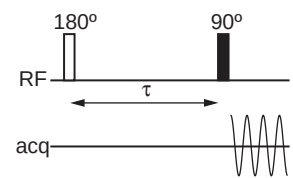


Fig. 3.23 The pulse sequence for the inversion recovery experiment used to measure longitudinal relaxation.

this gives a positive line in the spectrum, then having the magnetization along $+y$ will give a negative line. So, what we see for short values of τ is a negative line.

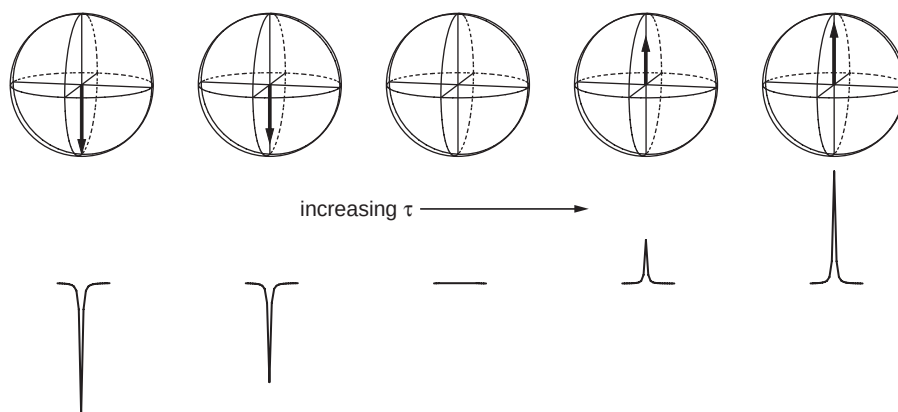


Fig. 3.24 Visualization of the outcome of an inversion recovery experiment. The size and sign of the z -magnetization is reflected in the spectra (shown underneath). By analysing the peak heights as a function of the delay τ it is possible to find the rate of recovery of the z -magnetization.

As τ gets longer more relaxation takes place and the magnetization shrinks towards zero; this result is a negative line in the spectrum, but one whose size is decreasing. Eventually the magnetization goes through zero and then starts to increase along $+z$ – this gives a positive line in the spectrum. Thus, by recording spectra with different values of the delay τ we can map out the recovery of the z -magnetization from the intensity of the observed lines. The whole process is visualized in Fig. 3.24.

3.11 Off-resonance effects and soft pulses

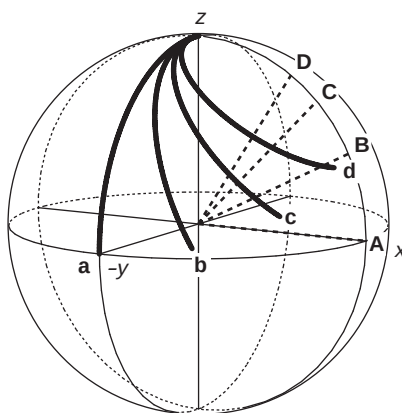


Fig. 3.25 Grapefruit diagram showing the path followed during a pulse for various different resonance offsets. Path **a** is for the on-resonance case; the effective field lies along x and is indicated by the dashed line **A**. Path **b** is for the case where the offset is half the RF field strength; the effective field is marked **B**. Paths **c** and **d** are for offsets equal to and 1.5 times the RF field strength, respectively. The effective field directions are labelled **C** and **D**.

So far we have only dealt with the case where the pulse is either on resonance or where the RF field strength is large compared to the offset (a hard

pulse) which is in effect the same situation. We now turn to the case where the offset is comparable to the RF field strength. The consequences of this are sometimes a problem to us, but they can also be turned to our advantage for *selective excitation*.

As the offset becomes comparable to the RF field, the effective field begins to move up from the x axis towards the z axis. As a consequence, rather than the magnetization moving in the yz plane from z to $-y$, the magnetization starts to follow a more complex curved path; this is illustrated in Fig. 3.25. The further off resonance we go, the further the vector ends up from the xy plane. Also, there is a significant component of magnetization generated along the x direction, something which does not occur in the on-resonance case.

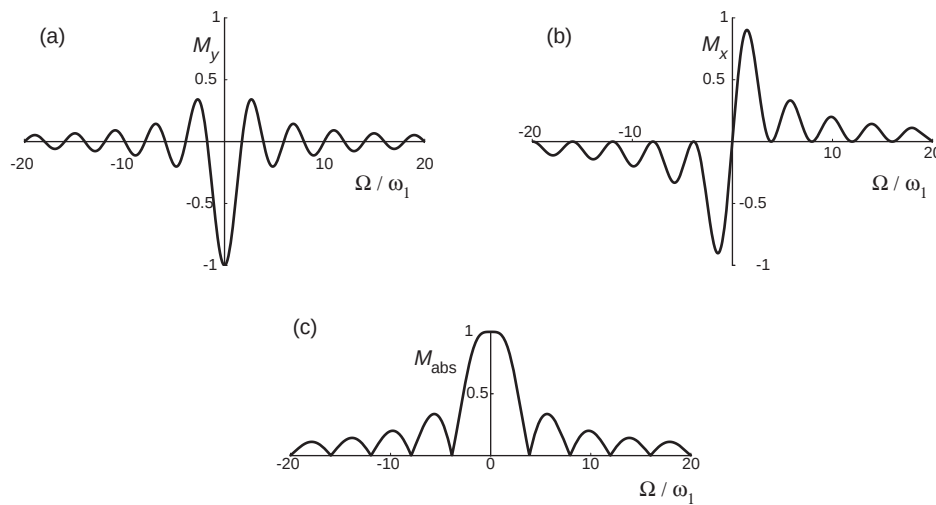


Fig. 3.26 Plots of the magnetization produced by a pulse as a function of the offset. The pulse length has been adjusted so that on resonance the flip angle is 90° . The horizontal axes of the plots is the offset expressed as a ratio of the RF field strength, ω_1 ; the equilibrium magnetization has been assumed to be of size 1.

We can see more clearly what is going on if we plot the x and y magnetization as a function of the offset; these are shown in Fig. 3.26. In (a) we see the y -magnetization and, as expected for a 90° pulse, on resonance the equilibrium magnetization ends up entirely along $-y$. However, as the offset increases the amount of y -magnetization generally decreases but imposed on this overall decrease there is an oscillation; at some offsets the magnetization is zero and at others it is positive. The plot of the x -magnetization, (b), shows a similar story with the magnetization generally falling off as the offset increases, but again with a strong oscillation.

Plot (c) is of the magnitude of the magnetization, which is given by

$$M_{abs} = \sqrt{M_x^2 + M_y^2}.$$

This gives the total transverse magnetization in any direction; it is, of course, always positive. We see from this plot the characteristic nulls and subsidiary maxima as the offset increases.

What plot (c) tells us is that although a pulse can excite magnetization over a wide range of offsets, the region over which it does so efficiently is really rather small. If we want at least 90% of the full intensity the offset must be less than about 1.6 times the RF field strength.

Excitation of a range of shifts

There are some immediate practical consequences of this observation. Suppose that we are trying to record the full range of carbon-13 shifts (200 ppm) on an spectrometer whose magnetic field gives a proton Larmor frequency of 800 MHz and hence a carbon-13 Larmor frequency of 200 MHz. If we place the transmitter frequency at 100 ppm, the maximum offset that a peak can have is 100 ppm which, at this Larmor frequency, translates to 20 kHz. According to our criterion above, if we accept a reduction to 90% of the full intensity at the edges of the spectrum we would need an RF field strength of $20/1.6 \approx 12.5$ kHz. This would correspond to a 90° pulse width of $20 \mu\text{s}$. If the manufacturer failed to provide sufficient power to produce this pulse width we can see that the excitation of the spectrum will fall below our (arbitrary) 90% mark.

We will see in a later section that the presence of a mixture of x - and y -magnetization leads to phase errors in the spectrum which can be difficult to correct.

Selective excitation

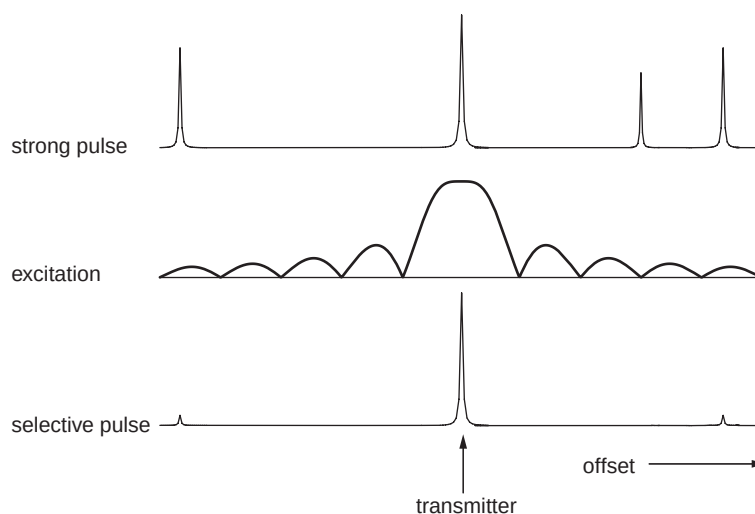


Fig. 3.27 Visualization of the use of selective excitation to excite just one line in the spectrum. At the top is shown the spectrum that would be excited using a hard pulse. If the transmitter is placed on resonance with one line and the strength of the RF field reduced then the pattern of excitation we expect is as shown in the middle. As a result, the peaks at non-zero offsets are attenuated and the spectrum which is excited will be as shown at the bottom.

Sometimes we want to excite just a portion of the spectrum, for example just the lines of a single multiplet. We can achieve this by putting the transmitter in the centre of the region we wish to excite and then reducing the RF field strength until the degree of excitation of the rest of the spectrum is es-

entially negligible. Of course, reducing the RF field strength lengthens the duration of a 90° pulse. The whole process is visualized in Fig. 3.27.

Such pulses which are designed to affect only part of the spectrum are called *selective pulses* or *soft pulses* (as opposed to non-selective or hard pulses). The value to which we need to reduce the RF field depends on the separation of the peak we want to excite from those we do not want to excite. The closer in the unwanted peaks are the weaker the RF field must become and hence the longer the 90° pulse. In the end a balance has to be made between making the pulse too long (and hence losing signal due to relaxation) and allowing a small amount of excitation of the unwanted signals.

Figure 3.27 does not portray one problem with this approach, which is that for peaks away from the transmitter a mixture of x - and y -magnetization is generated (as shown in Fig. 3.26). This is described as a phase error, more of which in a later section. The second problem that the figure does show is that the excitation only falls off rather slowly and “bounces” through a series of maximum and nulls; these are sometimes called “wiggles”. We might be lucky and have an unwanted peak fall on a null, or unlucky and have an unwanted peak fall on a maximum.

Much effort has been put into getting round both of these problems. The key feature of all of the approaches is to “shape” the envelope of the RF pulses i.e. not just switch it on and off abruptly, but with a smooth variation. Such pulses are called *shaped pulses*. The simplest of these are basically bell-shaped (like a gaussian function, for example). These suppress the “wiggles” at large offsets and give just a smooth decay; they do not, however, improve the phase properties. To attack this part of the problem requires an altogether more sophisticated approach.

Selective inversion

Sometimes we want to invert the magnetization associated with just one resonance while leaving all the others in the spectrum unaffected; such a pulse would be called a *selective inversion* pulse. Just as for selective excitation, all we need to do is to place the transmitter on the line we wish to invert and reduce the RF field until the other resonances in the spectrum are not affected significantly. Of course we need to adjust the pulse duration so that the on-resonance flip angle is 180° .

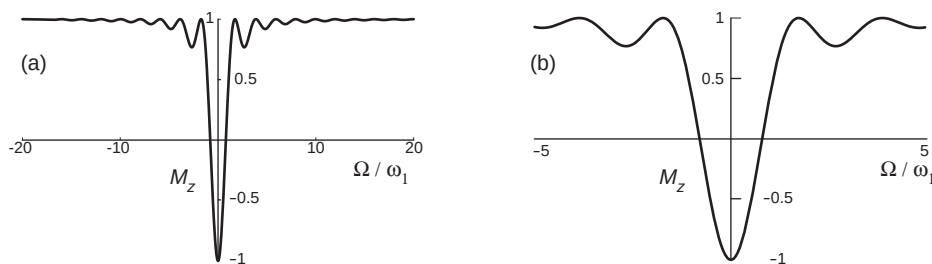


Fig. 3.28 Plots of the z -magnetization produced by a pulse as a function of the offset; the flip angle on-resonance has been set to 180° . Plot (b) covers a narrower range of offsets than plot (a). These plots should be compared with those in Fig. 3.26.

Figure 3.28 shows the z-magnetization generated as a function of offset for a pulse. We see that the range over which there is significant inversion is rather small, and that the oscillations are smaller in amplitude than for the excitation pulse.

This observation has two consequences: one “good” and one “bad”. The good consequence is that a selective 180° pulse is, for a given field strength, more selective than a corresponding 90° pulse. In particular, the weaker “bouncing” sidelobes are a useful feature. Do not forget, though, that the 180° pulse is longer, so some of the improvement may be illusory!

The bad consequence is that when it comes to hard pulses the range of offsets over which there is anything like complete inversion is much more limited than the range of offsets over which there is significant excitation. This can be seen clearly by comparing Fig. 3.28 with Figure 3.26. Thus, 180° pulses are often the source of problems in spectra with large offset ranges.

3.12 Exercises**E 3-1**

A spectrometer operates with a Larmor frequency of 600 MHz for protons. For a particular set up the RF field strength, $\omega_1/(2\pi)$ has been determined to be 25 kHz. Suppose that the transmitter is placed at 5 ppm; what is the offset (in Hz) of a peak at 10 ppm? Compute the tilt angle, θ , of a spin with this offset.

For the normal range of proton shifts (0 – 10 ppm), is this 25 kHz field strong enough to give what could be classed as hard pulses?

E 3-2

In an experiment to determine the pulse length an operator observed a positive signal for pulse widths of 5 and 10 μs ; as the pulse was lengthened further the intensity decreased going through a null at 20.5 μs and then turning negative.

Explain what is happening in this experiment and use the data to determine the RF field strength in Hz and the length of a 90° pulse.

A further null in the signal was seen at 41.0 μs ; to what do you attribute this?

E 3-3

Use vector diagrams to describe what happens during a spin echo sequence in which the 180° pulse (the refocusing pulse) is applied about the y axis (assume that the initial 90° pulse is still about the x-axis). Also, draw a phase evolution diagram appropriate for this pulse sequence.

In what way is the outcome different from the case where the refocusing pulse is applied about the x axis?

What would the effect of applying the refocusing pulse about the $-x$ axis be?

E 3-4

The gyromagnetic ratio of phosphorus-31 is $1.08 \times 10^8 \text{ rad s}^{-1} \text{ T}^{-1}$. This nucleus shows a wide range of shifts, covering some 700 ppm.

Estimate the minimum 90° pulse width you would need to excite peaks in this complete range to within 90% of their theoretical maximum for a spectrometer with a B_0 field strength of 9.4 T.

E 3-5

A spectrometer operates at a Larmor frequency of 400 MHz for protons and hence 100 MHz for carbon-13. Suppose that a 90° pulse of length 10 μs is applied to the protons. Does this have a significant effect of the carbon-13 nuclei? Explain your answer carefully.

E 3-6

Referring to the plots of Fig. 3.26 we see that there are some offsets at which the transverse magnetization goes to zero. Recall that the magnetization is rotating about the *effective field*, ω_{eff} ; it follows that these nulls in the excitation come about when the magnetization executes complete 360° rotations about the effective field. In such a rotation the magnetization is returned to the z axis. Make a sketch of a “grapefruit” showing this.

The effective field is given by

$$\omega_{\text{eff}} = \sqrt{\omega_1^2 + \Omega^2}.$$

Suppose that we express the offset as a multiple κ of the RF field strength:

$$\Omega = \kappa\omega_1.$$

Show that with this values of Ω the effective field is given by:

$$\omega_{\text{eff}} = \omega_1 \sqrt{1 + \kappa^2}.$$

(The reason for doing this is to reduce the number of variables.)

Let us assume that on-resonance the pulse flip angle is $\pi/2$, so the duration of the pulse, τ_p , is give from

$$\omega_1\tau_p = \pi/2 \quad \text{thus} \quad \tau_p = \frac{\pi}{2\omega_1}.$$

The angle of rotation about the effective field for a pulse of duration τ_p is $(\omega_{\text{eff}}\tau_p)$. Show that for the effective field given above this angle, β_{eff} is given by

$$\beta_{\text{eff}} = \frac{\pi}{2} \sqrt{1 + \kappa^2}.$$

The null in the excitation will occur when β_{eff} is 2π i.e. a complete rotation. Show that this occurs when $\kappa = \sqrt{15}$ i.e. when $(\Omega/\omega_1) = \sqrt{15}$. Does this agree with Fig. 3.26?

Predict other values of κ at which there will be nulls in the excitation.

E 3-7

When calibrating a pulse by looking for the null produced by a 180° rotation, why is it important to choose a line which is close to the transmitter frequency (i.e. one with a small offset)?

E 3-8

Use vector diagrams to predict the outcome of the sequence:

$$90^\circ - \text{delay } \tau - 90^\circ$$

applied to equilibrium magnetization; both pulses are about the x axis. In your answer, explain how the x , y and z magnetizations depend on the delay τ and the offset Ω .

E 3-9

Consider the spin echo sequence to which a 90° pulse has been added at the end:

$$90^\circ(x) - \text{delay } \tau - 180^\circ(x) - \text{delay } \tau - 90^\circ(\phi).$$

The axis about which the pulse is applied is given in brackets after the flip angle. Explain in what way the outcome is different depending on whether the phase ϕ of the pulse is chosen to be x , y , $-x$ or $-y$.

E 3-10

The so-called “1- $\bar{1}$ ” sequence is:

$$90^\circ(x) - \text{delay } \tau - 90^\circ(-x)$$

For a peak which is on resonance the sequence does not excite any observable magnetization. However, for a peak with an offset such that $\Omega\tau = \pi/2$ the sequence results in all of the equilibrium magnetization appearing along the x axis. Further, if the delay is such that $\Omega\tau = \pi$ no transverse magnetization is excited.

Explain these observations and make a sketch graph of the amount of transverse magnetization generated as a function of the offset for a fixed delay τ .

The sequence has been used for suppressing strong solvent signals which might otherwise overwhelm the spectrum. The solvent is placed on resonance, and so is not excited; τ is chosen so that the peaks of interest are excited. How does one go about choosing the value for τ ?

E 3-11

The so-called “1-1” sequence is:

$$90^\circ(x) - \text{delay } \tau - 90^\circ(y).$$

Describe the excitation that this sequence produces as a function of offset. How it could be used for observing spectra in the presence of strong solvent signals?

E 3-12

If there are two peaks in the spectrum, we can work out the effect of a pulse sequences by treating the two lines separately. *There is a separate magnetization vector for each line.*

(a) Suppose that a spectrum has two lines, A and B. Suppose also that line A is on resonance with the transmitter and that the offset of line B is 100 Hz.

Starting from equilibrium, we apply the following pulse sequence:

$$90^\circ(x) - \text{delay } \tau - 90^\circ(x)$$

Using the vector model, work out what happens to the magnetization from line A.

Assuming that the delay τ is set to 5 ms ($1 \text{ ms} = 10^{-3} \text{ s}$), work out what happens to the magnetization vector from line B.

(b) Suppose now that we move the transmitter so that it is exactly between the two lines. The offset of line A is now +50 Hz and of line B is -50 Hz.

Starting from equilibrium, we apply the following pulse sequence:

$$90^\circ(x) - \text{delay } \tau$$

Assuming that the delay τ is set to 5 ms, work out what happens to the two magnetization vectors.

If we record a FID after the delay τ and then Fourier transform it, what will the spectrum look like?

Clear diagrams are essential!

4 Fourier transformation and data processing

In the previous chapter we have seen how the precessing magnetization can be detected to give a signal which oscillates at the Larmor frequency – the free induction signal. We also commented that this signal will eventually decay away due to the action of relaxation; the signal is therefore often called the *free induction decay* or FID. The question is how do we turn this signal, which depends on *time*, into the a spectrum, in which the horizontal axis is *frequency*.

This conversion is made using a mathematical process known as *Fourier transformation*. This process takes the *time domain* function (the FID) and converts it into a *frequency domain* function (the spectrum); this is shown in Fig. 4.1. In this chapter we will start out by exploring some features of the spectrum, such as phase and lineshapes, which are closely associated with the Fourier transform and then go on to explore some useful manipulations of NMR data such as sensitivity and resolution enhancement.

4.1 The FID

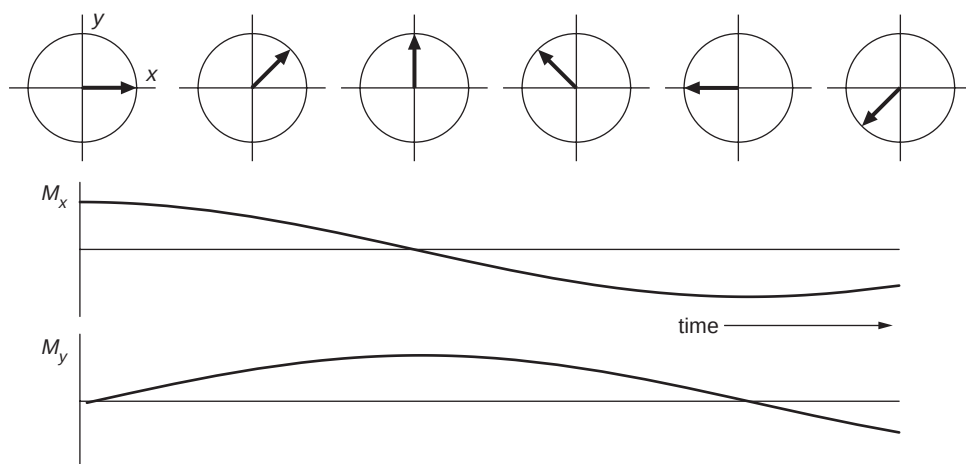


Fig. 4.2 Evolution of the magnetization over time; the offset is assumed to be positive and the magnetization starts out along the x axis.

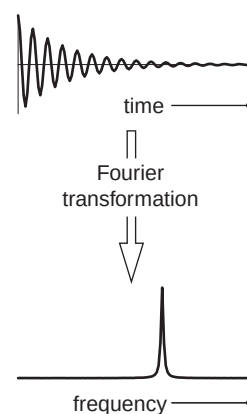


Fig. 4.1 Fourier transformation is the mathematical process which takes us from a function of time (the time domain) – such as a FID – to a function of frequency – the spectrum.

In section 3.6 we saw that the x and y components of the free induction signal could be computed by thinking about the evolution of the magnetization during the acquisition time. In that discussion we assumed that the magnetization started out along the $-y$ axis as this is where it would be rotated to by a 90° pulse. For the purposes of this chapter we are going to assume that the

magnetization starts out along x ; we will see later that this choice of starting position is essentially arbitrary.

If the magnetization does indeed start along x then Fig. 3.16 needs to be redrawn, as is shown in Fig. 4.2. From this we can easily see that the x and y components of the magnetization are:

$$M_x = M_0 \cos \Omega t$$

$$M_y = M_0 \sin \Omega t.$$

The signal that we detect is proportional to these magnetizations. The constant of proportion depends on all sorts of instrumental factors which need not concern us here; we will simply write the detected x and y signals, $S_x(t)$ and $S_y(t)$ as

$$S_x(t) = S_0 \cos \Omega t \quad \text{and} \quad S_y(t) = S_0 \sin \Omega t$$

where S_0 gives the overall size of the signal and we have reminded ourselves that the signal is a function of time by writing it as $S_x(t)$ etc.

It is convenient to think of this signal as arising from a vector of length S_0 rotating at frequency Ω ; the x and y components of the vector give S_x and S_y , as is illustrated in Fig. 4.3.

As a consequence of the way the Fourier transform works, it is also convenient to regard $S_x(t)$ and $S_y(t)$ as the real and imaginary parts of a complex signal $S(t)$:

$$\begin{aligned} S(t) &= S_x(t) + iS_y(t) \\ &= S_0 \cos \Omega t + iS_0 \sin \Omega t \\ &= S_0 \exp(i\Omega t). \end{aligned}$$

We need not concern ourselves too much with the mathematical details here, but just note that the time-domain signal is complex, with the real and imaginary parts corresponding to the x and y components of the signal.

We mentioned at the start of this section that the transverse magnetization decays over time, and this is most simply represented by an exponential decay with a time constant T_2 . The signal then becomes

$$S(t) = S_0 \exp(i\Omega t) \exp\left(\frac{-t}{T_2}\right). \quad (4.1)$$

A typical example is illustrated in Fig. 4.4. Another way of writing this is to define a (first order) rate constant $R_2 = 1/T_2$ and so $S(t)$ becomes

$$S(t) = S_0 \exp(i\Omega t) \exp(-R_2 t). \quad (4.2)$$

The shorter the time T_2 (or the larger the rate constant R_2) the more rapidly the signal decays.

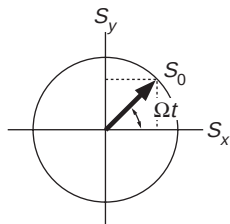


Fig. 4.3 The x and y components of the signal can be thought of as arising from the rotation of a vector S_0 at frequency Ω .

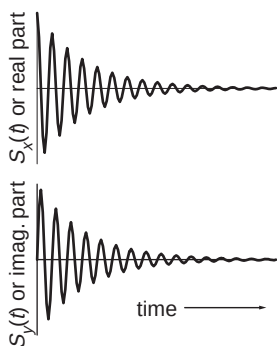


Fig. 4.4 Illustration of a typical FID, showing the real and imaginary parts of the signal; both decay over time.

4.2 Fourier transformation

Fourier transformation of a signal such as that given in Eq. 4.1 gives the frequency domain signal which we know as the spectrum. Like the time domain signal the frequency domain signal has a real and an imaginary part. The real part of the spectrum shows what we call an *absorption mode* line, in fact in the case of the exponentially decaying signal of Eq. 4.1 the line has a shape known as a *Lorentzian*, or to be precise the *absorption mode Lorentzian*. The imaginary part of the spectrum gives a lineshape known as the *dispersion mode Lorentzian*. Both lineshapes are illustrated in Fig. 4.5.

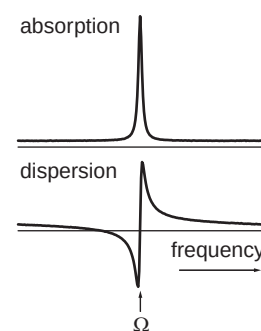


Fig. 4.5 Illustration of the absorption and dispersion mode Lorentzian lineshapes. Whereas the absorption lineshape is always positive, the dispersion lineshape has positive and negative parts; it also extends further.

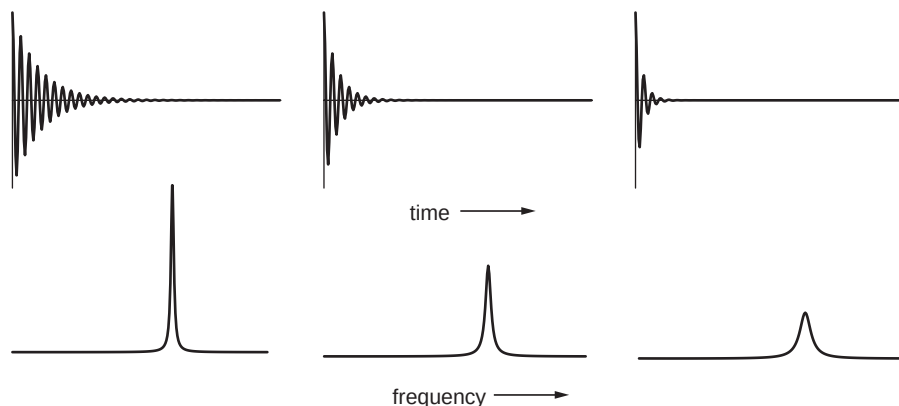


Fig. 4.6 Illustration of the fact that the more rapidly the FID decays the broader the line in the corresponding spectrum. A series of FIDs are shown at the top of the figure and below are the corresponding spectra, all plotted on the same vertical scale. The integral of the peaks remains constant, so as they get broader the peak height decreases.

This absorption lineshape has a width at half of its maximum height of $1/(\pi T_2)$ Hz or (R/π) Hz. This means that the faster the decay of the FID the broader the line becomes. However, the area under the line – that is the integral – remains constant so as it gets broader so the peak height reduces; these points are illustrated in Fig. 4.6.

If the size of the time domain signal increases, for example by increasing S_0 the height of the peak increases in direct proportion. These observations lead to the very important consequence that by integrating the lines in the spectrum we can determine the relative number of protons (typically) which contribute to each.

The dispersion line shape is not one that we would choose to use. Not only is it broader than the absorption mode, but it also has positive and negative parts. In a complex spectrum these might cancel one another out, leading to a great deal of confusion. If you are familiar with ESR spectra you might recognize the dispersion mode lineshape as looking like the *derivative* lineshape which is traditionally used to plot ESR spectra. Although these two lineshapes do look roughly the same, they are not in fact related to one another.

Positive and negative frequencies

As we discussed in section 3.5, the evolution we observe is at frequency Ω i.e. the apparent Larmor frequency in the rotating frame. This offset can be

positive or negative and, as we will see later, it turns out to be possible to determine the sign of the frequency. So, in our spectrum we have positive and negative frequencies, and it is usual to plot these with zero in the middle.

Several lines

What happens if we have more than one line in the spectrum? In this case, as we saw in section 3.5, the FID will be the sum of contributions from each line. For example, if there are three lines $S(t)$ will be:

$$S(t) = S_{0,1} \exp(i\Omega_1 t) \exp\left(\frac{-t}{T_2^{(1)}}\right) + S_{0,2} \exp(i\Omega_2 t) \exp\left(\frac{-t}{T_2^{(2)}}\right) + S_{0,3} \exp(i\Omega_3 t) \exp\left(\frac{-t}{T_2^{(3)}}\right).$$

where we have allowed each line to have a separate intensity, $S_{0,i}$, frequency, Ω_i , and relaxation time constant, $T_2^{(i)}$.

The Fourier transform is a *linear* process which means that if the time domain is a sum of functions the frequency domain will be a sum of Fourier transforms of those functions. So, as Fourier transformation of each of the terms in $S(t)$ gives a line of appropriate width and frequency, the Fourier transformation of $S(t)$ will be the sum of these lines – which is the complete spectrum, just as we require it.

4.3 Phase

So far we have assumed that at time zero (i.e. at the start of the FID) $S_x(t)$ is a maximum and $S_y(t)$ is zero. However, in general this need not be the case – it might just as well be the other way round or anywhere in between. We describe this general situation by saying that the signal is *phase shifted* or that it has a *phase error*. The situation is portrayed in Fig. 4.7.

In Fig. 4.7 (a) we see the situation we had before, with the signal starting out along x and precessing towards y . The real part of the FID (corresponding to S_x) is a damped cosine wave and the imaginary part (corresponding to S_y) is a damped sine wave. Fourier transformation gives a spectrum in which the real part contains the absorption mode lineshape and the imaginary part the dispersion mode.

In (b) we see the effect of a phase shift, ϕ , of 45° . S_y now starts out at a finite value, rather than at zero. As a result neither the real nor the imaginary part of the spectrum has the absorption mode lineshape; both are a mixture of absorption and dispersion.

In (c) the phase shift is 90° . Now it is S_y which takes the form of a damped cosine wave, whereas S_x is a sine wave. The Fourier transform gives a spectrum in which the absorption mode signal now appears in the imaginary part. Finally in (d) the phase shift is 180° and this gives a negative absorption mode signal in the real part of the spectrum.

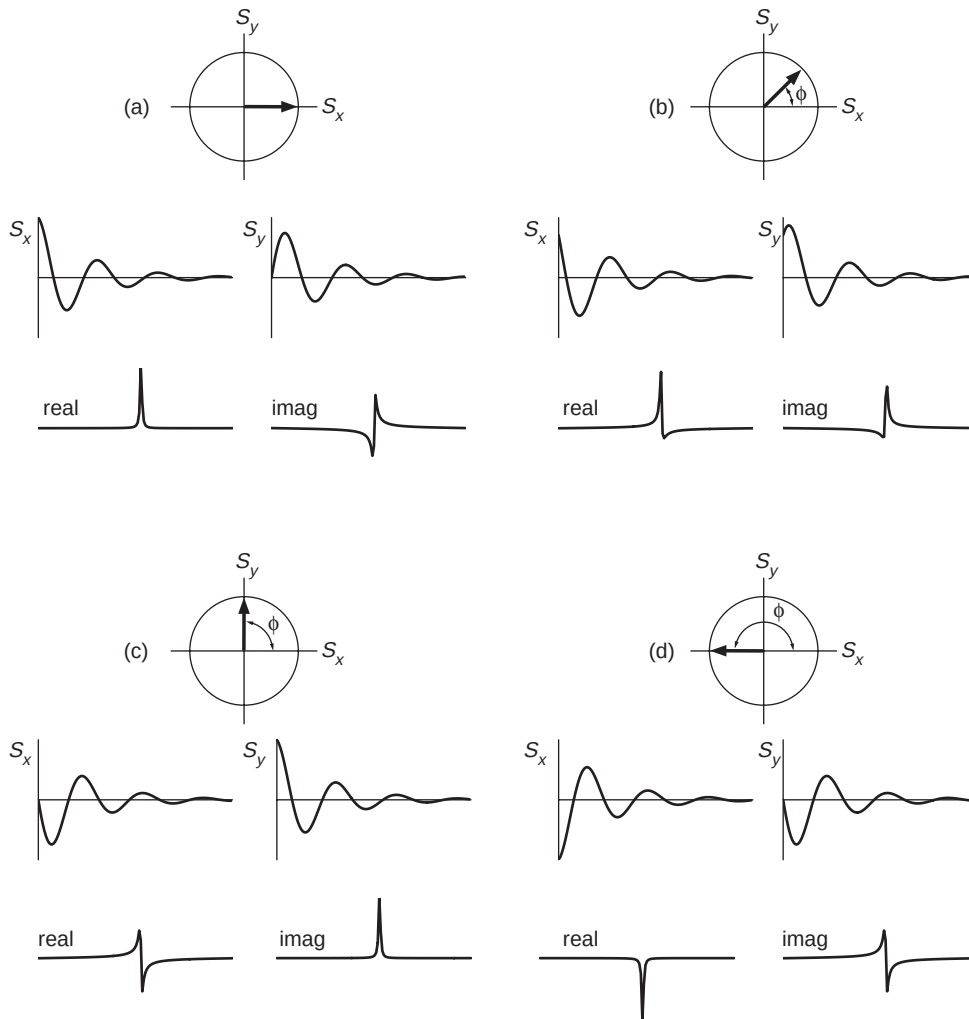


Fig. 4.7 Illustration of the effect of a phase shift of the time domain signal on the spectrum. In (a) the signal starts out along x and so the spectrum is the absorption mode in the real part and the dispersion mode in the imaginary part. In (b) there is a phase shift, ϕ , of 45° ; the real and imaginary parts of the spectrum are now mixtures of absorption and dispersion. In (c) the phase shift is 90° ; now the absorption mode appears in the imaginary part of the spectrum. Finally in (d) the phase shift is 180° giving a negative absorption line in the real part of the spectrum. The vector diagrams illustrate the position of the signal at time zero.

What we see is that in general the appearance of the spectrum depends on the position of the signal at time zero, that is on the phase of the signal at time zero. Mathematically, inclusion of this phase shift means that the (complex) signal becomes:

$$S(t) = S_0 \exp(i\phi) \exp(i\Omega t) \exp\left(\frac{-t}{T_2}\right). \quad (4.3)$$

Phase correction

It turns out that for instrumental reasons the axis along which the signal appears cannot be predicted, so in any practical situation there is an unknown phase shift. In general, this leads to a situation in which the real part of the spectrum (which is normally the part we display) does not show a pure ab-

sorption lineshape. This is undesirable as for the best resolution we require an absorption mode lineshape.

Luckily, restoring the spectrum to the absorption mode is easy. Suppose we take the FID, represented by Eq. 4.3, and multiply it by $\exp(i\phi_{\text{corr}})$:

$$\exp(i\phi_{\text{corr}})S(t) = \exp(i\phi_{\text{corr}}) \times \left[S_0 \exp(i\phi) \exp(i\Omega t) \exp\left(\frac{-t}{T_2}\right) \right].$$

This is easy to do as by now the FID is stored in computer memory, so the multiplication is just a mathematical operation on some numbers. Exponentials have the property that $\exp(A)\exp(B) = \exp(A+B)$ so we can re-write the time domain signal as

$$\exp(i\phi_{\text{corr}})S(t) = \exp(i(\phi_{\text{corr}} + \phi)) \left[S_0 \exp(i\Omega t) \exp\left(\frac{-t}{T_2}\right) \right].$$

Now suppose that we set $\phi_{\text{corr}} = -\phi$; as $\exp(0) = 1$ the time domain signal becomes:

$$\exp(i\phi_{\text{corr}})S(t) = S_0 \exp(i\Omega t) \exp\left(\frac{-t}{T_2}\right).$$

The signal now has no phase shift and so will give us a spectrum in which the real part shows the absorption mode lineshape – which is exactly what we want. All we need to do is find the correct ϕ_{corr} .

It turns out that the phase correction can just as easily be applied to the spectrum as it can to the FID. So, if the spectrum is represented by $S(\omega)$ (a function of frequency, ω) the phase correction is applied by computing

$$\exp(i\phi_{\text{corr}})S(\omega).$$

Such a correction is called a *frequency independent* or *zero order* phase correction as it is the same for all peaks in the spectrum, regardless of their offset.

In practice what happens is that we Fourier transform the FID and display the real part of the spectrum. We then adjust the phase correction (i.e. the value of ϕ_{corr}) until the spectrum appears to be in the absorption mode – usually this adjustment is made by turning a knob or by a “click and drag” operation with the mouse. The whole process is called *phasing the spectrum* and is something we have to do each time we record a spectrum.

In addition to the phase shifts introduced by the spectrometer we can of course deliberately introduce a shift of phase by, for example, altering the phase of a pulse. In a sense it does not matter what the phase of the signal is – we can always obtain an absorption spectrum by phase correcting the spectrum later on.

Frequency dependent phase errors

We saw in section 3.11 that if the offset becomes comparable with the RF field strength a 90° pulse about x results in the generation of magnetization along both the x and y axes. This is in contrast to the case of a hard pulse, where

Attempts have been made over the years to automate this phasing process; on well resolved spectra the results are usually good, but these automatic algorithms tend to have more trouble with poorly-resolved spectra. In any case, what constitutes a correctly phased spectrum is rather subjective.

the magnetization appears only along $-y$. We can now describe this mixture of x and y magnetization as resulting in a *phase shift* or *phase error* of the spectrum.

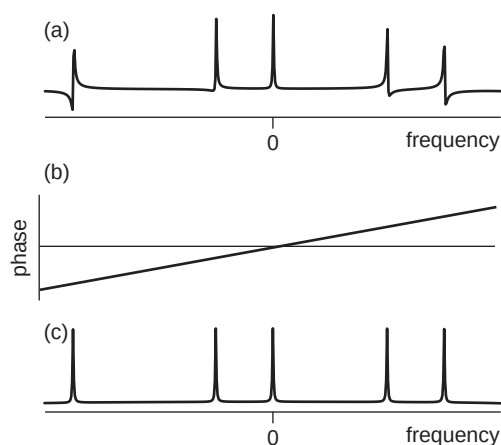


Fig. 4.8 Illustration of the appearance of a frequency dependent phase error in the spectrum. In (a) the line which is on resonance (at zero frequency) is in pure absorption, but as the offset increases the phase error increases. Such an frequency dependent phase error would result from the use of a pulse whose RF field strength was not much larger than the range of offsets. The spectrum can be returned to the absorption mode, (c), by applying a phase correction which varies with the offset in a linear manner, as shown in (b). Of course, to obtain a correctly phased spectrum we have to choose the correct slope of the graph of phase against offset.

Figure 3.25 illustrates very clearly how the x component increases as the offset increases, resulting in a phase error which also increases with offset. Therefore lines at different offsets in the spectrum will have different phase errors, the error increasing as the offset increases. This is illustrated schematically in the upper spectrum shown in Fig. 4.8.

If there were only one line in the spectrum it would be possible to ensure that the line appeared in the absorption mode simply by adjusting the phase in the way described above. However, if there is more than one line present in the spectrum the phase correction for each will be *different*, and so it will be impossible to phase all of the lines at once.

Luckily, it is often the case that the phase correction needed is directly proportional to the offset – called a *linear* or *first order* phase correction. Such a variation in phase with offset is shown in Fig. 4.8 (b). All we have to do is to vary the rate of change of phase with frequency (the slope of the line) until the spectrum appears to be phased; as with the zero-order phase correction the computer software usually makes it easy for us to do this by turning a knob or pushing the mouse. In practice, to phase the spectrum correctly usually requires some iteration of the zero- and first-order phase corrections.

The usual convention is to express the frequency dependent phase correction as the value that the phase takes at the extreme edges of the spectrum. So, for example, such a correction by 100° means that the phase correction is zero in the middle (at zero offset) and rises linearly to $+100^\circ$ at one edge and falls linearly to -100° at the opposite edge.

For a pulse the phase error due to these off-resonance effects for a peak with offset Ω is of the order of (Ωt_p) , where t_p is the length of the pulse. For

a carbon-13 spectrum recorded at a Larmor frequency of 125 MHz, the maximum offset is about 100 ppm which translates to 12500 Hz. Let us suppose that the 90° pulse width is 15 μ s, then the phase error is

$$2\pi \times 12500 \times 15 \times 10^{-6} \approx 1.2 \text{ radians}$$

which is about 68°; note that in the calculation we had to convert the offset from Hz to rad s⁻¹ by multiplying by 2 π . So, we expect the frequency dependent phase error to vary from zero in the middle of the spectrum (where the offset is zero) to 68° at the edges; this is a significant effect.

For reasons which we cannot go into here it turns out that the linear phase correction is sometimes only a first approximation to the actual correction needed. Provided that the lines in the spectrum are sharp a linear correction works very well, but for broad lines it is not so good. Attempting to use a first-order phase correction on such spectra often results in distortions of the baseline.

4.4 Sensitivity enhancement

Inevitably when we record a FID we also record noise at the same time. Some of the noise is contributed by the amplifiers and other electronics in the spectrometer, but the major contributor is the thermal noise from the coil used to detect the signal. Reducing the noise contributed by these two sources is largely a technical matter which will not concern us here. NMR is not a sensitive technique, so we need to take any steps we can to improve the signal-to-noise ratio in the spectrum. We will see that there are some manipulations we can perform on the FID which will give us some improvement in the signal-to-noise ratio (SNR).

By its very nature, the FID decays over time but in contrast the noise just goes on and on. Therefore, if we carry on recording data for long after the FID has decayed we will just measure noise and no signal. The resulting spectrum will therefore have a poor signal-to-noise ratio.

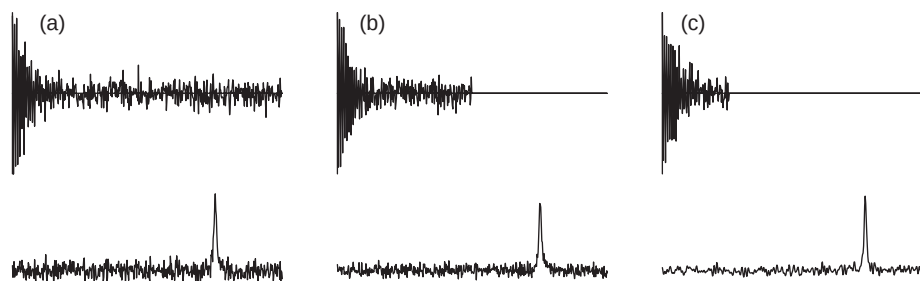


Fig. 4.9 Illustration of the effect of the time spent acquiring the FID on the signal-to-noise ratio (SNR) in the spectrum. In (a) the FID has decayed to next to nothing within the first quarter of the time, but the noise carries on unabated for the whole time. Shown in (b) is the effect of halving the time spent acquiring the data; the SNR improves significantly. In (c) we see that taking the first quarter of the data gives a further improvement in the SNR.

This point is illustrated in Fig. 4.9 where we see that by recording the FID for long after it has decayed all we end up doing is recording more noise

and no signal. Just shortening the time spent recording the signal (called the *acquisition time*) will improve the SNR since more or less all the signal is contained in the early part of the FID. Of course, we must not shorten the acquisition time too much or we will start to miss the FID, which would result in a reduction in SNR.

Sensitivity enhancement

Looking at the FID we can see that at the start the signal is strongest. As time progresses, the signal decays and so gets weaker but the noise remains at the same level. The idea arises, therefore, that the early parts of the FID are “more important” as it is here where the signal is the strongest.

This effect can be exploited by deliberately multiplying the FID by a function which starts at 1 and then steadily tails away to zero. The idea is that this function will cut off the later parts of the FID where the signal is weakest, but leave the early parts unaffected.

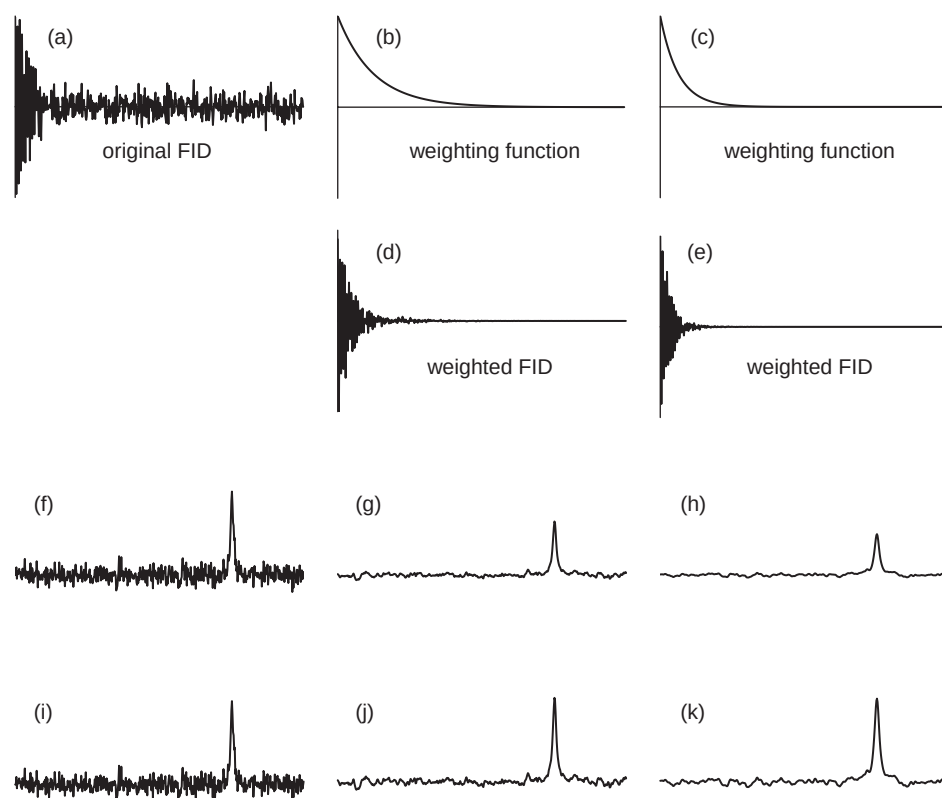


Fig. 4.10 Illustration of how multiplying a FID by a decaying function (a weighting function) can improve the SNR. The original FID is shown in (a) and the corresponding spectrum is (f). Multiplying the FID by a weighting function (b) gives (c); Fourier transformation of (c) gives the spectrum (g). Note the improvement in SNR of (g) compared to (f). Multiplying (a) by the more rapidly decaying weighting function (d) gives (e); the corresponding spectrum is (h). The improvement in SNR is less marked. Spectra (f) – (h) are all plotted on the same vertical scale so that the decrease in peak height can be seen. The same spectra are plotted in (i) – (k) but this time normalized so that the peak height is the same; this shows most clearly the improvement in the SNR.

A typical choice for this function – called a *weighting function* – is an

exponential:

$$W(t) = \exp(-R_{LB}t) \quad (4.4)$$

where R_{LB} is a rate constant which we are free to choose. Figure 4.10 illustrates the effect on the SNR of different choices of this decay constant.

Spectrum (g) shows a large improvement in the SNR when compared to spectrum (f) simply because the long tail of noise is suppressed. Using a more rapidly decaying weighting function, (d), gives a further small improvement in the SNR (h). It should not be forgotten that the weighting function also acts on the signal, causing it to decay more quickly. As was explained in section 4.2 a more rapidly decaying signal leads to a broader line. So, the use of the weighting function will not only attenuate the noise it will also broaden the lines; this is clear from Fig. 4.10.

Broadening the lines also reduces the peak height (remember that the integral remains constant) – something that is again evident from Fig. 4.10. Clearly, this reduction in peak height will reduce the SNR. Hence there is a trade-off we have to make: the more rapidly decaying the weighting function the more the noise in the tail of the FID is attenuated thus reducing the noise level in the spectrum. However, at the same time a more rapidly decaying function will cause greater line broadening, and this will reduce the SNR. It turns out that there is an optimum weighting function, called the *matched filter*.

Matched filter

Suppose that the FID can be represented by the exponentially decaying function introduced in Eq. 4.2:

$$S(t) = S_0 \exp(i\Omega t) \exp(-R_2 t).$$

The line in the corresponding spectrum is of width (R_2/π) Hz. Let us apply a weighting function of the kind described by Eq. 4.4 i.e. a decaying exponential:

$$W(t) \times S(t) = \exp(-R_{LB}t) [S_0 \exp(i\Omega t) \exp(-R_2 t)].$$

The decay due to the weighting function on its own would give a linewidth of (R_{LB}/π) Hz. Combining the two terms describing the exponential decays gives

$$W(t) \times S(t) = S_0 \exp(i\Omega t) \exp(-(R_{LB} + R_2)t).$$

From this we see that the weighted FID will give a linewidth of

$$(R_{LB} + R_2)/\pi.$$

In words, the linewidth in the spectrum is the sum of the linewidths in the original spectrum and the additional linebroadening imposed by the weighting function.

It is usual to specify the weighting function in terms of the extra line broadening it will cause. So, a “linebroadening of 1 Hz” is a function which

will increase the linewidth in the spectrum by 1 Hz. For example if 5 Hz of linebroadening is required then $(R_{LB}/\pi) = 5$ giving $R_{LB} = 15.7 \text{ s}^{-1}$.

It can be shown that the best SNR is obtained by applying a weighting function which matches the linewidth in the original spectrum – such a weighting function is called a *matched filter*. So, for example, if the linewidth is 2 Hz in the original spectrum, applying an additional line broadening of 2 Hz will give the optimum SNR.

We can see easily from this argument that if there is a range of linewidths in the spectrum we cannot find a value of the linebroadening which is the optimum for all the peaks. Also, the extra line broadening caused by the matched filter may not be acceptable on the grounds of the decrease in resolution it causes. Under these circumstances we may choose to use sufficient line broadening to cut off the excess noise in the tail of the FID, but still less than the matched filter.

4.5 Resolution enhancement

We noted in section 4.2 that the more rapidly the time domain signal decays the broader the lines become. A weighting function designed to improve the SNR inevitably leads to a broadening of the lines as such a function hastens the decay of the signal. In this section we will consider the opposite case, where the weighting function is designed to *narrow* the lines in the spectrum and so increase the resolution.

The basic idea is simple. All we need to do is to multiply the FID by a weighting function which *increases* with time, for example a rising exponential:

$$W(t) = \exp(+R_{RE}t) \quad R_{RE} > 0.$$

This function starts at 1 when $t = 0$ and then rises indefinitely.

The problem with multiplying the FID with such a function is that the noise in the tail of the FID is amplified, thus making the SNR in the spectrum very poor indeed. To get round this it is, after applying the positive exponential we multiply by a second decaying function to “clip” the noise at the tail of the FID. Usually, the second function is chosen to be one that decays relatively slowly over most of the FID and then quite rapidly at the end – a common choice is the Gaussian function:

$$W(t) = \exp(-\alpha t^2),$$

where α is a parameter which sets the decay rate. The larger α , the faster the decay rate.

The whole process is illustrated in Fig. 4.11. The original FID, (a), contains a significant amount of noise and has been recorded well beyond the point where the signal decays into the noise. Fourier transformation of (a) gives the spectrum (b). If (a) is multiplied by the rising exponential plotted in (c), the result is the FID (d); note how the decay of the signal has

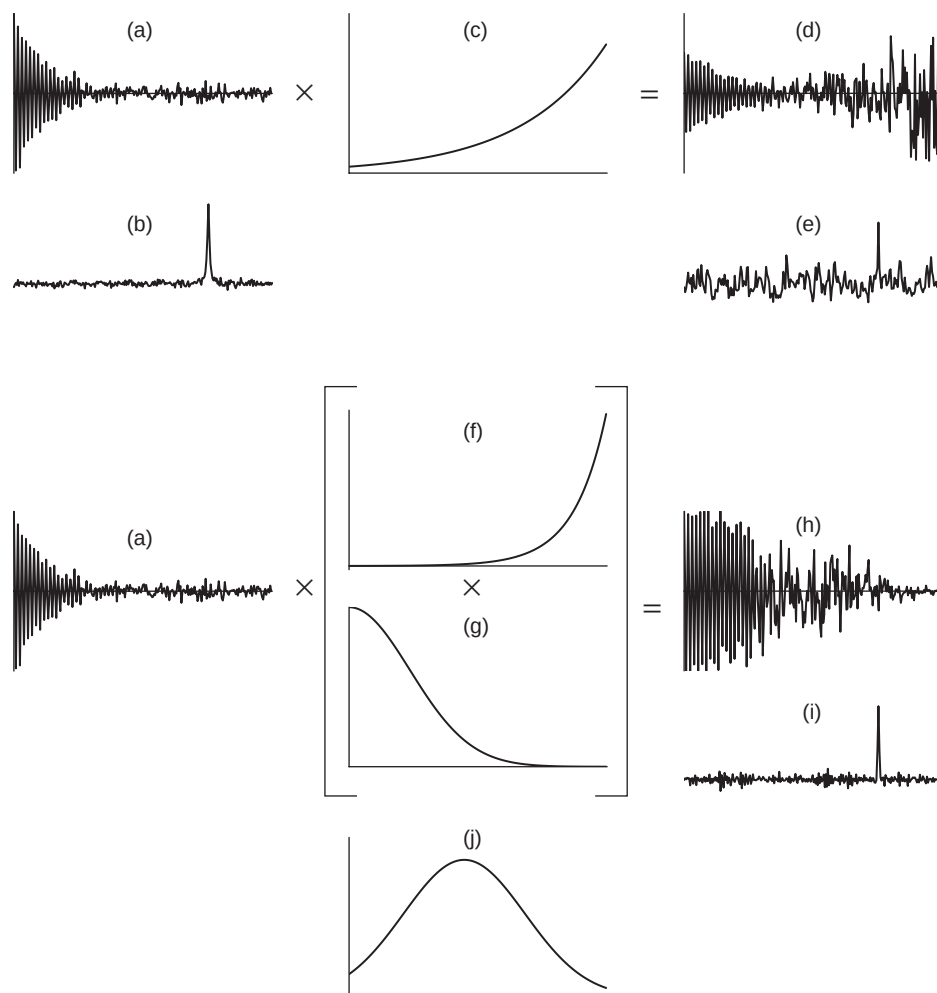


Fig. 4.11 Illustration of the use of weighting functions to enhance the resolution in the spectrum. Note that the scales of the plots have been altered to make the relevant features clear. See text for details.

been slowed, but the noise in the tail of the FID has been greatly magnified. Fourier transformation of (d) gives the spectrum (e); the resolution has clearly been improved, but at the expense of a large reduction in the SNR.

Referring now to the bottom part of Fig. 4.11 we can see the effect of introducing a Gaussian weighting function as well. The original FID (a) is multiplied by the rising exponential (f) and the decaying Gaussian (g); this gives the time-domain signal (h). Note that once again the signal decay has been slowed, but the noise in the tail of the FID is not as large as it is in (d). Fourier transformation of (h) gives the spectrum (i); the resolution has clearly been improved when compared to (b), but without too great a loss of SNR.

Finally, plot (j) shows the product of the two weighting functions (f) and (g). We can see clearly from this plot how the two functions combine together to first increase the time-domain function and then to attenuate it at longer times. Careful choice of the parameters R_{RE} and α are needed to obtain the optimum result. Usually, a process of trial and error is adopted.

If we can set R_{RE} to cancel exactly the original decay of the FID then the result of this process is to generate a time-domain function which only has a Gaussian decay. The resulting peak in the spectrum will have a Gaussian lineshape, which is often considered to be superior to the Lorentzian as it is narrower at the base; the two lineshapes are compared in Fig. 4.12. This transformation to a Gaussian lineshape is often called the *Lorentz-to-Gauss* transformation.

Parameters for the Lorentz-to-Gauss transformation

The combined weighting function for this transformation is

$$W(t) = \exp(R_{RE}t) \exp(-\alpha t^2).$$

The usual approach is to specify R_{RE} in terms of the linewidth that it would create on its own if it were used to specify a *decaying* exponential. Recall that in such a situation the linewidth, L , is given by R_{RE}/π . Thus $R_{RE} = \pi L$. The weighting function can therefore be rewritten

$$W(t) = \exp(-\pi Lt) \exp(-\alpha t^2).$$

For compatibility with the linebroadening role of a decaying exponential it is usual to define the exponential weighting function as $\exp(-\pi Lt)$ so that a positive value for L leads to line broadening and a negative value leads to resolution enhancement.

From Fig. 4.11 (j) we can see that the effect of the rising exponential and the Gaussian is to give an overall weighting function which has a maximum in it. It is usual to define the Gaussian parameter α from the position of this maximum. A little mathematics shows that this maximum occurs at time $t_{\max}$ given by:

$$t_{\max} = -\frac{L\pi}{2\alpha}$$

(recall that L is negative) so that

$$\alpha = -\frac{L\pi}{2t_{\max}}.$$

We simply select a value of $t_{\max}$ and use this to define the value of α . On some spectrometers, $t_{\max}$ is expressed as a fraction, f , of the acquisition time, t_{acq} : $t_{\max} = ft_{\text{acq}}$. In this case

$$\alpha = -\frac{L\pi}{2ft_{\text{acq}}}.$$

4.6 Other weighting functions

Many other weighting functions have been used for sensitivity enhancement and resolution enhancement. Perhaps the most popular are the *sine bell* variants on it, which are illustrated in Fig. 4.13.

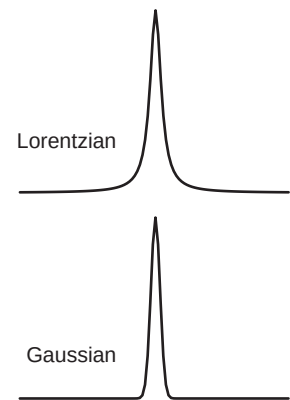


Fig. 4.12 Comparison of the Lorentzian and Gaussian lineshapes; the two peaks have been adjusted so that their peak heights and widths at half height are equal. The Gaussian is a more "compact" lineshape.

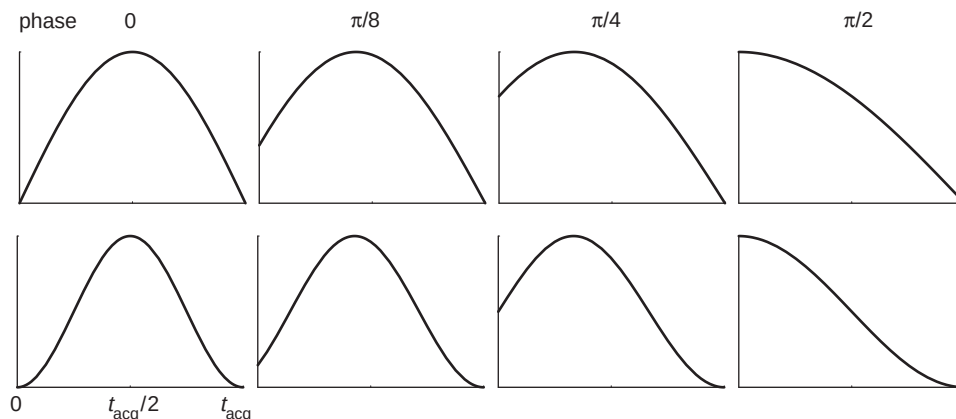


Fig. 4.13 The top row shows sine bell and the bottom row shows sine bell squared weighting functions for different choices of the phase parameter; see text for details.

The basic sine bell is just the first part of a $\sin \theta$ for $\theta = 0$ to $\theta = \pi$; this is illustrated in the top left-hand plot of Fig. 4.13. In this form the function will give resolution enhancement rather like the combination of a rising exponential and a Gaussian function (compare Fig. 4.11 (j)). The weighting function is chosen so that the sine bell fits exactly across the acquisition time; mathematically the required function is:

$$W(t) = \sin\left(\frac{\pi t}{t_{\text{acq}}}\right).$$

The sine bell can be modified by shifting it the left, as is shown in Fig. 4.13. The further the shift to the left the smaller the resolution enhancement effect will be, and in the limit that the shift is by $\pi/2$ or 90° the function is simply a decaying one and so will broaden the lines. The shift is usually expressed in terms of a phase ϕ (in radians); the resulting weighting function is:

$$W(t) = \sin\left(\frac{(\pi - \phi)t}{t_{\text{acq}}} + \phi\right).$$

Note that this definition of the function ensures that it goes to zero at t_{acq} .

The shape of all of these weighting functions are altered subtly by squaring them to give the *sine bell squared* functions; these are also shown in Fig. 4.13. The weighting function is then

$$W(t) = \sin^2\left(\frac{(\pi - \phi)t}{t_{\text{acq}}} + \phi\right).$$

Much of the popularity of these functions probably rests of the fact that there is only one parameter to adjust, rather than two in the case of the Lorentz-to-Gauss transformation.

4.7 Zero filling

Before being processed the FID must be converted into a digital form so that it can be stored in computer memory. We will have more to say about this

in Chapter 5 but for now we will just note that in this process the signal is sampled at regular intervals. The FID is therefore represented by a series of *data points*. When the FID is Fourier transformed the spectrum is also represented by a series of data points. So, although we plot the spectrum as a smooth line, it is in fact a series of closely spaced points.

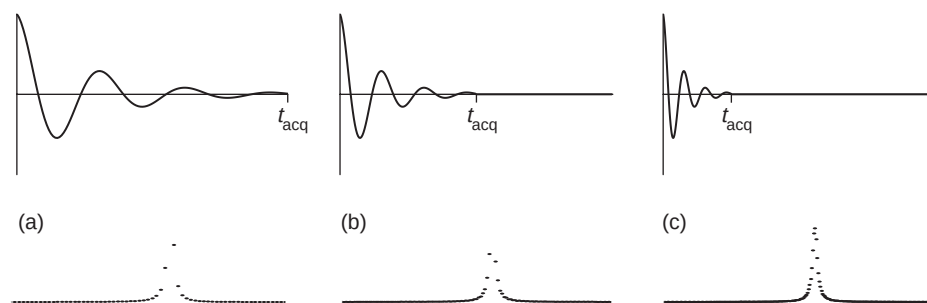


Fig. 4.14 Illustration of the results of zero filling. The FIDs along the top row have been supplemented with increasing numbers of zeroes and so contain more and more data points. Fourier transformation preserves the number of data points so the line in the spectrum is represented by more points as zeroes are added to the end of the FID. Note that the FID remains the same for all three cases; no extra data has been acquired.

This is illustrated in Fig. 4.14 (a) which shows the FID and the corresponding spectrum; rather than joining up the points which make up the spectrum we have just plotted the points. We can see that there are only a few data points which define the line in the spectrum.

If we take the original FID and add an equal number of zeroes to it, the corresponding spectrum has double the number of points and so the line is represented by more data points. This is illustrated in Fig. 4.14 (b). Adding a set of zeroes equal to the number of data points is called “one zero filling”.

We can carry on with this zero filling process. For example, having added one set of zeroes, we can add another to double the total number of data points (“two zero fillings”). This results in an even larger number of data points defining the line, as is shown in Fig. 4.14 (c).

Zero filling costs nothing in the sense that no extra data is required; it is just a manipulation in the computer. Of course, it does not improve the resolution as the measured signal remains the same, but the lines will be better defined in the spectrum. This is desirable, at least for aesthetic reasons if nothing else!

It turns out that the Fourier transform algorithm used by computer programs is most suited to a number of data points which is a power of 2. So, for example, $2^{14} = 16384$ is a suitable number of data points to transform, but 15000 is not. In practice, therefore, it is usual to zero fill the time domain data so that the total number of points is a power of 2; it is always an option, of course, to zero fill beyond this point.

4.8 Truncation

In conventional NMR it is virtually always possible to record the FID until it has decayed almost to zero (or into the noise). However, in multi-dimensional NMR this may not be the case, simply because of the restrictions on the amount of data which can be recorded, particularly in the “indirect dimension” (see Chapter 7 for further details). If we stop recording the signal before it has fully decayed the FID is said to be “truncated”; this is illustrated in Fig. 4.15.

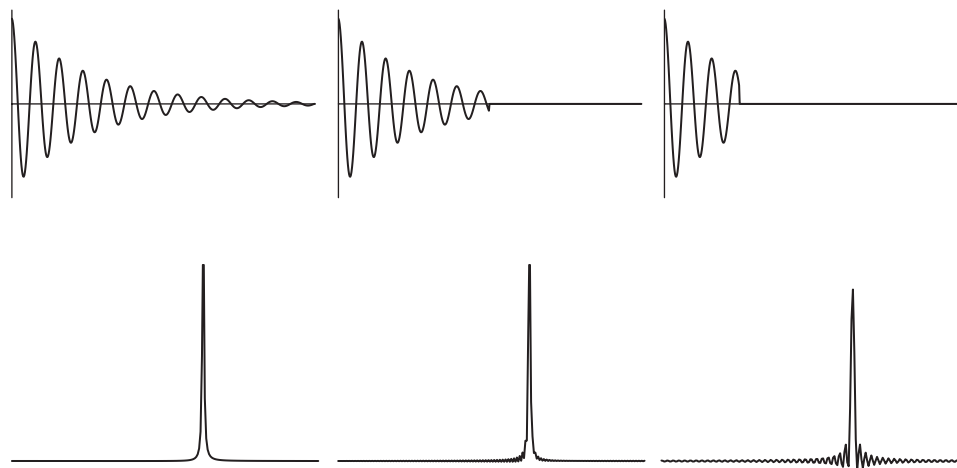


Fig. 4.15 Illustration of how truncation leads to artefacts (called *sinc wiggles*) in the spectrum. The FID on the left has been recorded for sufficient time that it has decayed almost to zero; the corresponding spectrum shows the expected lineshape. However, if data recording is stopped before the signal has fully decayed the corresponding spectra show oscillations around the base of the peak.

As is shown clearly in the figure, a truncated FID leads to oscillations around the base of the peak; these are usually called *sinc wiggles* or *truncation artefacts* – the name arises as the peak shape is related to a sinc function. The more severe the truncation, the larger the sinc wiggles. It is easy to show that the separation of successive maxima in these wiggles is $1/t_{\text{acq}}$ Hz.

Clearly these oscillations are undesirable as they may obscure nearby weaker peaks. Assuming that it is not an option to increase the acquisition time, the way forward is to apply a decaying weighting function to the FID so as to force the signal to go to zero at the end. Unfortunately, this will have the side effects of broadening the lines and reducing the SNR.

Highly truncated time domain signals are a feature of three- and higher-dimensional NMR experiments. Much effort has therefore been put into finding alternatives to the Fourier transform which will generate spectra without these truncation artefacts. The popular methods are *maximum entropy*, *linear prediction* and *FDM*. Each has its merits and drawbacks; they all need to be applied with great care.

4.9 Exercises

E 4-1

In a spectrum with just one line, the dispersion mode lineshape might be acceptable – in fact we can think of reasons why it might even be desirable (what might these be?). However, in a spectrum with many lines the dispersion mode lineshape is very undesirable – why?

E 4-2

Suppose that we record a spectrum with the simple pulse-acquire sequence using a 90° pulse applied along the x axis. The resulting FID is Fourier transformed and the spectrum is phased to give an absorption mode lineshape.

We then change the phase of the pulse from x to y , acquire an FID in the same way and phase the spectrum using the *same* phase correction as above. What lineshape would you expect to see in the spectrum; give the reasons for your answer.

How would the spectrum be affected by: (a) applying the pulse about $-x$; (b) changing the pulse flip angle to 270° about x ?

E 4-3

The gyromagnetic ratio of phosphorus-31 is $1.08 \times 10^8 \text{ rad s}^{-1} \text{ T}^{-1}$. This nucleus shows a wide range of shifts, covering some 700 ppm.

Suppose that the transmitter is placed in the middle of the shift range and that a 90° pulse of width $20 \mu\text{s}$ is used to excite the spectrum. Estimate the size of the phase correction which will be needed at the edges of the spectrum. (Assume that the spectrometer has a B_0 field strength of 9.4 T).

E 4-4

Why is it undesirable to continue to acquire the FID after the signal has decayed away?

How can weighting functions be used to improve the SNR of a spectrum? In your answer described how the parameters of a suitable weighting function can be chosen to optimize the SNR. Are there any disadvantages to the use of such weighting functions?

E 4-5

Describe how weighting functions can be used to improve the resolution in a spectrum. What sets the limit on the improvement that can be obtained in practice? Is zero filling likely to improve the situation?

E 4-6

Explain why use of a sine bell weighting function shifted by 45° may enhance the resolution but use of a sine bell shifted by 90° does not.

E 4-7

In a proton NMR spectrum the peak from TMS was found to show “wiggles” characteristic of truncation of the FID. However, the other peaks in the spectrum showed no such artefacts. Explain.

How can truncation artefacts be suppressed? Mention any difficulties with your solution to the problem.

5 How the spectrometer works

NMR spectrometers have now become very complex instruments capable of performing an almost limitless number of sophisticated experiments. However, the really important parts of the spectrometer are not that complex to understand in outline, and it is certainly helpful when using the spectrometer to have some understanding of how it works.

Broken down to its simplest form, the spectrometer consists of the following components:

- An intense, homogeneous and stable magnetic field
- A “probe” which enables the coils used to excite and detect the signal to be placed close to the sample
- A high-power RF transmitter capable of delivering short pulses
- A sensitive receiver to amplify the NMR signals
- A digitizer to convert the NMR signals into a form which can be stored in computer memory
- A “pulse programmer” to produce precisely times pulses and delays
- A computer to control everything and to process the data

We will consider each of these in turn.

5.1 The magnet

Modern NMR spectrometers use persistent superconducting magnets to generate the B_0 field. Basically such a magnet consists of a coil of wire through which a current passes, thereby generating a magnetic field. The wire is of a special construction such that at low temperatures (less than 6 K, typically) the resistance goes to zero – that is the wire is *superconducting*. Thus, once the current is set running in the coil it will persist for ever, thereby generating a magnetic field without the need for further electrical power. Superconducting magnets tend to be very stable and so are very useful for NMR.

To maintain the wire in its superconducting state the coil is immersed in a bath of liquid helium. Surrounding this is usually a “heat shield” kept at 77 K by contact with a bath of liquid nitrogen; this reduces the amount of (expensive) liquid helium which boils off due to heat flowing in from the surroundings. The whole assembly is constructed in a vacuum flask so as to further reduce the heat flow. The cost of maintaining the magnetic field

is basically the cost of liquid helium (rather expensive) and liquid nitrogen (cheap).

Of course, we do not want the sample to be at liquid helium temperatures, so a room temperature region – accessible to the outside world – has to be engineered as part of the design of the magnet. Usually this room temperature zone takes the form of a vertical tube passing through the magnet (called the *bore tube* of the magnet); the magnetic field is in the direction of this tube.

Shims

The lines in NMR spectra are very narrow – linewidths of 1 Hz or less are not uncommon – so the magnetic field has to be very homogeneous, meaning that it must not vary very much over space. The reason for this is easily demonstrated by an example.

Consider a proton spectrum recorded at 500 MHz, which corresponds to a magnetic field of 11.75 T. Recall that the Larmor frequency is given by

$$\nu_0 = -\frac{1}{2\pi}\gamma B_0 \quad (5.1)$$

where γ is the gyromagnetic ratio ($2.67 \times 10^8 \text{ rad s}^{-1} \text{ T}^{-1}$ for protons). We need to limit the variation in the magnetic field across the sample so that the corresponding variation in the Larmor frequency is much less than the width of the line, say by a factor of 10.

With this condition, the maximum acceptable change in Larmor frequency is 0.1 Hz and so using Eq. 5.1 we can compute the change in the magnetic field as $2\pi 0.1/\gamma = 2.4 \times 10^{-9} \text{ T}$. Expressed as a fraction of the main magnetic field this is about 2×10^{-10} . We can see that we need to have an extremely homogeneous magnetic field for work at this resolution.

On its own, no superconducting magnet can produce such a homogeneous field. What we have to do is to surround the sample with a set of *shim coils*, each of which produces a tiny magnetic field with a particular spatial profile. The current through each of these coils is adjusted until the magnetic field has the required homogeneity, something we can easily assess by recording the spectrum of a sample which has a sharp line. Essentially how this works is that the magnetic fields produced by the shims are cancelling out the small residual inhomogeneities in the main magnetic field.

Modern spectrometers might have up to 40 different shim coils, so adjusting them is a very complex task. However, once set on installation it is usually only necessary on a day to day basis to alter a few of the shims which generate the simplest field profiles.

The shims are labelled according to the field profiles they generate. So, for example, there are usually shims labelled x, y and z, which generate magnetic fields varying the corresponding directions. The shim z^2 generates a field that varies quadratically along the z direction, which is the direction of B_0 . There are more shims whose labels you will recognize as corresponding to the names of the hydrogen atomic orbitals. This is no coincidence; the magnetic

field profiles that the shims coils create are in fact the spherical harmonic functions, which are the angular parts of the atomic orbitals.

5.2 The probe

The probe is a cylindrical metal tube which is inserted into the bore of the magnet. The small coil used to both excite and detect the NMR signal (see sections 3.2 and 3.4) is held in the top of this assembly in such a way that the sample can come down from the top of the magnet and drop into the coil. Various other pieces of electronics are contained in the probe, along with some arrangements for heating or cooling the sample.

The key part of the probe is the small coil used to excite and detect the magnetization. To optimize the sensitivity this coil needs to be as close as possible to the sample, but of course the coil needs to be made in such a way that the sample tube can drop down from the top of the magnet into the coil. Extraordinary effort has been put into the optimization of the design of this coil.

The coil forms part of a *tuned circuit* consisting of the coil and a capacitor. The inductance of the coil and the capacitance of the capacitor are set such that the tuned circuit they form is resonant at the Larmor frequency. That the coil forms part of a tuned circuit is very important as it greatly increases the detectable current in the coil.

Spectroscopists talk about “tuning the probe” which means adjusting the capacitor until the tuned circuit is resonant at the Larmor frequency. Usually we also need to “match the probe” which involves further adjustments designed to maximize the power transfer between the probe and the transmitter and receiver; Fig. 5.1 shows a typical arrangement. The two adjustments tend to interact rather, so tuning the probe can be a tricky business. To aid us, the instrument manufacturers provide various indicators and displays so that the tuning and matching can be optimized. We expect the tuning of the probe to be particularly sensitive to changing solvent or to changing the concentration of ions in the solvent.

5.3 The transmitter

The radiofrequency transmitter is the part of the spectrometer which generates the pulses. We start with an RF source which produces a stable frequency which can be set precisely. The reason why we need to be able to set the frequency is that we might want to move the transmitter to different parts of the spectrum, for example if we are doing experiments involving selective excitation (section 3.11).

Usually a *frequency synthesizer* is used as the RF source. Such a device has all the desirable properties outlined above and is also readily controlled by a computer interface. It is also relatively easy to phase shift the output from such a synthesizer, which is something we will need to do in order to

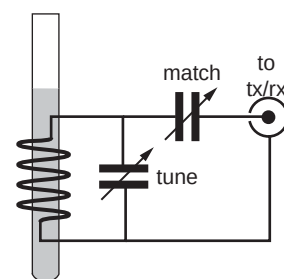


Fig. 5.1 Schematic of the key parts of the probe. The coil is shown on the left (with the sample tube in grey); a tuned circuit is formed by a capacitor (marked “tune”). The power transfer to the transmitter and receiver (tx. and rx.) is optimized by adjusting the capacitor marked “match”. Note that the coil geometry as shown is not suitable for a superconducting magnet in which the main field is parallel to the sample axis.

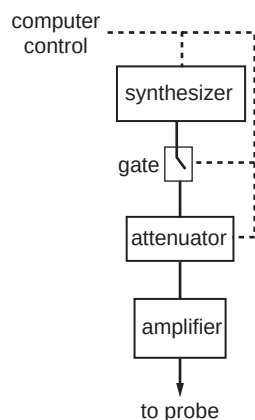


Fig. 5.2 Typical arrangement of the RF transmitter. The synthesizer which is the source of the RF, the attenuator and the gate used to create the pulses all under computer control.

create phase shifted pulses.

As we only need the RF to be applied for a short time, the output of the synthesizer has to be “gated” so as to create a pulse of RF energy. Such a gate will be under computer control so that the length and timing of the pulse can be controlled.

The RF source is usually at a low level (a few mW) and so needs to be boosted considerably before it will provide a useful B_1 field when applied to the probe; the complete arrangement is illustrated in Fig. 5.2. RF amplifiers are readily available which will boost this small signal to a power of 100 W or more. Clearly, the more power that is applied to the probe the more intense the B_1 field will become and so the shorter the 90° pulse length. However, there is a limit to the amount of power which can be applied because of the high voltages which are generated in the probe.

When the RF power is applied to the tuned circuit of which the coil is part, high voltages are generated across the tuning capacitor. Eventually, the voltage will reach a point where it is sufficient to ionize the air, thus generating a discharge or arc (like a lightning bolt). Not only does this *probe arcing* have the potential to destroy the coil and capacitor, but it also results in unpredictable and erratic B_1 fields. Usually the manufacturer states the power level which is “safe” for a particular probe.

Power levels and “dB”

The spectrometer usually provides us with a way of altering the RF power level and hence the strength of the B_1 field. This is useful as we may wish to set the B_1 field strength to a particular level, for example for selective excitation.

The usual way of achieving this control is add an *attenuator* between the RF source (the synthesizer) and the amplifier. As its name implies, the attenuator reduces the signal as it passes through.

The attenuation is normally expressed in *decibels* (abbreviated dB and pronounced “dee-bee”). If the input power is P_{in} and the output power is P_{out} the attenuation in dB is

$$10 \times \log_{10} \frac{P_{out}}{P_{in}};$$

note that the logarithm is to the base 10, not the natural logarithm. The factor of 10 is the “deci” part in the dB.

For example, if the output power is half the input power, i.e. $P_{in} = 2 \times P_{out}$, the power ratio in dB is

$$10 \times \log_{10} \frac{P_{out}}{P_{in}} = 10 \times \log_{10} \frac{1}{2} = -3.0$$

So, halving the power corresponds to a change of -3.0 dB, the minus indicating that there is a power reduction i.e. an attenuation. An attenuator which achieves this effect would be called “a 3 dB attenuator”.

Likewise, a power reduction by a factor of 4 corresponds to -6.0 dB. In fact, because of the logarithmic relationship we can see that each 3 dB of

attenuation will halve the power. So, a 12 dB attenuator will reduce the power by a factor of 16.

It turns out that the B_1 field strength is proportional to the square root of the power applied. The reason for this is that it is the current in the coil which is responsible for the B_1 field and, as we know from elementary electrical theory, power is equal to (resistance $\times$ current²). So, the current is proportional to the square root of the power.

Therefore, in order to double the B_1 field we need to double the current and this means multiplying the power by a factor of four. Recall from the above that a change in the power by a factor of 4 corresponds to 6 dB. Thus, 6 dB of attenuation causes the B_1 field to fall to half its original value; 12 dB would cause it to fall to one quarter of its initial value, and so on.

Usually the attenuator is under computer control and its value can be set in dB. This is very helpful to us as we can determine the attenuation needed for different B_1 field strengths. Suppose that with a certain setting of the attenuator we have determined the B_1 field strength to be $\omega_1^{\text{init}}/2\pi$, where we have expressed the field strength in frequency units as described in section 3.4. However, in our experiment we want the field strength to be $\omega_1^{\text{new}}/2\pi$. The ratio of the powers needed to achieve these two field strengths is equal to the square of the ratio of the field strengths:

$$\text{power ratio} = \left(\frac{\omega_1^{\text{new}}/2\pi}{\omega_1^{\text{init}}/2\pi} \right)^2.$$

Expressed in dB this is

$$\begin{aligned} \text{power ratio in dB} &= 10 \log_{10} \left(\frac{\omega_1^{\text{new}}/2\pi}{\omega_1^{\text{init}}/2\pi} \right)^2 \\ &= 20 \log_{10} \left(\frac{\omega_1^{\text{new}}/2\pi}{\omega_1^{\text{init}}/2\pi} \right). \end{aligned}$$

This expression can be used to find the correct setting for the attenuator.

As the duration of a pulse of a given flip angle is *inversely* proportional to $\omega_1/2\pi$ the relationship can be expressed in terms of the initial and new pulse widths, τ_{init} and τ_{new} :

$$\text{power ratio in dB} = 20 \log_{10} \left(\frac{\tau_{\text{init}}}{\tau_{\text{new}}} \right).$$

For example, suppose we have calibrated the pulse width for a 90° pulse to be 15 μs but what we actually wanted was a 90° pulse of 25 μs . The required attenuation would be:

$$\text{power ratio in dB} = 20 \log_{10} \left(\frac{15}{25} \right) = -4.4 \text{ dB}.$$

We would therefore need to increase the attenuator setting by 4.4 dB.

Phase shifted pulses

In section 3.9 we saw that pulses could be applied along different axes in the rotating frame, for example x and y . At first sight you might think that to achieve this we would need a coil oriented along the y axis to apply a pulse along y . However, this is not the case; all we need to do is to shift the phase of the RF by 90° .

In section 3.4 we described the situation where RF power at frequency ω_{RF} is applied to the coil. In this situation the B_1 field along the x axis can be written as

$$2B_1 \cos \omega_{\text{RF}}t.$$

If we phase shift the RF by 90° the field will be modulated by a sine function rather than a cosine:

$$2B_1 \sin \omega_{\text{RF}}t.$$

How this comes about is illustrated in Fig. 5.3.

The difference between the cosine and sine modulations is that in the former the field starts at a maximum and then decays, whereas in the latter it starts at zero and increases.

The resulting sine modulated field can still be thought of as arising from two counter-rotating fields, B_1^+ and B_1^- ; this is shown in Fig. 5.4. The crucial thing is that as the field is zero at time zero the two counter-rotating vectors start out opposed. As they evolve from this position the component along x grows, as required. Recall from section 3.4 that it is the component B_1^- which interacts with the magnetization. When we move to a rotating frame at ω_{RF} this component will be along the y axis (the position it has at time zero). This is in contrast to the case in Fig. 3.8 where B_1^- starts out on the x axis.

So, if the B_1 field is modulated by a cosine wave the result is a pulse about x , whereas if it is modulated by a sine wave the result is a pulse about y . We see that by phase shifting the RF we can apply pulses about any axis; there is no need to install more than one coil in the probe.

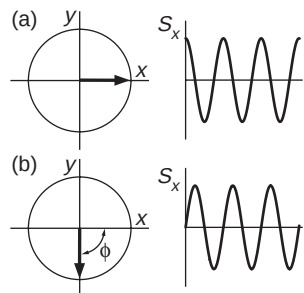


Fig. 5.3 Illustration of how a 90° phase shift takes us from cosine to sine modulation. In (a) the vector starts out along x ; as it rotates in a positive sense the x component, S_x , varies as a cosine wave, as is shown in the graph on the right. In (b) the vector starts on $-y$; the x component now takes the form of a sine wave, as is shown in the graph. Situation (b) is described as a phase shift, ϕ , of -90° compared to (a).

5.4 The receiver

The NMR signal emanating from the probe is very small (of the order of μV) but for modern electronics there is no problem in amplifying this signal to a level where it can be digitized. These amplifiers need to be designed so that they introduce a minimum of extra noise (they should be *low-noise* amplifiers).

The first of these amplifiers, called the *pre-amplifier* or *pre-amp* is usually placed as close to the probe as possible (you will often see it resting by the foot of the magnet). This is so that the weak signal is boosted before being sent down a cable to the spectrometer console.

One additional problem which needs to be solved comes about because the coil in the probe is used for both exciting the spins and detecting the signal. This means that at one moment hundreds of Watts of RF power are being

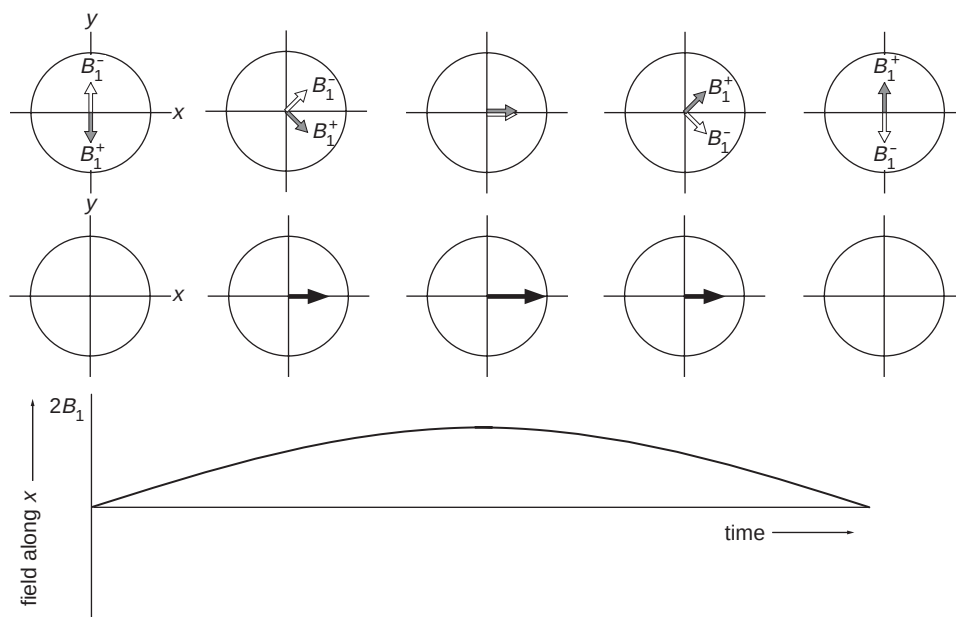


Fig. 5.4 Illustration of the decomposition of a field oscillating along the x axis according to a sine wave into two counter rotating components. This figure should be compared to Fig. 3.8. Note that the component B_1^- starts out along the y axis at time zero in contrast to the case of Fig. 3.8 where it starts out along x .

applied and the next we are trying to detect a signal of a few μV . We need to ensure that the high-power pulse does not end up in the sensitive receiver, thereby destroying it!

This separation of the receiver and transmitter is achieved by a gadget known as a *diplexer*. There are various different ways of constructing such a device, but at the simplest level it is just a fast acting switch. When the pulse is on the high power RF is routed to the probe and the receiver is protected by disconnecting it or shorting it to ground. When the pulse is off the receiver is connected to the probe and the transmitter is disconnected.

Some diplexers are passive in the sense that they require no external power to achieve the required switching. Other designs use fast electronic switches (rather like the gate in the transmitter) and these are under the command of the pulse programmer so that the receiver or transmitter are connected to the probe at the right times.

5.5 Digitizing the signal

The analogue to digital converter

A device known as an *analogue to digital converter* or ADC is used to convert the NMR signal from a voltage to a binary number which can be stored in computer memory. The ADC samples the signal at regular intervals, resulting in a representation of the FID as *data points*.

The output from the ADC is just a number, and the largest number that the ADC can output is set by the number of binary “bits” that the ADC uses. For example with only three bits the output of the ADC could take just 8

values: the binary numbers 000, 001, 010, 011, 100, 101, 110 and 111. The smallest number is 0 and the largest number is decimal 7 (the total number of possibilities is 8, which is 2 raised to the power of the number of bits). Such an ADC would be described as a “3 bit ADC”.

The waveform which the ADC is digitizing is varying continuously, but output of the 3 bit ADC only has 8 levels so what it has to do is simply output the level which is closest to the current input level; this is illustrated in Fig. 5.5. The numbers that the ADC outputs are therefore an approximation to the actual waveform.

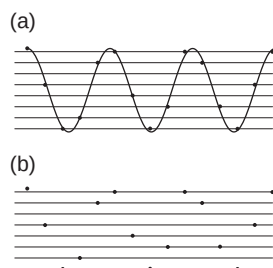


Fig. 5.5 Digitization of a waveform using an ADC with 8 levels (3 bits). The output of the ADC can only be one of the 8 levels, so the smoothly varying waveform has to be approximated by data points at one of the 8 levels. The data points, indicated by filled circles, are therefore an approximation to the true waveform.

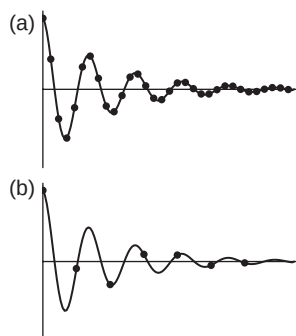


Fig. 5.7 Illustration of the effect of sampling rate on the representation of the FID. In (a) the data points (shown by dots) are quite a good representation of the signal (shown by the continuous line). In (b) the data points are too widely separated and so are a very poor representation of the signal.

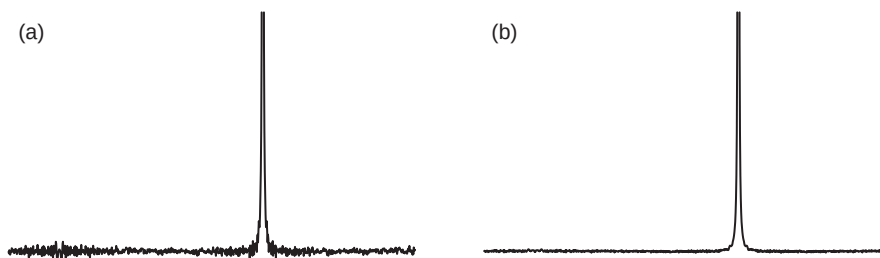


Fig. 5.6 Spectrum (a) is from an FID which has been digitized using a 6 bit ADC (i.e. 64 levels); the vertical scale has been expanded 10 fold so that the digitization sidebands are clearly visible. Spectrum (b) is from an FID which has been digitized using an 8 bit ADC (256 levels); the improvement over (a) is evident.

The approximation can be improved by increasing the number of bits; this gives more output levels. At present, ADC with between 16 and 32 bits are commonly in use in NMR spectrometers. The move to higher numbers of bits is limited by technical considerations.

The main consequence of the approximation process which the ADC uses is the generation of a forest of small sidebands – called *digitization sidebands* – around the base of the peaks in the spectrum. Usually these are not a problem as they are likely to be swamped by noise. However, if the spectrum contains a very strong peak the sidebands from it can swamp a nearby weak peak. Improving the resolution of the ADC, i.e. the number of bits it uses, reduces the digitization sidebands, as is shown in Fig. 5.6.

Sampling rates

Given that the ADC is only going to sample the signal at regular intervals the question arises as to how frequently it is necessary to sample the FID i.e. what should the time interval between the data points be. Clearly, if the time interval is too long we will miss crucial features of the waveform and so the digitized points will be a poor representation of the signal. This is illustrated in Fig. 5.7.

It turns out that if the interval between the points is Δ the highest frequency which can be represented correctly, $f_{\max}$, is given by

$$f_{\max} = \frac{1}{2\Delta};$$

$f_{\max}$ is called the *Nyquist frequency*. Usually we think of this relationship the other way round i.e. if we wish to represent correctly frequencies up to $f_{\max}$

the sampling interval is given by:

$$\Delta = \frac{1}{2f_{\max}}.$$

This sampling interval is often called the *dwell time*. A signal at $f_{\max}$ will have two data points per cycle.

We will see shortly that we are able to distinguish positive and negative frequencies, so a dwell time of Δ means that the range of frequencies from $-f_{\max}$ to $+f_{\max}$ are represented correctly.

A signal at greater than $f_{\max}$ will still appear in the spectrum but not at the correct frequency; such a peak is said to be *folded*. For example a peak at $(f_{\max} + F)$ will appear in the spectrum at $(-f_{\max} + F)$ as is illustrated in Fig. 5.8.

This Nyquist condition quickly brings us to a problem. A typical NMR frequency is of the order of hundreds of MHz but there simply are no ADCs available which work fast enough to digitize such a waveform with the kind of accuracy (i.e. number of bits) we need for NMR. The solution to this problem is to mix down the signal to a lower frequency, as is described in the next section.

Mixing down to a lower frequency

Luckily for us, although the NMR frequency is quite high, the range of frequencies that any typical spectrum covers is rather small. For example, the 10 ppm of a proton spectrum recorded at 400 MHz covers just 4000 Hz. Such a small range of frequencies is easily within the capability of suitable ADCs.

The procedure is to take the NMR signals and then *subtract* a frequency from them all such that we are left with a much lower frequency. Typically, the frequency we subtract is set somewhere in the middle of the spectrum; this frequency is called the *receiver reference frequency* or just the *receiver frequency*.

For example, suppose that the 10 ppm range of proton shifts runs from 400.000 MHz to 400.004 MHz. If we set the receiver reference frequency at 400.002 MHz and then subtract this from the NMR frequencies we end up with a range -2000 to $+2000$ Hz.

The subtraction process is carried out by a device called a *mixer*. There are various ways of actually making a mixer – some are “active” and contain the usual transistors and so on, some are passive and are mainly constructed from transformers and diodes. The details need not concern us – all we need to know is that a mixer takes two signal inputs at frequencies f_1 and f_2 and produces an output signal which contains the signals at the sum of the two inputs, $(f_1 + f_2)$ and the difference $(f_1 - f_2)$. The process is visualized in Fig. 5.9.

Usually one of the inputs is generated locally, typically by a synthesizer. This signal is often called the *local oscillator*. The other input is the NMR signal the probe. In the spectrometer the local oscillator is set to a required

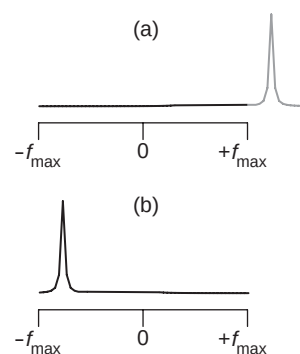


Fig. 5.8 Illustration of the concept of folding. In spectrum (a) the peak (shown in grey) is at a higher frequency than the maximum set by the Nyquist condition. In practice, such a peak would appear in the position shown in (b).

receiver frequency so that the difference frequency will be quite small. The other output of the mixer is the sum frequency, which will be at the order of twice the Larmor frequency. This high frequency signal is easily separated from the required low frequency signal by passing the output of the mixer through a *low pass filter*. The filtered signal is then passed to the ADC.

5.6 Quadrature detection

In the previous section we saw that the signal which is passed to the ADC contains positive and negative frequencies – essentially this corresponds to positive and negative offsets. The question therefore arises as to whether or not we are able to discriminate between frequencies which only differ in their sign; the answer is yes, but only provided we use a method known as *quadrature detection*.

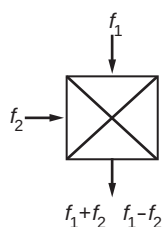


Fig. 5.9 A radiofrequency mixer takes inputs at two different frequencies and produces an output which contains the sum and difference of the two inputs.

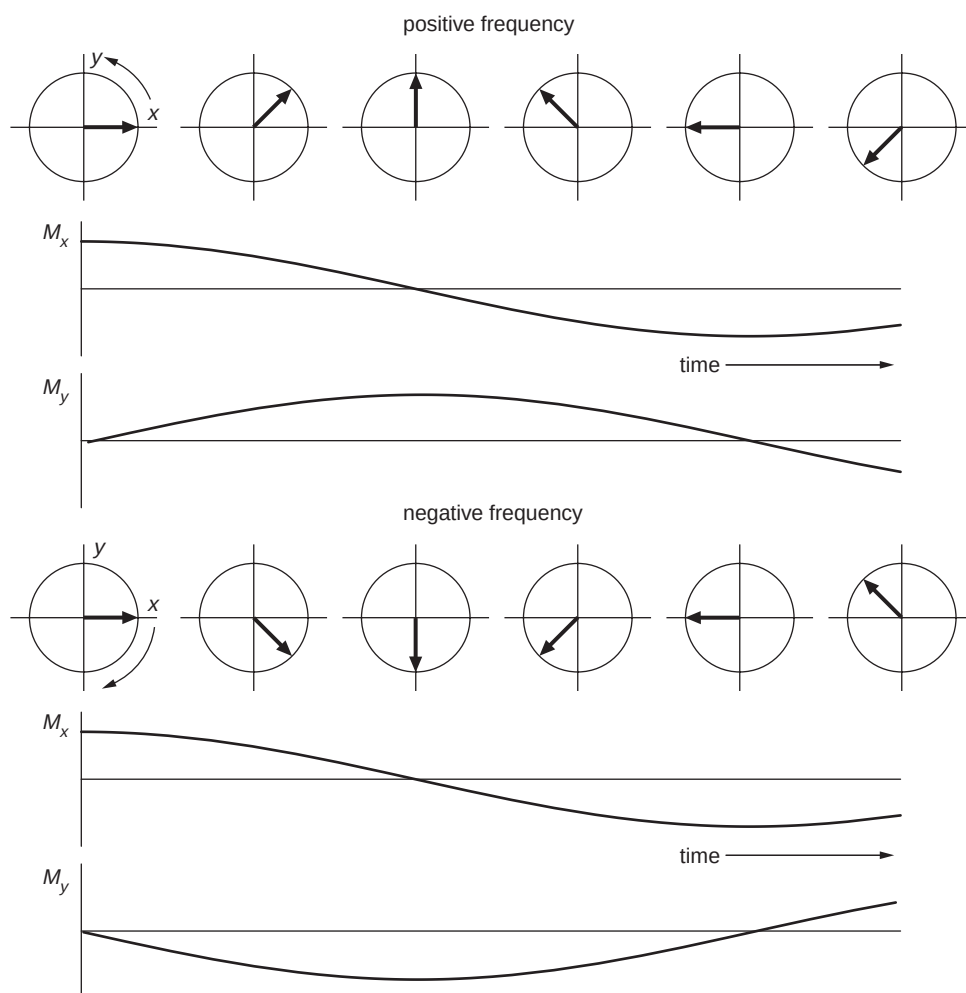


Fig. 5.10 Illustration of the x and y components from a vector precessing as a positive and negative offset. The only difference between the two is in the y component.

Figure 5.10 illustrates the problem. The upper part shows a vector precessing with a positive offset frequency (in the rotating frame); the x component is

a cosine wave and the y component is a sine wave. The lower part shows the evolution of a vector with a negative frequency. The x component remains the same but the y component has changed sign when compared with evolution at a positive frequency. We thus see that to distinguish between positive and negative frequencies we need to know *both* the x and y components of the magnetization.

It is tempting to think that it would be sufficient just to know the y component as positive frequencies give a positive sine wave and negative ones give a negative sine wave. In the corresponding spectrum this would mean that peaks at positive offsets would be positive and those at negative offsets would be negative. If there were a few well separated peaks in the spectrum such a spectrum might be interpretable, but if there are many peaks in the spectrum the result would be an uninterpretable mess.

How then can we detect both the x and y components of the magnetization? One idea is that we should have two coils, one along x and one along y : these would certainly detect the x and y components. In practice, it turns out to be very hard to achieve such an arrangement, partly because of the confined space in the probe and partly because of the difficulties in making the two coils electrically isolated from one another.

Luckily, though, it turns out to be easy to achieve the same effect by phase shifting the receiver reference frequency; we describe how this works in the next section.

The mixing process

In section 5.5 we described the use of a mixer to subtract a locally generated frequency from that of the NMR signal. The mixer achieves this by *multiplying* together the two input signals, and it is instructive to look at the consequence of this in more detail.

Suppose that the signal coming from the probe can be written as $A \cos \omega_0 t$ where A gives the overall intensity and as usual ω_0 is the Larmor frequency. Let us write the local oscillator signal as $\cos \omega_{\text{rx}} t$, where ω_{rx} is the receiver frequency. Multiplying these two together gives:

$$A \cos \omega_0 t \times \cos \omega_{\text{rx}} t = \frac{1}{2} A [\cos(\omega_0 + \omega_{\text{rx}})t + \cos(\omega_0 - \omega_{\text{rx}})t]$$

where we have used the well known formula $\cos A \cos B = \frac{1}{2}(\cos(A + B) + \cos(A - B))$.

After the low pass filter only the term $\cos(\omega_0 - \omega_{\text{rx}})t$ will survive. We recognize the difference frequency $(\omega_0 - \omega_{\text{rx}})$ as the separation between the receiver frequency and the Larmor frequency. If the receiver frequency is made to be the same as the transmitter frequency (as is usually the case) this difference frequency is the offset we identified in section 3.4.

In summary, the output of the mixer is at the offset frequency in the rotating frame and takes the form of a cosine modulation; it is thus equivalent to detecting the x component of the magnetization in the rotating frame.

Now consider the case where the local oscillator signal is $-\sin \omega_{\text{rx}}t$; as we saw in section 5.3 such a signal can be created by a phase shift of 90° . The output of the mixer is now given by

$$A \cos \omega_0 t \times -\sin \omega_{\text{rx}} t = \frac{1}{2}A [-\sin(\omega_0 + \omega_{\text{rx}})t + \sin(\omega_0 - \omega_{\text{rx}})t]$$

where we have used the well known formula $\cos A \sin B = \frac{1}{2}(\sin(A + B) - \sin(A - B))$. We now see that the difference frequency term is sine modulated at the offset frequency; it is the equivalent of the y component in the rotating frame. Thus, simply by shifting the phase of the local oscillator signal we can detect both the x and y components in the rotating frame.

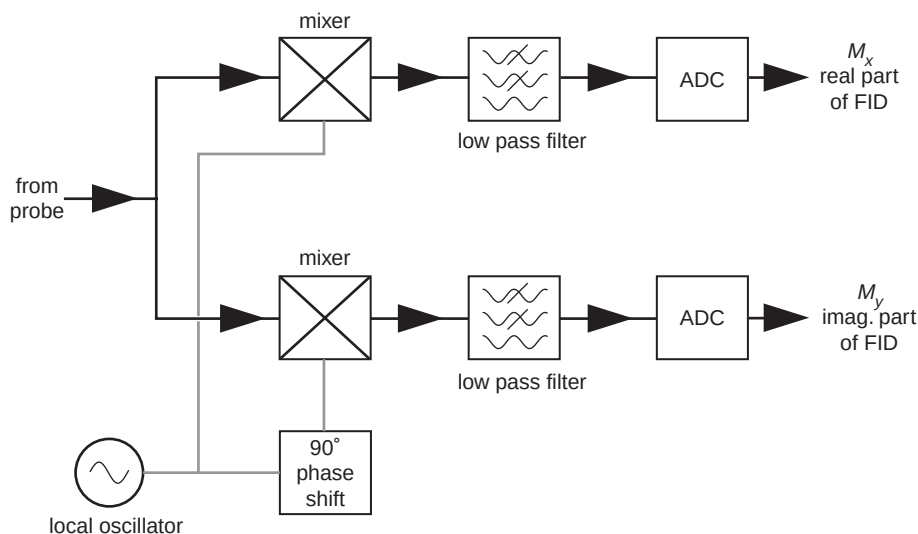


Fig. 5.11 Schematic of the arrangement for quadrature detection; see text for details.

The complete arrangement for quadrature detection is outlined in Fig. 5.11. The NMR signal from the probe is split into two and fed to two separate mixers. The local oscillator signals fed to the two mixers are phase shifted by 90° relative to one another; as a result, the outputs of the two mixers are the equivalents of the x and y components in the rotating frame. These two outputs are digitized separately and become the real and imaginary parts of a complex time domain signal (as described in section 4.1). Such a quadrature detection scheme is the norm for all modern NMR spectrometers.

5.7 Time and frequency

We are now in a position to tidy up a number of points of terminology relating to the FID and spectrum. It is important to understand these in order to report correctly the parameters used to record a spectrum.

The range of frequencies correctly represented in the spectrum (i.e. not folded) is called the *spectral width*, f_{sw} . As quadrature detection is used, the frequency scale on the spectrum runs from $-\frac{1}{2}f_{\text{sw}}$ to $+\frac{1}{2}f_{\text{sw}}$ i.e. positive and negative frequencies are discriminated.

For a given spectral width, the sampling interval or dwell time, Δ , is computed from:

$$\Delta = \frac{1}{f_{\text{SW}}}$$

which means that the Nyquist frequency is $\frac{1}{2}f_{\text{SW}}$ – the edges of the spectrum.

If we record the FID for a time t_{acq} , called the acquisition time, then the number of data points, N , is

$$N = \frac{t_{\text{acq}}}{\Delta}.$$

Each data point is complex, consisting of a real and an imaginary part.

To be clear about the conditions under which an FID is recorded it is important to quote the acquisition time and the spectral width; quoting the number of data points on its own is meaningless.

5.8 The pulse programmer

The pulse programmer has become an immensely sophisticated piece of computer hardware, controlling as it does all of the functions of the spectrometer. As the pulse programmer needs to produce very precisely timed events, often in rapid succession, it is usual for it to run independently of the main computer. Typically, the pulse program is specified in the main computer and then, when the experiment is started, the instructions are loaded into the pulse programmer and then executed there.

The acquisition of data is usually handled by the pulse programmer, again separately from the main computer.

5.9 Exercises

E 5-1

You have been offered a superconducting magnet which claims to have a homogeneity of “1 part in 10^8 ”. Your intention is to use it to record phosphorus-31 spectra at Larmor frequency of 180 MHz, and you know that your typical linewidths are likely to be of the order of 25 Hz. Is the magnet sufficiently homogeneous to be of use?

E 5-2

A careful pulse calibration experiment determines that the 180° pulse is $24.8 \mu\text{s}$. How much attenuation, in dB, would have to be introduced into the transmitter in order to give a field strength, $(\omega_1/2\pi)$, of 2 kHz?

E 5-3

A spectrometer is equipped with a transmitter capable of generating a maximum of 100 W of RF power at the frequency of carbon-13. Using this transmitter at full power the 90° pulse width is found to be $20 \mu\text{s}$. What power would be needed to reduce the 90° pulse width to $7.5 \mu\text{s}$? Would you have any reservations about using this amount of power?

E 5-4

Explain what is meant by “a two bit ADC” and draw a diagram to illustrate the outcome of such a ADC being used to digitize a sine wave.

Why is it generally desirable to improve the number of bits that the ADC uses?

E 5-5

A spectrometer operates at 800 MHz for proton and it is desired to cover a shift range of 15 ppm. Assuming that the receiver frequency is placed in the middle of this range, what spectral width would be needed and what would the sampling interval (dwell time) have to be?

If we recorded a FID for 2 s with the spectral width set as you have determined, how many data points will have been collected?

E 5-6

Suppose that the spectral width is set to 100 Hz, but that a peak is present whose offset from the receiver is +60 Hz. Where will the peak appear in the spectrum? Can you explain why?

6 Product Operators[†]

The vector model, introduced in Chapter 3, is very useful for describing basic NMR experiments but unfortunately is not applicable to coupled spin systems. When it comes to two-dimensional NMR many of the experiments are only of interest in coupled spin systems, so we really must have some way of describing the behaviour of such systems under multiple-pulse experiments.

The tools we need are provided by quantum mechanics, specifically in the form of density matrix theory which is the best way to formulate quantum mechanics for NMR. However, we do not want to get involved in a great deal of complex quantum mechanics! Luckily, there is a way of proceeding which we can use without a deep knowledge of quantum mechanics: this is the *product operator* formalism.

The product operator formalism is a complete and rigorous quantum mechanical description of NMR experiments and is well suited to calculating the outcome of modern multiple-pulse experiments. One particularly appealing feature is the fact that the operators have a clear physical meaning and that the effects of pulses and delays can be thought of as geometrical rotations, much in the same way as we did for the vector model in Chapter 3.

6.1 A quick review of quantum mechanics

In this section we will review a few key concepts before moving on to a description of the product operator formalism.

In quantum mechanics, two mathematical objects – wavefunctions and operators – are of central importance. The wavefunction describes the system of interest (such as a spin or an electron) completely; if the wavefunction is known it is possible to calculate all the properties of the system. The simplest example of this that is frequently encountered is when considering the wavefunctions which describe electrons in atoms (atomic orbitals) or molecules (molecular orbitals). One often used interpretation of such electronic wavefunctions is to say that the square of the wavefunction gives the probability of finding the electron at that point.

Wavefunctions are simply mathematical functions of position, time *etc.* For example, the 1s electron in a hydrogen atom is described by the function $\exp(-ar)$, where r is the distance from the nucleus and a is a constant.

In quantum mechanics, operators represent "observable quantities" such as position, momentum and energy; each observable has an operator associated with it.

Operators "operate on" functions to give new functions, hence their name

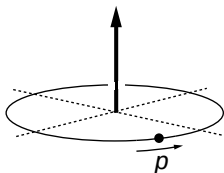
$$\text{operator} \times \text{function} = (\text{new function})$$

An example of an operator is (d/dx) ; in words this operator says "differentiate with respect to x ". Its effect on the function $\sin x$ is

[†] Chapter 6 "Product Operators" © James Keeler, 1998, 2002 & 2004

$$\frac{d}{dx}(\sin x) = \cos x$$

the "new function" is $\cos x$. Operators can also be simple functions, so for example the operator x^2 just means "multiply by x^2 ".



A mass going round a circular path possesses angular momentum, represented by a vector which points perpendicular to the plane of rotation.

6.1.1 Spin operators

A mass going round a circular path (an orbit) possesses *angular momentum*; it turns out that this is a vector quantity which points in a direction perpendicular to the plane of the rotation. The x -, y - and z -components of this vector can be specified, and these are the angular momenta in the x -, y - and z -directions. In quantum mechanics, there are operators which represent these three components of the angular momentum.

Nuclear spins also have angular momentum associated with them – called *spin angular momentum*. The three components of this spin angular momentum (along x , y and z) are represented by the operators I_x , I_y and I_z .

6.1.2 Hamiltonians

The Hamiltonian, H , is the special name given to the operator for the *energy* of the system. This operator is exceptionally important as its eigenvalues and eigenfunctions are the "energy levels" of the system, and it is transitions between these energy levels which are detected in spectroscopy. To understand the spectrum, therefore, it is necessary to have a knowledge of the energy levels and this in turn requires a knowledge of the Hamiltonian operator.

In NMR, the Hamiltonian is seen as having a more subtle effect than simply determining the energy levels. This comes about because the Hamiltonian also affects how the spin system evolves in time. By altering the Hamiltonian the time evolution of the spins can be manipulated and it is precisely this that lies at the heart of multiple-pulse NMR.

The precise mathematical form of the Hamiltonian is found by first writing down an expression for the energy of the system using classical mechanics and then "translating" this into quantum mechanical form according to a set of rules. In this chapter the form of the relevant Hamiltonians will simply be stated rather than derived.

In NMR the Hamiltonian changes depending on the experimental situation. There is one Hamiltonian for the spin or spins in the presence of the applied magnetic field, but this Hamiltonian changes when a radio-frequency pulse is applied.

6.2 Operators for one spin

6.2.1 Operators

In quantum mechanics operators represent observable quantities, such as energy, angular momentum and magnetization. For a single spin-half, the x -, y - and z -components of the magnetization are represented by the spin angular momentum operators I_x , I_y and I_z respectively. Thus at any time the state of the spin system, in quantum mechanics the density operator, σ , can be

represented as a sum of different amounts of these three operators

$$\sigma(t) = a(t)I_x + b(t)I_y + c(t)I_z$$

The amounts of the three operators will vary with time during pulses and delays. This expression of the density operator as a combination of the spin angular momentum operators is exactly analogous to specifying the three components of a magnetization vector.

At equilibrium the density operator is proportional to I_z (there is only z-magnetization present). The constant of proportionality is usually unimportant, so it is usual to write $\sigma_{\text{eq}} = I_z$

6.1.2 Hamiltonians for pulses and delays

In order to work out how the density operator varies with time we need to know the Hamiltonian (which is also an operator) which is acting during that time.

The free precession Hamiltonian (*i.e.* that for a delay), H_{free} , is

$$H_{\text{free}} = \Omega I_z$$

In the vector model free precession involves a rotation at frequency Ω about the z-axis; in the quantum mechanical picture the Hamiltonian involves the z-angular momentum operator, I_z – there is a direct correspondence.

The Hamiltonian for a pulse about the x-axis, H_{pulse} , is

$$H_{\text{pulse},x} = \omega_1 I_x$$

and for a pulse about the y-axis it is

$$H_{\text{pulse},y} = \omega_1 I_y$$

Again there is a clear connection to the vector model where pulses result in rotations about the x- or y-axes.

6.1.3 Equation of motion

The density operator at time t , $\sigma(t)$, is computed from that at time 0, $\sigma(0)$, using the following relationship

$$\sigma(t) = \exp(-iHt) \sigma(0) \exp(iHt)$$

where H is the relevant hamiltonian. If H and σ are expressed in terms of the angular momentum operators it turns out that this equation can be solved easily with the aid of a few rules.

Suppose that an x-pulse, of duration t_p , is applied to equilibrium magnetization. In this situation $H = \omega_1 I_x$ and $\sigma(0) = I_z$; the equation to be solved is

$$\sigma(t_p) = \exp(-i\omega_1 t_p I_x) I_z \exp(i\omega_1 t_p I_x)$$

Such equations involving angular momentum operators are common in quantum mechanics and the solution to them are already all know. The identity required here to solve this equation is

$$\exp(-i\beta I_x) I_z \exp(i\beta I_x) \equiv \cos \beta I_z - \sin \beta I_y \quad [6.1]$$

This is interpreted as a *rotation* of I_z by an angle β about the x-axis. By

putting $\beta = \omega_1 t_p$ this identity can be used to solve Eqn. [6.1]

$$\sigma(t_p) = \cos \omega_1 t_p I_z - \sin \omega_1 t_p I_y$$

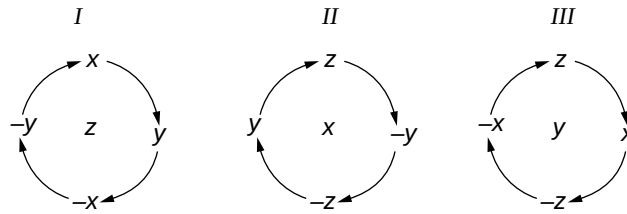
The result is exactly as expected from the vector model: a pulse about the x -axis rotates z -magnetization towards the $-y$ -axis, with a sinusoidal dependence on the flip angle, β .

6.1.4 Standard rotations

Given that there are only three operators, there are a limited number of identities of the type of Eqn. [6.1]. They all have the same form

$$\begin{aligned} \exp(-i\theta I_a) \{ \text{old operator} \} \exp(i\theta I_a) \\ \equiv \cos \theta \{ \text{old operator} \} + \sin \theta \{ \text{new operator} \} \end{aligned}$$

where $\{ \text{old operator} \}$, $\{ \text{new operator} \}$ and I_a are determined from the three possible angular momentum operators according to the following diagrams; the label in the centre indicates which axis the rotation is about



Angle of rotation = Ωt for offsets and $\omega_1 t_p$ for pulses

First example: find the result of rotating the operator I_y by θ about the x -axis, that is

$$\exp(-i\theta I_x) I_y \exp(i\theta I_x)$$

For rotations about x the middle diagram *II* is required. The diagram shows that I_y (the "old operator") is rotated to I_z (the "new operator"). The required identity is therefore

$$\exp(-i\theta I_x) I_y \exp(i\theta I_x) \equiv \cos \theta I_y + \sin \theta I_z$$

Second example: find the result of

$$\exp(-i\theta I_y) \{-I_z\} \exp(i\theta I_y)$$

This is a rotation about y , so diagram *III* is required. The diagram shows that $-I_z$ (the "old operator") is rotated to $-I_x$ (the "new operator"). The required identity is therefore

$$\begin{aligned} \exp(-i\theta I_y) \{-I_z\} \exp(i\theta I_y) &\equiv \cos \theta \{-I_z\} + \sin \theta \{-I_x\} \\ &\equiv -\cos \theta I_z - \sin \theta I_x \end{aligned}$$

Finally, note that a rotation of an operator about its own axis has no effect *e.g.* a rotation of I_x about x leaves I_x unaltered.

6.1.5 Shorthand notation

To save writing, the arrow notation is often used. In this, the term Ht is written over an arrow which connects the old and new density operators. So, for example, the following

$$\sigma(t_p) = \exp(-i\omega_1 t_p I_x) \sigma(0) \exp(i\omega_1 t_p I_x)$$

is written

$$\sigma(0) \xrightarrow{\omega_1 t_p I_x} \sigma(t_p)$$

For the case where $\sigma(0) = I_z$

$$I_z \xrightarrow{\omega_1 t_p I_x} \cos \omega_1 t_p I_z - \sin \omega_1 t_p I_y$$

6.1.6 Example calculation: spin echo

$$90^\circ(x) \xrightarrow{a} \text{delay } \tau \xrightarrow{b} 180^\circ(x) \xrightarrow{e} \text{delay } \tau \xrightarrow{f} \text{acquire}$$

At a the density operator is $-I_y$. The transformation from a to b is free precession, for which the Hamiltonian is ΩI_z ; the delay τ therefore corresponds to a rotation about the z -axis at frequency Ω . In the short-hand notation this is

$$-I_y \xrightarrow{\Omega \tau I_z} \sigma(b)$$

To solve this diagram I above is needed with the angle $= \Omega \tau$, the "new operator" is I_x

$$-I_y \xrightarrow{\Omega \tau I_z} -\cos \Omega \tau I_y + \sin \Omega \tau I_x$$

In words this says that the magnetization precesses from $-y$ towards $+x$.

The pulse about x has the Hamiltonian $\omega_1 I_x$; the pulse therefore corresponds to a rotation about x for a time t_p such that the angle, $\omega_1 t_p$, is π radians. In the shorthand notation

$$-\cos \Omega \tau I_y + \sin \Omega \tau I_x \xrightarrow{\omega_1 t_p I_x} \sigma(e) \quad [6.2]$$

Each term on the left is dealt with separately. The first term is a rotation of y about x ; the relevant diagram is thus II

$$-\cos \Omega \tau I_y \xrightarrow{\omega_1 t_p I_x} -\cos \Omega \tau \cos \omega_1 t_p I_y - \cos \Omega \tau \sin \omega_1 t_p I_z$$

However, the flip angle of the pulse, $\omega_1 t_p$, is π so the second term on the right is zero and the first term just changes sign ($\cos \pi = -1$); overall the result is

$$-\cos \Omega \tau I_y \xrightarrow{\pi I_x} \cos \Omega \tau I_y$$

The second term on the left of Eqn. [6.2] is easy to handle as it is unaffected by a rotation about x . Overall, the effect of the 180° pulse is then

$$-\cos \Omega \tau I_y + \sin \Omega \tau I_x \xrightarrow{\pi I_x} \cos \Omega \tau I_y + \sin \Omega \tau I_x \quad [6.3]$$

As was shown using the vector model, the y -component just changes sign. The next stage is the evolution of the offset for time τ . Again, each term on the right of Eqn. [6.3] is considered separately

$$\cos \Omega \tau I_y \xrightarrow{\Omega \tau I_z} \cos \Omega \tau \cos \Omega \tau I_y - \sin \Omega \tau \cos \Omega \tau I_x$$

$$\sin \Omega \tau I_x \xrightarrow{\Omega \tau I_z} \cos \Omega \tau \sin \Omega \tau I_x + \sin \Omega \tau \sin \Omega \tau I_y$$

Collecting together the terms in I_x and I_y the final result is

$$(\cos \Omega \tau \cos \Omega \tau + \sin \Omega \tau \sin \Omega \tau) I_y + (\cos \Omega \tau \sin \Omega \tau - \sin \Omega \tau \cos \Omega \tau) I_x$$

The bracket multiplying I_x is zero and the bracket multiplying I_y is $=1$ because of the identity $\cos^2 \theta + \sin^2 \theta = 1$. Thus the overall result of the spin echo sequence can be summarised

$$I_z \xrightarrow{90^\circ(x) - \tau - 180^\circ(x) - \tau} I_y$$

In words, the outcome is independent of the offset, Ω , and the delay τ , even though there is evolution during the delays. The offset is said to be *refocused* by the spin echo. This is exactly the result we found in section 3.8.

In general the sequence

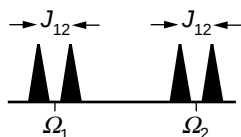
$$- \tau - 180^\circ(x) - \tau - \quad [6.4]$$

refocuses any evolution due to offsets; this is a very useful feature which is much used in multiple-pulse NMR experiments.

One further point is that as far as the offset is concerned the spin echo sequence of Eqn. [6.4] is just equivalent to $180^\circ(x)$.

6.2 Operators for two spins

The product operator approach comes into its own when coupled spin systems are considered; such systems cannot be treated by the vector model. However, product operators provide a clean and simple description of the important phenomena of coherence transfer and multiple quantum coherence.



The spectrum from two coupled spins, with offsets Ω_1 and Ω_2 (rad s^{-1}) and mutual coupling J_{12} (Hz).

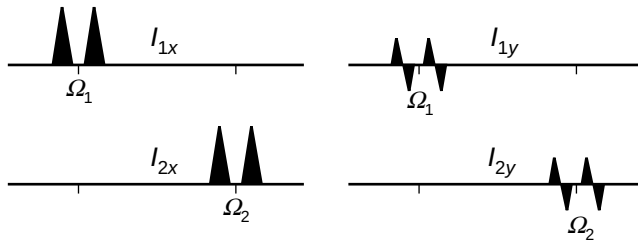
6.1.1 Product operators for two spins

For a single spin the three operators needed for a complete description are I_x , I_y and I_z . For two spins, three such operators are needed for each spin; an additional subscript, 1 or 2, indicates which spin they refer to.

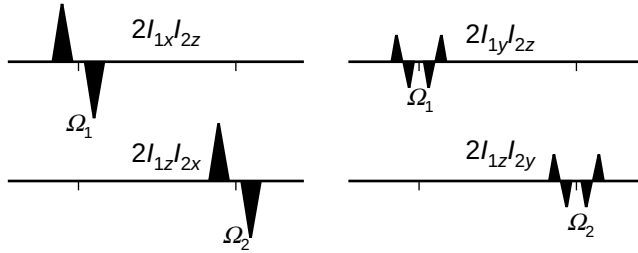
$$\text{spin 1: } I_{1x} \quad I_{1y} \quad I_{1z} \quad \text{spin 2: } I_{2x} \quad I_{2y} \quad I_{2z}$$

I_{1z} represents z-magnetization of spin 1, and I_{2z} likewise for spin 2. I_{1x} represents x-magnetization on spin 1. As spin 1 and 2 are coupled, the spectrum consists of two doublets and the operator I_{1x} can be further identified with the two lines of the spin-1 doublet. In the language of product operators I_{1x} is said to represent *in-phase* magnetization of spin 1; the description in-phase means that the two lines of the spin 1 doublet have the same sign and lineshape.

Following on in the same way I_{2x} represents in-phase magnetization on spin 2. I_{1y} and I_{2y} also represent in-phase magnetization on spins 1 and 2, respectively, but this magnetization is aligned along y and so will give rise to a different lineshape. Arbitrarily, an absorption mode lineshape will be assigned to magnetization aligned along x and a dispersion mode lineshape to magnetization along y.



There are four additional operators which represent *anti-phase* magnetization: $2I_{1x}I_{2z}$, $2I_{1y}I_{2z}$, $2I_{1z}I_{2x}$, $2I_{1z}I_{2y}$ (the factors of 2 are needed for normalization purposes). The operator $2I_{1x}I_{2z}$ is described as magnetization on spin 1 which is anti-phase with respect to the coupling to spin 2.



Note that the two lines of the spin-1 multiplet are associated with different spin states of spin-2, and that in an anti-phase multiplet these two lines have different signs. Anti-phase terms are thus sensitive to the spin states of the coupled spins.

There are four remaining product operators which contain two transverse (*i.e.* x- or y-operators) terms and correspond to multiple-quantum coherences; they are not observable

$$\text{multiple quantum : } 2I_{1x}I_{2y} \quad 2I_{1y}I_{2x} \quad 2I_{1x}I_{2x} \quad 2I_{1y}I_{2y}$$

Finally there is the term $2I_{1z}I_{2z}$ which is also not observable and corresponds to a particular kind of non-equilibrium population distribution.

6.1.2 Evolution under offsets and pulses

The operators for two spins evolve under offsets and pulses in the same way as do those for a single spin. The rotations have to be applied separately to each spin and it must be remembered that rotations of spin 1 do not affect spin 2, and *vice versa*.

For example, consider I_{1x} evolving under the offset of spin 1 and spin 2. The relevant Hamiltonian is

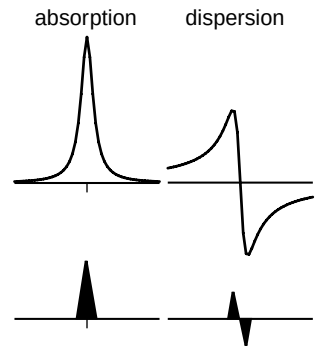
$$H_{\text{free}} = \Omega_1 I_{1z} + \Omega_2 I_{2z}$$

where Ω_1 and Ω_2 are the offsets of spin 1 and spin 2 respectively. Evolution under this Hamiltonian can be considered by applying the two terms sequentially (the order is immaterial)

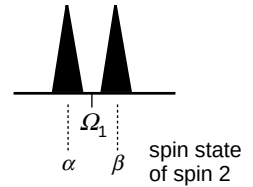
$$\begin{aligned} I_{1x} &\xrightarrow{H_{\text{free}}t} \\ I_{1x} &\xrightarrow{\Omega_1 t I_{1z} + \Omega_2 t I_{2z}} \\ I_{1x} &\xrightarrow{\Omega_1 t I_{1z}} \xrightarrow{\Omega_2 t I_{2z}} \end{aligned}$$

The first "arrow" is a rotation about z

$$I_{1x} \xrightarrow{\Omega_1 t I_{1z}} \cos \Omega_1 t I_{1x} + \sin \Omega_1 t I_{1y} \xrightarrow{\Omega_2 t I_{2z}}$$



The absorption and dispersion lineshapes. The absorption lineshape is a maximum on resonance, whereas the dispersion goes through zero at this point. The "cartoon" forms of the lineshapes are shown in the lower part of the diagram.



The two lines of the spin-1 doublet can be associated with different spin states of spin 2.

The second arrow leaves the intermediate state unaltered as spin-2 operators have not effect on spin-1 operators. Overall, therefore

$$I_{1x} \xrightarrow{\Omega_1 t_{1z} + \Omega_2 t_{2z}} \cos \Omega_1 t I_{1x} + \sin \Omega_1 t I_{1y}$$

A second example is the term $2I_{1x}I_{2z}$ evolving under a 90° pulse about the y-axis applied to both spins. The relevant Hamiltonian is

$$H = \omega_1 I_{1y} + \omega_2 I_{2y}$$

The evolution can be separated into two successive rotations

$$2I_{1x}I_{2z} \xrightarrow{\omega_1 t_{1y}} \xrightarrow{\omega_2 t_{2y}}$$

The first arrow affects only the spin-1 operators; a 90° rotation of I_{1x} about y gives $-I_{1z}$ (remembering that $\omega_1 t = \pi/2$ for a 90° pulse)

$$2I_{1x}I_{2z} \xrightarrow{\omega_1 t_{1y}} \cos \omega_1 t 2I_{1x}I_{2z} - \sin \omega_1 t 2I_{1z}I_{2z} \xrightarrow{\omega_2 t_{2y}} \\ 2I_{1x}I_{2z} \xrightarrow{\pi/2 I_{1y}} -2I_{1z}I_{2z} \xrightarrow{\pi/2 I_{2y}}$$

The second arrow only affects the spin 2 operators; a 90° rotation of z about y takes it to x

$$2I_{1x}I_{2z} \xrightarrow{\pi/2 I_{1y}} -2I_{1z}I_{2z} \xrightarrow{\pi/2 I_{2y}} -2I_{1z}I_{2x}$$

The overall result is that anti-phase magnetization of spin 1 has been transferred into anti-phase magnetization of spin 2. Such a process is called *coherence transfer* and is exceptionally important in multiple-pulse NMR.

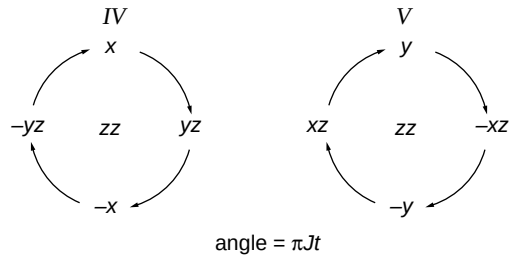
6.1.3 Evolution under coupling

The new feature which arises when considering two spins is the effect of coupling between them. The Hamiltonian representing this coupling is itself a product of two operators:

$$H_J = 2\pi J_{12} I_{1z}I_{2z}$$

where J_{12} is the coupling in Hz.

Evolution under coupling causes the interconversion of in-phase and anti-phase magnetization according to the following diagrams



For example, in-phase magnetization along x becomes anti-phase along y according to the diagram IV

$$I_{1x} \xrightarrow{2\pi J_{12} t I_{1z}I_{2z}} \cos \pi J_{12} t I_{1x} + \sin \pi J_{12} t 2I_{1y}I_{2z}$$

note that the angle is $\pi J_{12} t$ i.e. half the angle for the other rotations, I–III.

Anti-phase magnetization along x becomes in-phase magnetization along y; using diagram V:

$$2I_{1x}I_{2z} \xrightarrow{2\pi J_{12} t I_{1z}I_{2z}} \cos \pi J_{12} t 2I_{1x}I_{2z} + \sin \pi J_{12} t I_{1y}$$

The diagrams apply equally well to spin-2; for example

$$-2I_{1z}I_{2y} \xrightarrow{2\pi J_{12}t I_{1z}I_{2z}} -\cos \pi J_{12}t 2I_{1z}I_{2y} + \sin \pi J_{12}t I_{2x}$$

Complete interconversion of in-phase and anti-phase magnetization requires a delay such that $\pi J_{12}t = \pi/2$ i.e. a delay of $1/(2J_{12})$. A delay of $1/J_{12}$ causes in-phase magnetization to change its sign:

$$I_{1x} \xrightarrow{2\pi J_{12}t I_{1z}I_{2z} \quad t=1/2J_{12}} 2I_{1y}I_{2z} \quad I_{2y} \xrightarrow{2\pi J_{12}t I_{1z}I_{2z} \quad t=1/J_{12}} -I_{2y}$$

6.3 Spin echoes

It was shown in section 6.2.6 that the offset is refocused in a spin echo. In this section it will be shown that the evolution of the scalar coupling is not necessarily refocused.

6.3.1 Spin echoes in homonuclear spin system

In this kind of spin echo the 180° pulse affects both spins i.e. it is a non-selective pulse:

$$-\tau - 180^\circ(x, \text{ to spin 1 and spin 2}) - \tau -$$

At the start of the sequence it will be assumed that only in-phase x-magnetization on spin 1 is present: I_{1x} . In fact the starting state is not important to the overall effect of the spin echo, so this choice is arbitrary.

It was shown in section 6.2.6 that the spin echo applied to one spin refocuses the offset; this conclusion is not altered by the presence of a coupling so the offset will be ignored in the present calculation. This greatly simplifies things.

For the first delay τ only the effect of evolution under coupling need be considered therefore:

$$I_{1x} \xrightarrow{2\pi J_{12}\tau I_{1z}I_{2z}} \cos \pi J_{12}\tau I_{1x} + \sin \pi J_{12}\tau 2I_{1y}I_{2z}$$

The 180° pulse affects both spins, and this can be calculated by applying the 180° rotation to each in succession

$$\cos \pi J_{12}\tau I_{1x} + \sin \pi J_{12}\tau 2I_{1y}I_{2z} \xrightarrow{\pi I_{1x}} \xrightarrow{\pi I_{2x}}$$

where it has already been written in that $\omega_1 t_p = \pi$, for a 180° pulse. The 180° rotation about x for spin 1 has no effect on the operator I_{1x} and I_{2z} , and it simply reverses the sign of the operator I_{1y}

$$\cos \pi J_{12}\tau I_{1x} + \sin \pi J_{12}\tau 2I_{1y}I_{2z} \xrightarrow{\pi I_{1x}} \cos \pi J_{12}\tau I_{1x} - \sin \pi J_{12}\tau 2I_{1y}I_{2z} \xrightarrow{\pi I_{2x}}$$

The 180° rotation about x for spin 2 has no effect on the operators I_{1x} and I_{1y} , but simply reverses the sign of the operator I_{2z} . The final result is thus

$$\begin{aligned} \cos \pi J_{12}\tau I_{1x} + \sin \pi J_{12}\tau 2I_{1y}I_{2z} &\xrightarrow{\pi I_{1x}} \cos \pi J_{12}\tau I_{1x} - \sin \pi J_{12}\tau 2I_{1y}I_{2z} \\ &\xrightarrow{\pi I_{2x}} \cos \pi J_{12}\tau I_{1x} + \sin \pi J_{12}\tau 2I_{1y}I_{2z} \end{aligned}$$

Nothing has happened; the 180° pulse has left the operators unaffected! So, for the purposes of the calculation it is permissible to ignore the 180° pulse and simply allow the coupling to evolve for 2τ . The final result can therefore just be written down:

$$I_{1x} \xrightarrow{\tau - 180^\circ(x) - \tau} \cos 2\pi J_{12} \tau I_{1x} + \sin 2\pi J_{12} \tau 2I_{1y}I_{2z}$$

From this it is easy to see that complete conversion to anti-phase magnetization requires $2\pi J_{12} \tau = \pi/2$ i.e. $\tau = 1/(4 J_{12})$.

The calculation is not quite as simple if the initial state is chosen as I_{1y} , but the final result is just the same – the coupling evolves for 2τ :

$$I_{1y} \xrightarrow{\tau - 180^\circ(x) - \tau} -\cos 2\pi J_{12} \tau I_{1y} + \sin 2\pi J_{12} \tau 2I_{1x}I_{2z}$$

In fact, the general result is that the sequence

$$-\tau - 180^\circ(x, \text{ to spin 1 and spin 2}) - \tau -$$

is equivalent to the sequence

$$-2\tau - 180^\circ(x, \text{ to spin 1 and spin 2})$$

in which the offset is ignored and coupling is allowed to act for time 2τ .

6.1.2 Interconverting in-phase and anti-phase states

So far, spin echoes have been demonstrated as being useful for generating anti-phase terms, independent of offsets. For example, the sequence

$$90^\circ(x) - 1/(4J_{12}) - 180^\circ(x) - 1/(4J_{12}) -$$

generates pure anti-phase magnetization.

Equally useful is the sequence

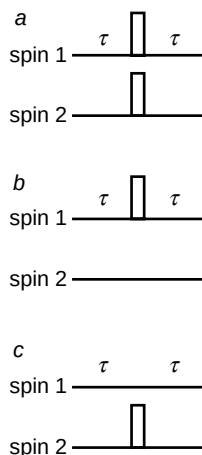
$$-1/(4J_{12}) - 180^\circ(x) - 1/(4J_{12}) -$$

which will convert pure anti-phase magnetization, such as $2I_{1x}I_{2z}$ into in-phase magnetization, I_{1y} .

6.1.3 Spin echoes in heteronuclear spin systems

If spin 1 and spin 2 are different nuclear species, such as ^{13}C and ^1H , it is possible to choose to apply the 180° pulse to either or both spins; the outcome of the sequence depends on the pattern of 180° pulses.

Sequence *a* has already been analysed: the result is that the offset is refocused but that the coupling evolves for time 2τ . Sequence *b* still refocuses the offset of spin 1, but it turns out that the coupling is also refocused. Sequence *c* refocuses the coupling but leaves the evolution of the offset unaffected.



Three different spin echo sequences that can be applied to heteronuclear spin systems. The open rectangles represent 180° pulses.

Sequence b

It will be assumed that the offset is refocused, and attention will therefore be restricted to the effect of the coupling

$$I_{1x} \xrightarrow{2\pi J_{12} \tau I_{1x}I_{2z}} \cos \pi J_{12} \tau I_{1x} + \sin \pi J_{12} \tau 2I_{1y}I_{2z}$$

The $180^\circ(x)$ pulse is only applied to spin 1

$$\cos \pi J_{12} \tau I_{1x} + \sin \pi J_{12} \tau 2I_{1y}I_{2z} \xrightarrow{\pi I_{1x}} \cos \pi J_{12} \tau I_{1x} - \sin \pi J_{12} \tau 2I_{1y}I_{2z} \quad [6.5]$$

The two terms on the right each evolve under the coupling during the second delay:

$$\begin{aligned}
& \cos \pi J_{12} \tau I_{1x} \xrightarrow{2\pi J_{12} \tau I_{1z} I_{2z}} \\
& \cos \pi J_{12} \tau \cos \pi J_{12} \tau I_{1x} + \sin \pi J_{12} \tau \cos \pi J_{12} \tau 2I_{1y} I_{2z} \\
& - \sin \pi J_{12} \tau 2I_{1y} I_{2z} \xrightarrow{2\pi J_{12} \tau I_{1z} I_{2z}} \\
& - \cos \pi J_{12} \tau \sin \pi J_{12} \tau 2I_{1y} I_{2z} + \sin \pi J_{12} \tau \sin \pi J_{12} \tau I_{1x}
\end{aligned}$$

Collecting the terms together and noting that $\cos^2 \theta + \sin^2 \theta = 1$ the final result is just I_{1x} . In words, the effect of the coupling has been refocused.

Sequence c

As there is no 180° pulse applied to spin 1, the offset of spin 1 is not refocused, but continues to evolve for time 2τ . The evolution of the coupling is easy to calculate:

$$I_{1x} \xrightarrow{2\pi J_{12} \tau I_{1z} I_{2z}} \cos \pi J_{12} \tau I_{1x} + \sin \pi J_{12} \tau 2I_{1y} I_{2z}$$

This time the $180^\circ(x)$ pulse is applied to spin 2

$$\cos \pi J_{12} \tau I_{1x} + \sin \pi J_{12} \tau 2I_{1y} I_{2z} \xrightarrow{\pi I_{2x}} \cos \pi J_{12} \tau I_{1x} - \sin \pi J_{12} \tau 2I_{1y} I_{2z}$$

The results is exactly as for sequence *b* (Eqn. [6.5]), so the final result is the same *i.e.* the coupling is refocused.

Summary

In heteronuclear systems it is possible to choose whether or not to allow the offset and the coupling to evolve; this gives great freedom in generating and manipulating anti-phase states which play a key role in multiple pulse NMR experiments.

6.4 Multiple quantum terms

6.4.1 Coherence order

In NMR the directly observable quantity is the transverse magnetization, which in product operators is represented by terms such as I_{1x} and $2I_{1z}I_{2y}$. Such terms are examples of single quantum coherences, or more generally coherences with order, p , $= \pm 1$. Other product operators can also be classified according to coherence order *e.g.* $2I_{1z}I_{2z}$ has $p = 0$ and $2I_{1x}I_{2y}$ has both $p = 0$ (zero-quantum coherence) and ± 2 (double quantum coherence). Only single quantum coherences are observable.

In heteronuclear systems it is sometimes useful to classify operators according to their coherence orders with respect to each spin. So, for example, $2I_{1z}I_{2y}$ has $p = 0$ for spin 1 and $p = \pm 1$ for spin 2.

6.4.2 Raising and lowering operators

The classification of operators according to coherence order is best carried out by re-expressing the Cartesian operators I_x and I_y in terms of the raising and lowering operators, I_+ and I_- , respectively. These are defined as follows

$$I_+ = I_x + iI_y \quad I_- = I_x - iI_y \quad [6.6]$$

where i is the square root of -1 . I_+ has coherence order $+1$ and I_- has coherence order -1 ; coherence order is a *signed* quantity.

Using the definitions of Eqn. [6.6] I_x and I_y can be expressed in terms of the raising and lowering operators

$$I_x = \frac{1}{2}(I_+ + I_-) \quad I_y = \frac{1}{2i}(I_+ - I_-) \quad [6.7]$$

from which it is seen that I_x and I_y are both mixtures of coherences with $p = +1$ and -1 .

The operator product $2I_{1x}I_{2x}$ can be expressed in terms of the raising and lowering operators in the following way (note that separate operators are used for each spin: $I_{1\pm}$ and $I_{2\pm}$)

$$\begin{aligned} 2I_{1x}I_{2x} &= 2 \times \frac{1}{2}(I_{1+} + I_{1-}) \times \frac{1}{2}(I_{2+} + I_{2-}) \\ &= \frac{1}{2}(I_{1+}I_{2+} + I_{1-}I_{2-}) + \frac{1}{2}(I_{1+}I_{2-} + I_{1-}I_{2+}) \end{aligned} \quad [6.8]$$

The first term on the right of Eqn. [6.8] has $p = (+1+1) = 2$ and the second term has $p = (-1-1) = -2$; both are double quantum coherences. The third and fourth terms both have $p = (+1-1) = 0$ and are zero quantum coherences. The value of p can be found simply by noting the number of raising and lowering operators in the product.

The pure double quantum part of $2I_{1x}I_{2x}$ is, from Eqn. [6.8],

$$\text{double quantum part}[2I_{1x}I_{2x}] = \frac{1}{2}(I_{1+}I_{2+} + I_{1-}I_{2-}) \quad [6.9]$$

The raising and lowering operators on the right of Eqn. [6.9] can be re-expressed in terms of the Cartesian operators:

$$\begin{aligned} \frac{1}{2}(I_{1+}I_{2+} + I_{1-}I_{2-}) &= \frac{1}{2}[(I_{1x} + iI_{1y})(I_{2x} + iI_{2y}) + (I_{1x} - iI_{1y})(I_{2x} - iI_{2y})] \\ &= \frac{1}{2}[2I_{1x}I_{2x} + 2I_{1y}I_{2y}] \end{aligned}$$

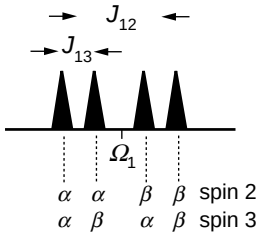
So, the pure double quantum part of $2I_{1x}I_{2x}$ is $\frac{1}{2}(2I_{1x}I_{2x} + 2I_{1y}I_{2y})$; by a similar method the pure zero quantum part can be shown to be $\frac{1}{2}(2I_{1x}I_{2x} - 2I_{1y}I_{2y})$. Some further useful relationships are given in section 6.9

6.5 Three spins

The product operator formalism can be extended to three or more spins. No really new features arise, but some of the key ideas will be highlighted in this section. The description will assume that spin 1 is coupled to spins 2 and 3 with coupling constants J_{12} and J_{13} ; in the diagrams it will be assumed that $J_{12} > J_{13}$.

6.5.1 Types of operators

I_{1x} represents in-phase magnetization on spin 1; $2I_{1x}I_{2z}$ represents magnetization anti-phase with respect to the coupling to spin 2 and $2I_{1x}I_{3z}$ represents magnetization anti-phase with respect to the coupling to spin 3. $4I_{1x}I_{2z}I_{3z}$ represents magnetization which is doubly anti-phase with respect to the couplings to both spins 2 and 3.



The doublet of doublets from spin 1 coupled to two other spins. The spin states of the coupled spins are also indicated.

As in the case of two spins, the presence of more than one transverse operator in the product represents multiple quantum coherence. For example, $2I_{1x}I_{2x}$ is a mixture of double- and zero-quantum coherence between spins 1 and 2. The product $4I_{1x}I_{2x}I_{3z}$ is the same mixture, but anti-phase with respect to the coupling to spin 3. Products such as $4I_{1x}I_{2x}I_{3x}$ contain, amongst other things, triple-quantum coherences.

6.5.2 Evolution

Evolution under offsets and pulses is simply a matter of applying sequentially the relevant rotations for each spin, remembering that rotations of spin 1 do not affect operators of spins 2 and 3. For example, the term $2I_{1x}I_{2z}$ evolves under the offset in the following way:

$$2I_{1x}I_{2z} \xrightarrow{\Omega_1 t I_{1z}} \xrightarrow{\Omega_2 t I_{2z}} \xrightarrow{\Omega_3 t I_{3z}} \cos \Omega_1 t \, 2I_{1x}I_{2z} + \sin \Omega_1 t \, 2I_{1y}I_{2z}$$

The first arrow, representing evolution under the offset of spin 1, affects only the spin 1 operator I_{1x} . The second arrow has no effect as the spin 2 operator I_{2z} and this is unaffected by a z-rotation. The third arrow also has no effect as there are no spin 3 operators present.

The evolution under coupling follows the same rules as for a two-spin system. For example, evolution of I_{1x} under the influence of the coupling to spin 3 generates $2I_{1y}I_{3z}$

$$I_{1x} \xrightarrow{2\pi J_{13} t I_{1z} I_{3z}} \cos \pi J_{13} t \, I_{1x} + \sin \pi J_{13} t \, 2I_{1y}I_{3z}$$

Further evolution of the term $2I_{1y}I_{3z}$ under the influence of the coupling to spin 2 generates a double anti-phase term

$$2I_{1y}I_{3z} \xrightarrow{2\pi J_{12} t I_{1z} I_{2z}} \cos \pi J_{12} t \, 2I_{1y}I_{3z} - \sin \pi J_{13} t \, 4I_{1x}I_{2z}I_{3z}$$

In this evolution the spin 3 operator is unaffected as the coupling does not involve this spin. The connection with the evolution of I_{1y} under a coupling can be made more explicit by writing $2I_{3z}$ as a "constant" γ

$$\gamma I_{1y} \xrightarrow{2\pi J_{12} t I_{1z} I_{2z}} \cos \pi J_{12} t \, \gamma I_{1y} - \sin \pi J_{13} t \, 2\gamma I_{1x}I_{2z}$$

which compares directly to

$$I_{1y} \xrightarrow{2\pi J_{12} t I_{1z} I_{2z}} \cos \pi J_{12} t \, I_{1y} - \sin \pi J_{13} t \, 2I_{1x}I_{2z}$$

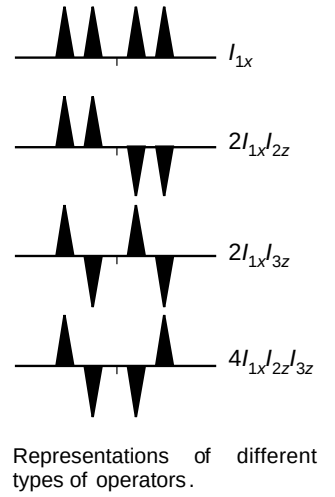
6.6 Alternative notation

In this chapter different spins have been designated with a subscript 1, 2, 3 ... Another common notation is to distinguish the spins by using a different letter to represent their operators; commonly I and S are used for two of the symbols

$$2I_{1x}I_{2z} \equiv 2I_x S_z$$

Note that the order in which the operators are written is not important, although it is often convenient (and tidy) always to write them in the same sequence.

In heteronuclear experiments a notation is sometimes used where the letter represents the nucleus. So, for example, operators referring to protons are given the letter H , carbon-13 atoms the letter C and nitrogen-15 atoms



the letter N ; carbonyl carbons are sometimes denoted C' . For example, $4C_xH_zN_z$ denotes magnetization on carbon-13 which is anti-phase with respect to coupling to both proton and nitrogen-15.

6.7 Conclusion

The product operator method as described here only applies to spin-half nuclei. It can be extended to higher spins, but significant extra complexity is introduced; details can be found in the article by Sørensen *et al.* (*Prog. NMR Spectrosc.* **16**, 163 (1983)).

The main difficulty with the product operator method is that the more pulses and delays that are introduced the greater becomes the number of operators and the more complex the trigonometrical expressions multiplying them. If pulses are either 90° or 180° then there is some simplification as such pulses do not increase the number of terms. As will be seen in chapter 7, it is important to try to simplify the calculation as much as possible, for example by recognizing when offsets or couplings are refocused by spin echoes.

A number of computer programs are available for machine computation using product operators within programs such as *Mathematica*. These can be very labour saving.

6.8 Multiple -quantum coherence

6.8.1 Multiple-quantum terms

In the product operator representation of multiple quantum coherences it is usual to distinguish between *active* and *passive* spins. Active spins contribute transverse operators, such as I_x , I_y and I_+ , to the product; passive spins contribute only z-operators, I_z . In a sense the spins contributing transverse operators are "involved" in the coherence, while those contributing z-operators are simply spectators.

For double- and zero-quantum coherence in which spins i and j are active it is convenient to define the following set of operators which represent pure multiple quantum states of given order. The operators can be expressed in terms of the Cartesian or raising and lowering operators.

double quantum, $p = \pm 2$

$$DQ_x^{(ij)} \equiv \frac{1}{2} (2I_{ix}I_{jx} - 2I_{iy}I_{jy}) \equiv \frac{1}{2} (I_{i+}I_{j+} + I_{i-}I_{j-})$$

$$DQ_y^{(ij)} \equiv \frac{1}{2} (2I_{ix}I_{jy} + 2I_{iy}I_{jx}) \equiv \frac{1}{2i} (I_{i+}I_{j+} - I_{i-}I_{j-})$$

zero quantum, $p = 0$

$$ZQ_x^{(ij)} \equiv \frac{1}{2} (2I_{ix}I_{jx} + 2I_{iy}I_{jy}) \equiv \frac{1}{2} (I_{i+}I_{j-} + I_{i-}I_{j+})$$

$$ZQ_y^{(ij)} \equiv \frac{1}{2} (2I_{iy}I_{jx} - 2I_{ix}I_{jy}) \equiv \frac{1}{2i} (I_{i+}I_{j-} - I_{i-}I_{j+})$$

6.8.2 Evolution of multiple -quantum terms

Evolution under offsets

The double- and zero-quantum operators evolve under offsets in a way

which is entirely analogous to the evolution of I_x and I_y under free precession except that the frequencies of evolution are $(\Omega_i + \Omega_j)$ and $(\Omega_i - \Omega_j)$ respectively:

$$\begin{aligned} \text{DQ}_x^{(ij)} &\xrightarrow{\Omega_i t I_{iz} + \Omega_j t I_{jz}} \cos(\Omega_i + \Omega_j)t \text{DQ}_x^{(ij)} + \sin(\Omega_i + \Omega_j)t \text{DQ}_y^{(ij)} \\ \text{DQ}_y^{(ij)} &\xrightarrow{\Omega_i t I_{iz} + \Omega_j t I_{jz}} \cos(\Omega_i + \Omega_j)t \text{DQ}_y^{(ij)} - \sin(\Omega_i + \Omega_j)t \text{DQ}_x^{(ij)} \\ \text{ZQ}_x^{(ij)} &\xrightarrow{\Omega_i t I_{iz} + \Omega_j t I_{jz}} \cos(\Omega_i - \Omega_j)t \text{ZQ}_x^{(ij)} + \sin(\Omega_i - \Omega_j)t \text{ZQ}_y^{(ij)} \\ \text{ZQ}_y^{(ij)} &\xrightarrow{\Omega_i t I_{iz} + \Omega_j t I_{jz}} \cos(\Omega_i - \Omega_j)t \text{ZQ}_y^{(ij)} - \sin(\Omega_i - \Omega_j)t \text{ZQ}_x^{(ij)} \end{aligned}$$

Evolution under couplings

Multiple quantum coherence between spins i and j does not evolve under the influence of the coupling between the two active spins, i and j .

Double- and zero-quantum operators evolve under passive couplings in a way which is entirely analogous to the evolution of I_x and I_y ; the resulting multiple quantum terms can be described as being anti-phase with respect to the effective couplings:

$$\begin{aligned} \text{DQ}_x^{(ij)} &\longrightarrow \cos \pi J_{\text{DQ,eff}} t \text{DQ}_x^{(ij)} + \sin \pi J_{\text{DQ,eff}} t 2I_{kz} \text{DQ}_y^{(ij)} \\ \text{DQ}_y^{(ij)} &\longrightarrow \cos \pi J_{\text{DQ,eff}} t \text{DQ}_y^{(ij)} - \sin \pi J_{\text{DQ,eff}} t 2I_{kz} \text{DQ}_x^{(ij)} \\ \text{ZQ}_x^{(ij)} &\longrightarrow \cos \pi J_{\text{ZQ,eff}} t \text{ZQ}_x^{(ij)} + \sin \pi J_{\text{ZQ,eff}} t 2I_{kz} \text{ZQ}_y^{(ij)} \\ \text{ZQ}_y^{(ij)} &\longrightarrow \cos \pi J_{\text{ZQ,eff}} t \text{ZQ}_y^{(ij)} - \sin \pi J_{\text{ZQ,eff}} t 2I_{kz} \text{ZQ}_x^{(ij)} \end{aligned}$$

$J_{\text{DQ,eff}}$ is the sum of the couplings between spin i and all other spins *plus* the sum of the couplings between spin j and all other spins. $J_{\text{ZQ,eff}}$ is the sum of the couplings between spin i and all other spins *minus* the sum of the couplings between spin j and all other spins.

For example in a three-spin system the zero-quantum coherence between spins 1 and 2, anti-phase with respect to spin 3, evolves according to

$$2I_{3z} \text{ZQ}_y^{(12)} \longrightarrow \cos \pi J_{\text{ZQ,eff}} t 2I_{3z} \text{ZQ}_y^{(12)} - \sin \pi J_{\text{ZQ,eff}} t \text{ZQ}_x^{(12)}$$

where $J_{\text{ZQ,eff}} = J_{12} - J_{23}$

Further details of multiple-quantum evolution can be found in section 5.3 of Ernst, Bodenhausen and Wokaun *Principles of NMR in One and Two Dimensions* (Oxford University Press, 1987).

Exercises for Chapters 6, 7, 8 & 9

The more difficult and challenging problems are marked with an asterisk, *

Chapter 6: Product operators

E6-1 Using the standard rotations from section 6.1.4, express the following rotations in terms of sines and cosines:

$$\begin{aligned} & \exp(-i\theta I_x) I_y \exp(i\theta I_x) \quad \exp(-i\theta I_z) (-I_y) \exp(i\theta I_z) - \exp(-i\theta I_y) I_z \exp(i\theta I_y) \\ & \exp(-i\frac{1}{2}\theta I_y) I_x \exp(i\frac{1}{2}\theta I_y) \quad \exp(-i\theta I_x) I_x \exp(i\theta I_x) \quad \exp(i\theta I_x) (-I_z) \exp(-i\theta I_x) \end{aligned}$$

Express all of the above transformations in the shorthand notation of section 6.1.5.

E6-2. Repeat the calculation in section 6.1.6 for a spin echo with the 180° pulse about the y -axis. You should find that the magnetization refocuses onto the $-y$ axis.

E6-3. Assuming that magnetization *along the y -axis* gives rise to an absorption mode lineshape, draw sketches of the spectra which arise from the following operators

$$I_{1y} \quad I_{2x} \quad 2I_{1y}I_{2z} \quad 2I_{1x}I_{2x}$$

E6-4. Describe the following terms in words:

$$I_{1y} \quad I_{2z} \quad 2I_{1y}I_{2z} \quad 2I_{1x}I_{2x}$$

E6-5. Give the outcome of the following rotations

$$\begin{aligned} & I_{1x} \xrightarrow{\omega_1 t_p I_{1y}} \\ & 2I_{1x}I_{2z} \xrightarrow{(\pi/2)(I_{1y}+I_{2y})} \\ & 2I_{1x}I_{2z} \xrightarrow{-\pi I_{2y}} \\ & I_{1x} \xrightarrow{\Omega_1 t I_{1z}} \xrightarrow{\Omega_2 t I_{2z}} \\ & -I_{1x} \xrightarrow{2\pi J_{12} t I_{1z} I_{2z}} \\ & -2I_{1z}I_{2y} \xrightarrow{2\pi J_{12} t I_{1z} I_{2z}} \end{aligned}$$

Describe the outcome in words in each case.

E6-6. Consider the spin echo sequence

$$-\tau - 180^\circ(x, \text{ to spin 1 and spin 2}) - \tau -$$

applied to a two-spin system. Starting with magnetization along y , represented by I_{1y} , show that overall effect of the sequence is

$$I_{1y} \xrightarrow{\text{spin echo}} -\cos(2\pi J_{12} \tau) I_{1y} + \sin(2\pi J_{12} \tau) 2I_{1x}I_{2z}$$

You should ignore the effect of offsets, which are refocused, are just consider evolution due to coupling.

Is your result consistent with the idea that this echo sequence is equivalent to

$$-2\tau - 180^\circ(x, \text{ to spin 1 and spin 2})$$

[This calculation is rather more complex than that in section 6.4.1. You will need the identities

$$\cos 2\theta = \cos^2 \theta - \sin^2 \theta \quad \text{and} \quad \sin 2\theta = 2 \cos \theta \sin \theta]$$

E6-7. For a two-spin system, what delay, τ , in a spin echo sequence would you use to achieve the following overall transformations (do not worry about signs)? [$\cos \pi/4 = \sin \pi/4 = 1/\sqrt{2}$]

$$I_{2y} \longrightarrow 2I_{1z}I_{2x}$$

$$2I_{1z}I_{2x} \longrightarrow I_{2x}$$

$$I_{1x} \longrightarrow \left(\frac{1}{\sqrt{2}}\right)I_{1x} + \left(\frac{1}{\sqrt{2}}\right)2I_{1y}I_{2z}$$

$$I_{1x} \longrightarrow -I_{1x}$$

E6-8. Confirm by a calculation that spin echo sequence *c* shown on page 6–11 does not refocus the evolution of the offset of spin 1. [Start with a state I_{1x} or I_{1y} ; you may ignore the evolution due to coupling].

***E6-9.** Express $2I_{1x}I_{2y}$ in terms of raising and lowering operators: see section 6.5.2. Take the *zero-quantum part* of your expression and then re-write this in terms of Cartesian operators using the procedure shown in section 6.5.2.

E6-10. Consider three coupled spins in which $J_{23} > J_{12}$. Following section 6.6, draw a sketch of the doublet and doublets expected for the multiplet on spin 2 and label each line with the spin states of the coupled spins, 1 and 3. Label the splittings, too.

Assuming that magnetization along x gives an absorption mode lineshape, sketch the spectra from the following operators:

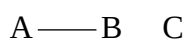
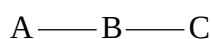
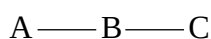
$$I_{2x} \quad 2I_{1z}I_{2x} \quad 2I_{2y}I_{3z} \quad 4I_{1z}I_{2x}I_{3z}$$

E6-11. Complete the following rotations.

$$\begin{aligned}
 I_{2x} &\xrightarrow{\Omega_1 t_{1z}} \xrightarrow{\Omega_2 t_{2z}} \xrightarrow{\Omega_3 t_{3z}} \\
 2I_{1y}I_{2z} &\xrightarrow{(\pi/2)(I_{1y}+I_{2y}+I_{3y})} \\
 4I_{1x}I_{2z}I_{3z} &\xrightarrow{(\pi/2)(I_{1y}+I_{2y}+I_{3y})} \\
 2I_{1z}I_{2x} &\xrightarrow{2\pi J_{12}t_{1z}I_{2z}} \\
 2I_{1z}I_{2x} &\xrightarrow{2\pi J_{23}t_{2z}I_{3z}} \\
 4I_{1z}I_{2y}I_{3z} &\xrightarrow{2\pi J_{23}t_{2z}I_{3z}} \\
 I_{1x} &\xrightarrow{2\pi J_{12}t_{1z}I_{2z}} \xrightarrow{2\pi J_{13}t_{1z}I_{3z}}
 \end{aligned}$$

Chapter 7: Two-dimensional NMR

E7-1. Sketch the COSY spectra you would expect from the following arrangements of spins. In the diagrams, a line represents a coupling. Assume that the spins have well separated shifts; do not concern yourself with the details of the multiplet structures of the cross- and diagonal-peaks.



E7-2. Sketch labelled two-dimensional spectra which have peaks arising from the following transfer processes

	frequencies in F_1 / Hz		frequencies in F_2 / Hz
<i>a</i>	30	Transferred to	30
<i>b</i>	30 and 60	transferred to	30
<i>c</i>	60	transferred to	30 and 60
<i>d</i>	30	transferred to	20, 30 and 60
<i>e</i>	30 and 60	transferred to	30 and 60

E7-3. What would the diagonal-peak multiplet of a COSY spectrum of two spins look like if we assigned the absorption mode lineshape in F_2 to magnetization along x and the absorption mode lineshape in F_1 to sine modulated data in t_1 ?

What would the cross-peak multiplet look like with these assignments?

E7-4. The smallest coupling that will gives rise to a discernible cross-peak in a COSY spectra depends on both the linewidth and the signal-to-noise ratio of the spectrum. Explain this observation.

***E7-5.** Complete the analysis of the DQF COSY spectrum by showing in detail that both the cross and diagonal-peak multiplets have the same lineshape and are in anti-phase in both dimensions. Start from the expression in the middle of page 7–10, section 7.4.21. [You will need the identity $\cos A \sin B = \frac{1}{2}[\sin(B + A) + \sin(B - A)]$].

***E7-6.** Consider the COSY spectrum for a three-spin system. Start with magnetization just on spin 1. The effect of the first pulse is

$$I_{1z} \xrightarrow{(\pi/2)(I_{1x}+I_{2x}+I_{3x})} -I_{1y}$$

Then, only the offset of spin 1 has an effect

$$-I_{1y} \xrightarrow{\Omega_1 t_1 I_{1z}} -\cos \Omega_1 t_1 I_{1y} + \sin \Omega_1 t_1 I_{1x}$$

Only the term in I_{1x} leads to cross- and diagonal peaks, so consider this term only from now on.

First allow it to evolve under the coupling to spin 2 and then the coupling to spin 3

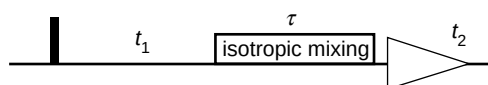
$$\sin \Omega_1 t_1 I_{1x} \xrightarrow{2\pi J_{12} t_1 I_{1x} I_{2z}} \xrightarrow{2\pi J_{13} t_1 I_{1x} I_{3z}}$$

Then, consider the effect of a $90^\circ(x)$ pulse applied to all three spins. After this pulse, you should find one term which represents a diagonal-peak multiplet, one which represents a cross-peak multiplet between spin 1 and spin 2, and one which represents a cross-peak multiplet between spin 1 and spin 3. What does the fourth term represent?

[More difficult] Determine the form of the cross-peak multiplets, using the approach adopted in section 6.4.1. Sketch the multiplets for the case $J_{12} \approx J_{23} > J_{13}$. [You will need the identity

$$\sin A \sin B = \frac{1}{2}[\cos(A + B) - \cos(A - B)] \quad]$$

***E7-7.** The pulse sequence for two-dimensional TOCSY (total correlation spectroscopy) is shown below



The mixing time, of length τ , is a period of *isotropic mixing*. This is a multiple-pulse sequence which results in the transfer of *in-phase* magnetization from one spin to another. In a two spin system the mixing goes as follows:

$$I_{1x} \xrightarrow{\text{isotropic mixing for time } \tau} \cos^2 \pi J_{12} \tau I_{1x} + \sin^2 \pi J_{12} \tau I_{2x}$$

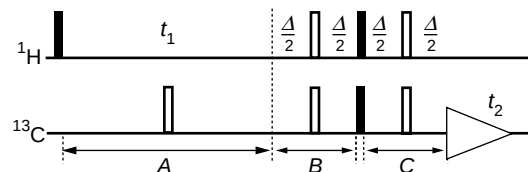
We can assume that all terms other than I_{1x} do not survive the isotropic mixing sequence, and so can be ignored.

Predict the form of the two-dimensional TOCSY spectrum for a two-spin system. What is the value of τ which gives the strongest cross peaks? For this optimum value of τ , what happens to the diagonal peaks? Can you think of any advantages that TOCSY might have over COSY?

E7-8. Repeat the analysis for the HMQC experiment, section 7.4.3.1, with the phase of the first spin-2 (carbon-13) pulse set to $-x$ rather than $+x$. Confirm that the observable signals present at the end of the sequence do indeed change sign.

E7-9. Why must the phase of the second spin-1 (proton) 90° pulse in the HSQC sequence, section 7.4.3.2, be y rather than x ?

E7-10. Below is shown the pulse sequence for the HETCOR (heteronuclear correlation) experiment



This sequence is closely related to HSQC, but differs in that the signal is observed on carbon-13, rather than being transferred back to proton for observation. Like HSQC and HMQC the resulting spectrum shows cross peaks whose co-ordinates are the shifts of directly attached carbon-13 proton pairs. However, in contrast to these sequences, in HETCOR the proton shift is in F_1 and the carbon-13 shift is in F_2 . In the early days of two-dimensional NMR this was a popular sequence for shift correlation as it is less demanding of the spectrometer; there are no strong signals from protons not coupled to carbon-13 to suppress.

We shall assume that spin 1 is proton, and spin 2 is carbon-13. During period A, t_1 , the offset of spin 1 evolves but the coupling between spins 1 and 2 is refocused by the centrally placed 180° pulse. During period B the coupling evolves, but the offset is refocused. The optimum value for the time Δ is $1/(2J_{12})$, as this leads to complete conversion into anti-phase. The two 90° pulses transfer the anti-phase magnetization to spin 2.

During period C the anti-phase magnetization rephases (the offset is refocused) and if Δ is $1/(2J_{12})$ the signal is purely in-phase at the start of t_2 .

Make an informal analysis of this sequence, along the lines of that given in section 7.4.3.2, and hence predict the form of the spectrum. In the first instance assume that Δ is set to its optimum value. Then, make the analysis slightly more complex and show that for an arbitrary value of Δ the signal intensity goes as $\sin^2 \pi J_{12} \Delta$.

Does altering the phase of the second spin-1 (proton) 90° pulse from x to y make any difference to the spectrum?

[Harder] What happens to carbon-13 magnetization, I_{2z} , present at the beginning of the sequence? How could the contribution from this be removed?

Chapter 8: Relaxation

E8-1. In an inversion-recovery experiment the following peak heights (S , arbitrary units) were measured as a function of the delay, t , in the sequence:

t / s	0.1	0.5	0.9	1.3	1.7	2.1	2.5	2.9
S	-98.8	-3.4	52.2	82.5	102.7	115.2	120.7	125.1

The peak height after a single 90° pulse was measured as 130.0. Use a graphical method to analyse these data and hence determine a value for the longitudinal relaxation rate constant and the corresponding value of the relaxation time, T_1 .

E8-2. In an experiment to estimate T_1 using the sequence $[180^\circ - \tau - 90^\circ]$ acquire three peaks in the spectrum were observed to go through a null at 0.5, 0.6 and 0.8 s respectively. Estimate T_1 for each of these resonances.

A solvent resonance was still inverted after a delay of 1.5 s; what does this tell you about the relaxation time of the solvent?

***E8-3.** Using the diagram at the top of page 8-5, write down expressions for dn_1/dt , dn_2/dt etc. in terms of the rate constants W and the populations n_i . [Do this without looking at the expressions given on page 8-6 and then check carefully to see that you have the correct expressions].

***E8-4.** Imagine a modified experiment, designed to record a transient NOE enhancement, in which rather than spin S being inverted at the beginning of the experiment, it is *saturated*. The initial conditions are thus

$$I_z(0) = I_z^0 \quad S_z(0) = 0$$

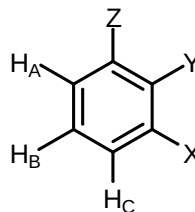
Using these starting conditions rather than those of Eq. [16] on page 8-10, show that in the initial rate limit the NOE enhancement builds up at a rate proportional to σ_{IS} rather than $2\sigma_{IS}$. You should use the method given in Section 8.4.1 for your analysis.

Without detailed calculation, sketch a graph, analogous to that given on page 8-12, for the behaviour of I_z and S_z for these new initial conditions as a function of mixing time.

E8-5. Why is it that in a two spin system the size of transient NOE enhancements depends on R_I , R_S and σ_{IS} , whereas in a steady state experiment the enhancement only depends on R_I and σ_{IS} ? [Spin S is the target].

In a particular two-spin system, S relaxes quickly and I relaxes slowly. Which experiment would you choose in order to measure the NOE enhancement between these two spins? Include in your answer an explanation of which spin you would irradiate.

E8-6. For the molecule shown on the right, a transient NOE experiment in which H_B is inverted gave equal initial NOE build-up rates on H_A and H_C . If H_A was inverted the initial build-up rate on H_B was the same as in the first experiment; no enhancement is seen of H_C . In steady state experiments, irradiation of H_B gave equal enhancements on H_A and H_B . However, irradiation of H_A gave a much smaller enhancement on H_B than for the case where H_B was the irradiated spin and the enhancement was observed on H_A . Explain.



E8-7. What do you understand by the terms *correlation time* and *spectral density*? Why are these quantities important in determining NMR relaxation rate constants?

E8-8. The simplest form of the spectral density, $J(\omega)$, is the Lorentzian:

$$J(\omega) = \frac{2\tau_c}{1 + \omega^2\tau_c^2}$$

Describe how this spectral density varies with both ω and τ_c . For a given frequency, ω_0 , at what correlation time is the spectral density a maximum?

Show how this form of the spectral density leads to the expectation that, for a given Larmor frequency, T_1 will have a minimum value at a certain value of the correlation time, τ_c .

E8-9. Suppose that the Larmor frequency (for proton) is 800 MHz. What correlation time will give the minimum value for T_1 ? What kind of molecule might have such a correlation time?

E8-10. Explain why the NOE enhancements observed in small molecules are positive whereas those observed for large molecules are negative.

E8-11. Explain how it is possible for the sign of an NOE enhancement to change when the magnetic field strength used by the spectrometer is changed.

E8-12. What is transverse relaxation and how it is different from longitudinal relaxation? Explain why it is that the rate constant for transverse relaxation increases with increasing correlation times, whereas that for longitudinal relaxation goes through a maximum.

Chapter 9: Coherence selection: phase cycling and gradient pulses

E9-1. (a) Show, using vector diagrams like those of section 9.1.6, that in a pulse-acquire experiment a phase cycle in which the pulse goes $x, y, -x, -y$ and in which the receiver phase is fixed leads to no signal after four transients have been co-added.

(b) In a simple spin echo sequence

$$90^\circ - \tau - 180^\circ - \tau -$$

the EXORCYCLE sequence involves cycling the 180° pulse $x, y, -x, -y$ and the receiver $x, -x, x, -x$. Suppose that, by accident, the 180° pulse has been omitted. Use vector pictures to show that the four step phase cycle cancels all the signal.

(c) In the simple echo sequence, suppose that there is some z -magnetization present at the end of the first τ delay; also suppose that the 180° pulse is imperfect so that some of the z -magnetization is made transverse. Show that the four steps of EXORCYCLE cancels the signal arising from this magnetization.

***E9-2.** (a) For the INEPT pulse sequence of section 9.1.8, confirm with product operator calculations that: [You should ignore the evolution of offsets as this is refocused by the spin echo; assume that the spin echo delay is $1/(2J_{IS})$].

(i) the sign of the signal transferred from I to S is altered by changing the phase of the second I spin 90° pulse from y to $-y$;

(ii) the signs of both the transferred signal and the signal originating from equilibrium S spin magnetization, S_z , are altered by changing the phase of the first S spin 90° pulse by 180° .

On the basis of your answers to (i) and (ii), suggest a suitable phase cycle, different to that given in the notes, for eliminating the contribution from the equilibrium S spin magnetization.

(b) Imagine that in the INEPT sequence the first I spin 180° pulse is cycled $x, y, -x, -y$. Without detailed calculations, deduce the effect of this cycle on the transferred signal and hence determine a suitable phase cycle for the receiver [hint - this 180° pulse is just forming a spin echo]. Does your cycle eliminate the contribution from the equilibrium S spin magnetization?

(c) Suppose now that the first S spin 180° pulse is cycled $x, y, -x, -y$; what effect does this have on the signal transferred from I to S ?

E9-3. Determine the coherence order or orders of each of the following operators [you will need to express I_x and I_y in terms of the raising and lowering operators, see section 9.3.1]

$$I_{1+}I_{2-} \quad 4I_{1+}I_{2+}I_{3z} \quad I_{1x} \quad I_{1y} \quad 2I_{1x}I_{2z} \quad (2I_{1x}I_{2x} + 2I_{1y}I_{2y})$$

In a heteronuclear system a coherence order can be assigned to each spin separately. If I and S represent different nuclei, assign separate coherence orders for the I and S spins to the following operators

$$I_x \quad S_y \quad 2I_xS_z \quad 2I_xS_x$$

E9-4. (a) Consider the phase cycle devised in section 9.5.1 which was designed to select $\Delta p = -3$: the pulse phase goes 0, 90, 180, 270 and the receiver phase goes 0, 270, 180, 90. Complete the following table and use it to show that such a cycle cancels signals arising from a pathway with $\Delta p = 0$.

step	pulse phase	phase shift experienced by pathway with $\Delta p = 0$	equivalent phase	rx. phase for $\Delta p = -3$	difference
1	0			0	
2	90			270	
3	180			180	
4	270			90	

Construct a similar table to show that a pathway with $\Delta p = -1$ is cancelled, but that one with $\Delta p = +5$ is selected by this cycle.

(b) Bodenhausen *et al.* have introduced a notation in which the sequence of possible Δp values is written out in a line; the values of Δp which are selected by the cycle are put into **bold** print, and those that are rejected are put into parenthesis, viz (1). Use this notation to describe the pathways selected and rejected by the cycle given above for pathways with Δp between -5 and $+5$ [the fate of several pathways is given in section 9.5.1, you have worked out two more in part (a) and you may also assume that the pathways with $\Delta p = -5, -4, -2, 3$ and 4 are rejected]. Confirm that, as expected for this four-step cycle, the selected values of Δp are separated by 4.

(c) Complete the following table for a *three-step* cycle designed to select $\Delta p = +1$.

step	pulse phase	phase shift experienced by pathway with $\Delta p = +1$	equivalent phase	rx. phase
1	0			
2	120°			
4	240°			

(d) Without drawing up further tables, use the general rules of section 9.5.2 to show that, in Bodenhausen's notation, the selectivity of the cycle devised in (c) can be written:

$$-2 \quad (-1) \quad (0) \quad \mathbf{1} \quad (2)$$

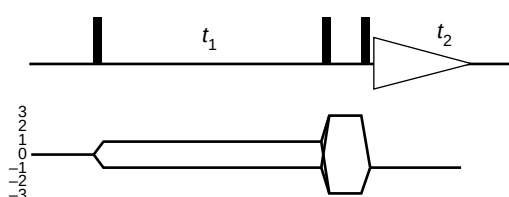
(e) Use Bodenhausen's notation to describe the selectivity of a 6 step cycle

designed to select $\Delta p = +1$; consider Δp values in the range -6 to $+6$.

E9-5. Draw coherence transfer pathways for (a) four-quantum filtered COSY [use the sequence in section 9.5.5.4 as a model]; (b) a 180° pulse used to refocus double-quantum coherence; (c) *N*-type NOESY.

E9-6. Write down four-step phase cycles to select (a) $\Delta p = -1$ and (b) $\Delta p = +2$. Suppose that the cycles are applied to different pulses. Combine them to give a 16-step cycle as was done in section 9.5.4; give the required receiver phase shifts.

***E9-7.** The CTP and pulse sequence of triple-quantum filtered COSY are



(a) Write down the values of Δp brought about by (i) the first pulse, (iii) the second pulse, (iii) the first and second pulses acting together and considered as a group, (iv) the last pulse.

(b) Imagine the spin system under consideration is able to support multiple quantum coherence up to and including $p = \pm 4$. Consider a phase cycle in which the first two pulses are cycled as a group and, having in mind your answer to (a) (iii), use the notation of Bodenhausen to indicate which pathways need to be selected and which blocked. Hence argue that the required phase cycle must have 6 steps. Draw up such a six-step phase cycle and confirm that the proposed sequence of pulse and receiver phases does indeed discriminate against filtration through double-quantum coherence.

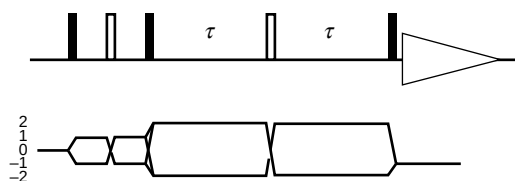
(c) Consider an alternative phase cycle in which just the phase of the last pulse is cycled. As in (b), use Bodenhausen's notation to describe the wanted and unwanted values of Δp . Devise a suitable phase cycle to select the required values of Δp .

(d) Try to devise a cycle which involves shifting just the phase of the second pulse.

E9-8. (a) Write down a 16 step cycle which selects the pathways shown in the double quantum spectroscopy pulse sequence of section 9.5.9.1; include in your cycle double-quantum selection and EXORCYCLE phase cycling of the 180° pulse.

(b) Write down an 8 step cycle which selects the pathway for NOESY shown in section 9.5.9.2; include in your cycle explicit axial peak suppression steps.

***E9-9.** The following sequence is one designed to measure the relaxation-induced decay of double-quantum coherence as a function of the time 2τ



The 180° pulse placed in the middle of the double-quantum period is used to refocus evolution due to offsets and inhomogeneous line broadening.

Devise a suitable phase cycle for the second 180° pulse bearing in mind that the 180° pulse may be imperfect. [hint: are four steps sufficient?]

E9-10. Use the formula given in section 9.6.2 for the overall decay of magnetization during a gradient

$$\frac{2}{\gamma G r_{\max}}$$

to calculate how long a gradient is needed to dephase magnetization to (a) 10% and (b) 1% of its initial value assuming that: $G = 0.1 \text{ T m}^{-1}$ (10 G cm^{-1}), $r_{\max} = 0.005 \text{ m}$ (0.5 cm) and $\gamma = 2.8 \times 10^8 \text{ rad s}^{-1}$. [Put all the quantities in SI units].

E9-11. Imagine that two gradients, G_1 and G_2 , are placed before and after a radiofrequency pulse. For each of the following ratios between the two gradients, identify *two* coherence transfer pathways which will be refocused:

(a) $G_1:G_2 = 1:1$ (b) $G_1:G_2 = -1:1$ (c) $G_1:G_2 = 1:-2$ (d) $G_1:G_2 = 0.5:1$

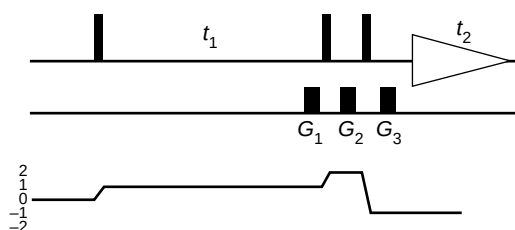
E9-12. Determine the ratio $G_1:G_2$ needed to select the following pathways:

(a) $p = 2 \rightarrow 1$ (b) $p = 3 \rightarrow 1$ (c) $p = -3 \rightarrow 1$

(d) $p = -1 \rightarrow 1$ and $p = 1 \rightarrow -1$ (e) $p = 0 \rightarrow 1$

Comment on the way in which case (e) differs from all the others.

E9-13. Consider a gradient selected N -type DQF COSY



We use three gradients; G_1 in t_1 so as to select $p = +1$ during t_1 ; G_2 to select double quantum during the filter delay and G_3 to refocus prior to acquisition.

(a) Show that the pathway shown, which can be denoted $0 \rightarrow 1 \rightarrow 2 \rightarrow -1$, is selected by gradients in the ratio $G_1:G_2:G_3 = 1:1:3$.

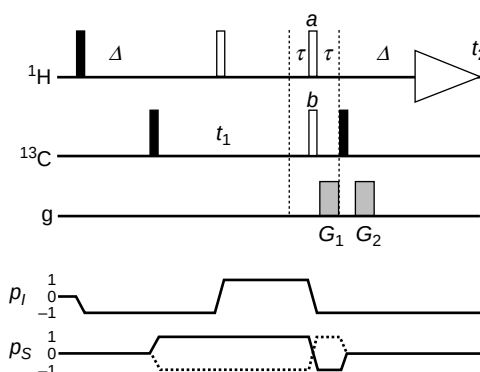
- (b) Show that this set of gradients also selects the pathway $0 \rightarrow -1 \rightarrow 4 \rightarrow -1$. What kind of spectrum does such a pathway give rise to?
- (c) Consider gradients in the ratios $G_1:G_2:G_3$ (i) $1:\frac{1}{2}:2$ and (ii) $1:\frac{1}{3}:\frac{5}{3}$. Show that these combinations select the DQ filtered pathway desired. In each case, give another possible pathway which has $p = \pm 1$ during t_1 and $p = -1$ during t_2 that these gradient combinations select.
- (d) In the light of (b) and (c), consider the utility of the gradient ratio $1.0 : 0.8 : 2.6$

***E9-14.**

Devise a gradient selected version of the triple-quantum filtered COSY experiment, whose basic pulse sequence and CTP was given in **E9-7**. Your sequence should include recommendations for the relative size of the gradients used. The resulting spectrum must have pure phase (*i.e.* $p = \pm 1$ must be preserved in t_1) and phase errors due to the evolution of offsets during the gradients must be removed.

How would you expect the sensitivity of your sequence to compare with its phase-cycled counterpart?

***E9-1.** A possible sequence for P/N selected HMQC is



The intention is to recombine P - and N -type spectra so as to obtain absorption mode spectra.

- (a) Draw CTPs for the P -type and N -type versions of this experiment [refer to section 9.6.7.3 for some hints].
- (b) What is the purpose of the two 180° pulses a and b , and why are such pulses needed for both proton and carbon-13?
- (c) Given that $\gamma_H/\gamma_C = 4$, what ratios of gradients are needed for the P -type and the N -type spectra?
- (d) Compare this sequence with that given in section 9.6.7.3, pointing out any advantages and disadvantages that each has.

7 Two-dimensional NMR[†]

7.1 Introduction

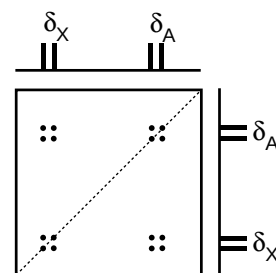
The basic ideas of two-dimensional NMR will be introduced by reference to the appearance of a COSY spectrum; later in this chapter the product operator formalism will be used to predict the form of the spectrum.

Conventional NMR spectra (one-dimensional spectra) are plots of intensity vs. frequency; in two-dimensional spectroscopy intensity is plotted as a function of two frequencies, usually called F_1 and F_2 . There are various ways of representing such a spectrum on paper, but the one most usually used is to make a contour plot in which the intensity of the peaks is represented by contour lines drawn at suitable intervals, in the same way as a topographical map. The position of each peak is specified by two frequency co-ordinates corresponding to F_1 and F_2 . Two-dimensional NMR spectra are always arranged so that the F_2 co-ordinates of the peaks correspond to those found in the normal one-dimensional spectrum, and this relation is often emphasized by plotting the one-dimensional spectrum alongside the F_2 axis.

The figure shows a schematic COSY spectrum of a hypothetical molecule containing just two protons, A and X, which are coupled together. The one-dimensional spectrum is plotted alongside the F_2 axis, and consists of the familiar pair of doublets centred on the chemical shifts of A and X, δ_A and δ_X respectively. In the COSY spectrum, the F_1 co-ordinates of the peaks in the two-dimensional spectrum also correspond to those found in the normal one-dimensional spectrum and to emphasize this point the one-dimensional spectrum has been plotted alongside the F_1 axis. It is immediately clear that this COSY spectrum has some symmetry about the diagonal $F_1 = F_2$ which has been indicated with a dashed line.

In a one-dimensional spectrum scalar couplings give rise to multiplets in the spectrum. In two-dimensional spectra the idea of a multiplet has to be expanded somewhat so that in such spectra a multiplet consists of an array of individual peaks often giving the impression of a square or rectangular outline. Several such arrays of peaks can be seen in the schematic COSY spectrum shown above. These two-dimensional multiplets come in two distinct types: *diagonal-peak multiplets* which are centred around the same F_1 and F_2 frequency co-ordinates and *cross-peak multiplets* which are centred around different F_1 and F_2 co-ordinates. Thus in the schematic COSY spectrum there are two diagonal-peak multiplets centred at $F_1 = F_2 = \delta_A$ and $F_1 = F_2 = \delta_X$, one cross-peak multiplet centred at $F_1 = \delta_A, F_2 = \delta_X$ and a second cross-peak multiplet centred at $F_1 = \delta_X, F_2 = \delta_A$.

The appearance in a COSY spectrum of a cross-peak multiplet $F_1 = \delta_A, F_2 = \delta_X = \delta_X$ indicates that the two protons at shifts δ_A and δ_X have a scalar coupling between them. This statement is all that is required for the analysis of a COSY spectrum, and it is this simplicity which is the key to the great utility



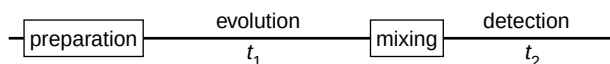
Schematic COSY spectrum for two coupled spins, A and X

[†] Chapter 7 "Two-Dimensional NMR" © James Keeler 1998, 2002 and 2004

of such spectra. From a single COSY spectrum it is possible to trace out the whole coupling network in the molecule

7.1.1 General Scheme for Two-Dimensional NMR

In one-dimensional pulsed Fourier transform NMR the signal is recorded as a function of one time variable and then Fourier transformed to give a spectrum which is a function of one frequency variable. In two-dimensional NMR the signal is recorded as a function of two time variables, t_1 and t_2 , and the resulting data Fourier transformed twice to yield a spectrum which is a function of two frequency variables. The general scheme for two-dimensional spectroscopy is



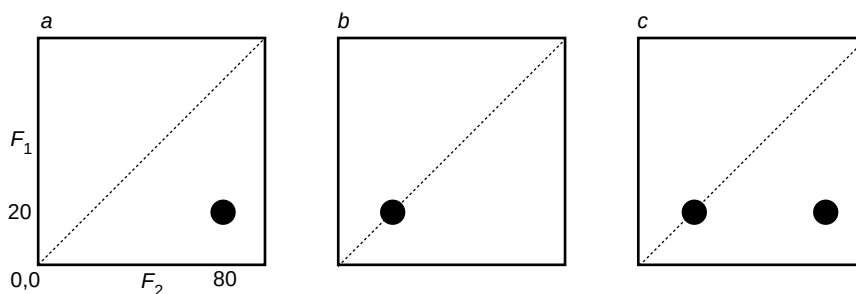
In the first period, called the preparation time, the sample is excited by one or more pulses. The resulting magnetization is allowed to evolve for the first time period, t_1 . Then another period follows, called the mixing time, which consists of a further pulse or pulses. After the mixing period the signal is recorded as a function of the second time variable, t_2 . This sequence of events is called a *pulse sequence* and the exact nature of the preparation and mixing periods determines the information found in the spectrum.

It is important to realize that the signal is not recorded during the time t_1 , but only during the time t_2 at the end of the sequence. The data is recorded at regularly spaced intervals in both t_1 and t_2 .

The two-dimensional signal is recorded in the following way. First, t_1 is set to zero, the pulse sequence is executed and the resulting free induction decay recorded. Then the nuclear spins are allowed to return to equilibrium. t_1 is then set to Δ_1 , the sampling interval in t_1 , the sequence is repeated and a free induction decay is recorded and stored separately from the first. Again the spins are allowed to equilibrate, t_1 is set to $2\Delta_1$, the pulse sequence repeated and a free induction decay recorded and stored. The whole process is repeated again for $t_1 = 3\Delta_1, 4\Delta_1$ and so on until sufficient data is recorded, typically 50 to 500 increments of t_1 . Thus recording a two-dimensional data set involves repeating a pulse sequence for increasing values of t_1 and recording a free induction decay as a function of t_2 for each value of t_1 .

7.1.2 Interpretation of peaks in a two-dimensional spectrum

Within the general framework outlined in the previous section it is now possible to interpret the appearance of a peak in a two-dimensional spectrum at particular frequency co-ordinates.



Suppose that in some unspecified two-dimensional spectrum a peak appears at $F_1 = 20$ Hz, $F_2 = 80$ Hz (spectrum *a* above). The interpretation of this peak is that a signal was present during t_1 which evolved with a frequency of 20 Hz. During the mixing time this *same* signal was transferred in some way to another signal which evolved at 80 Hz during t_2 .

Likewise, if there is a peak at $F_1 = 20$ Hz, $F_2 = 20$ Hz (spectrum *b*) the interpretation is that there was a signal evolving at 20 Hz during t_1 which was unaffected by the mixing period and continued to evolve at 20 Hz during t_2 . The processes by which these signals are transferred will be discussed in the following sections.

Finally, consider the spectrum shown in *c*. Here there are two peaks, one at $F_1 = 20$ Hz, $F_2 = 80$ Hz and one at $F_1 = 20$ Hz, $F_2 = 20$ Hz. The interpretation of this is that some signal was present during t_1 which evolved at 20 Hz and that during the mixing period part of it was transferred into another signal which evolved at 80 Hz during t_2 . The other part remained unaffected and continued to evolve at 20 Hz. On the basis of the previous discussion of COSY spectra, the part that changes frequency during the mixing time is recognized as leading to a cross-peak and the part that does not change frequency leads to a diagonal-peak. This kind of interpretation is a very useful way of thinking about the origin of peaks in a two-dimensional spectrum.

It is clear from the discussion in this section that the mixing time plays a crucial role in forming the two-dimensional spectrum. In the absence of a mixing time, the frequencies that evolve during t_1 and t_2 would be the same and only diagonal-peaks would appear in the spectrum. To obtain an interesting and useful spectrum it is essential to arrange for some process during the mixing time to transfer signals from one spin to another.

7.2 EXSY and NOESY spectra in detail

In this section the way in which the EXSY (EXchange SpectroscopY) sequence works will be examined; the pulse sequence is shown opposite. This experiment gives a spectrum in which a cross-peak at frequency coordinates $F_1 = \delta_A$, $F_2 = \delta_B$ indicates that the spin resonating at δ_A is chemically exchanging with the spin resonating at δ_B .

The pulse sequence for EXSY is shown opposite. The effect of the sequence will be analysed for the case of two spins, 1 and 2, but without any coupling between them. The initial state, before the first pulse, is equilibrium magnetization, represented as $I_{1z} + I_{2z}$; however, for simplicity only magnetization from the first spin will be considered in the calculation.

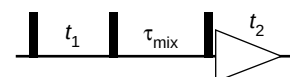
The first 90° pulse (of phase x) rotates the magnetization onto $-y$

$$I_{1z} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -I_{1y}$$

(the second arrow has no effect as it involves operators of spin 2). Next follows evolution for time t_1

$$-I_{1y} \xrightarrow{\Omega_1 t_1 I_{1z}} \xrightarrow{\Omega_2 t_1 I_{2z}} -\cos \Omega_1 t_1 I_{1y} + \sin \Omega_1 t_1 I_{1x}$$

again, the second arrow has no effect. The second 90° pulse turns the first



The pulse sequence for EXSY (and NOESY). All pulses have 90° flip angles.

term onto the z-axis and leaves the second term unaffected

$$\begin{aligned} -\cos \Omega_1 t_1 I_{1y} &\xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -\cos \Omega_1 t_1 I_{1z} \\ \sin \Omega_1 t_1 I_{1x} &\xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} \sin \Omega_1 t_1 I_{1x} \end{aligned}$$

Only the I_{1z} term leads to cross-peaks by chemical exchange, so the other term will be ignored (in an experiment this is achieved by appropriate coherence pathway selection). The effect of the first part of the sequence is to generate, at the start of the mixing time, τ_{mix} , some z-magnetization on spin 1 whose size depends, via the cosine term, on t_1 and the frequency, Ω_1 , with which the spin 1 evolves during t_1 . The magnetization is said to be *frequency labelled*.

During the mixing time, τ_{mix} , spin 1 may undergo chemical exchange with spin 2. If it does this, it carries with it the frequency label that it acquired during t_1 . The extent to which this transfer takes place depends on the details of the chemical kinetics; it will be assumed simply that during τ_{mix} a fraction f of the spins of type 1 chemically exchange with spins of type 2. The effect of the mixing process can then be written

$$-\cos \Omega_1 t_1 I_{1z} \xrightarrow{\text{mixing}} -(1-f) \cos \Omega_1 t_1 I_{1z} - f \cos \Omega_1 t_1 I_{2z}$$

The final 90° pulse rotates this z-magnetization back onto the y-axis

$$\begin{aligned} -(1-f) \cos \Omega_1 t_1 I_{1z} &\xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} (1-f) \cos \Omega_1 t_1 I_{1y} \\ -f \cos \Omega_1 t_1 I_{2z} &\xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} f \cos \Omega_1 t_1 I_{2y} \end{aligned}$$

Although the magnetization started on spin 1, at the end of the sequence there is magnetization present on spin 2 – a process called *magnetization transfer*. The analysis of the experiment is completed by allowing the I_{1y} and I_{2y} operators to evolve for time t_2 .

$$\begin{aligned} (1-f) \cos \Omega_1 t_1 I_{1y} &\xrightarrow{\Omega_1 t_2 I_{1z}} \xrightarrow{\Omega_2 t_2 I_{2z}} \\ &(1-f) \cos \Omega_1 t_2 \cos \Omega_1 t_1 I_{1y} - (1-f) \sin \Omega_1 t_2 \cos \Omega_1 t_1 I_{1x} \\ f \cos \Omega_1 t_1 I_{2y} &\xrightarrow{\Omega_1 t_2 I_{1z}} \xrightarrow{\Omega_2 t_2 I_{2z}} \\ &f \cos \Omega_2 t_2 \cos \Omega_1 t_1 I_{2y} - f \sin \Omega_2 t_2 \cos \Omega_1 t_1 I_{2x} \end{aligned}$$

If it is assumed that the y-magnetization is detected during t_2 (this is an arbitrary choice, but a convenient one), the time domain signal has two terms:

$$(1-f) \cos \Omega_1 t_2 \cos \Omega_1 t_1 + f \cos \Omega_2 t_2 \cos \Omega_1 t_1$$

The crucial thing is that the amplitude of the signal recorded during t_2 is modulated by the evolution during t_1 . This can be seen more clearly by imagining the Fourier transform, with respect to t_2 , of the above function. The $\cos \Omega_1 t_2$ and $\cos \Omega_2 t_2$ terms transform to give absorption mode signals centred at Ω_1 and Ω_2 respectively in the F_2 dimension; these are denoted $A_1^{(2)}$ and $A_2^{(2)}$ (the subscript indicates which spin, and the superscript which dimension). The time domain function becomes

$$(1-f) A_1^{(2)} \cos \Omega_1 t_1 + f A_2^{(2)} \cos \Omega_1 t_1$$

If a series of spectra recorded as t_1 progressively increases are inspected it

would be found that the $\cos \Omega_1 t_2$ term causes a change in size of the peaks at Ω_1 and Ω_2 – this is the modulation referred to above.

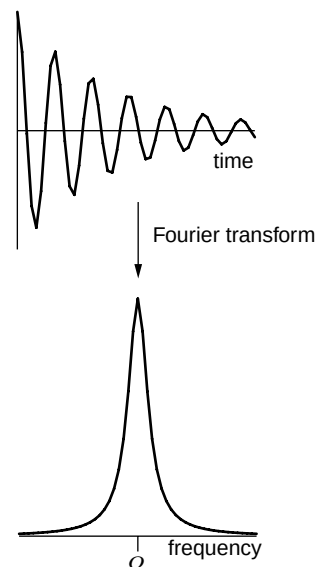
Fourier transformation with respect to t_1 gives peaks with an absorption lineshape, but this time in the F_1 dimension; an absorption mode signal at Ω_1 in F_1 is denoted $A_1^{(1)}$. The time domain signal becomes, after Fourier transformation in each dimension

$$(1-f)A_1^{(2)}A_1^{(1)} + fA_2^{(2)}A_1^{(1)}$$

Thus, the final two-dimensional spectrum is predicted to have two peaks. One is at $(F_1, F_2) = (\Omega_1, \Omega_1)$ – this is a diagonal peak and arises from those spins of type 1 which did not undergo chemical exchange during τ_{mix} . The second is at $(F_1, F_2) = (\Omega_1, \Omega_2)$ – this is a cross peak which indicates that part of the magnetization from spin 1 was transferred to spin 2 during the mixing time. It is this peak that contains the useful information. If the calculation were repeated starting with magnetization on spin 2 it would be found that there are similar peaks at (Ω_2, Ω_2) and (Ω_2, Ω_1) .

The NOESY (Nuclear Overhauser Effect Spectroscopy) spectrum is recorded using the same basic sequence. The only difference is that during the mixing time the cross-relaxation is responsible for the exchange of magnetization between different spins. Thus, a cross-peak indicates that two spins are experiencing mutual cross-relaxation and hence are close in space.

Having completed the analysis it can now be seen how the EXCSY/NOESY sequence is put together. First, the $90^\circ - t_1 - 90^\circ$ sequence is used to generate frequency labelled z-magnetization. Then, during τ_{mix} , this magnetization is allowed to migrate to other spins, carrying its label with it. Finally, the last pulse renders the z-magnetization observable.



The Fourier transform of a decaying cosine function $\cos \Omega t \exp(-t/T_2)$ is an absorption mode Lorentzian centred at frequency Ω ; the real part of the spectrum has been plotted.

7.3 More about two-dimensional transforms

From the above analysis it was seen that the signal observed during t_2 has an amplitude proportional to $\cos(\Omega_1 t_1)$; the amplitude of the signal observed during t_2 depends on the evolution during t_1 . For the first increment of t_1 ($t_1 = 0$), the signal will be a maximum, the second increment will have size proportional to $\cos(\Omega_1 \Delta_1)$, the third proportional to $\cos(\Omega_1 2\Delta_1)$, the fourth to $\cos(\Omega_1 3\Delta_1)$ and so on. This modulation of the amplitude of the observed signal by the t_1 evolution is illustrated in the figure below.

In the figure the first column shows a series of free induction decays that would be recorded for increasing values of t_1 and the second column shows the Fourier transforms of these signals. The final step in constructing the two-dimensional spectrum is to Fourier transform the data along the t_1 dimension. This process is also illustrated in the figure. Each of the spectra shown in the second column are represented as a series of data points, where each point corresponds to a different F_2 frequency. The data point corresponding to a particular F_2 frequency is selected from the spectra for $t_1 = 0$, $t_1 = \Delta_1$, $t_1 = 2\Delta_1$ and so on for all the t_1 values. Such a process results in a function, called an *interferogram*, which has t_1 as the running variable.

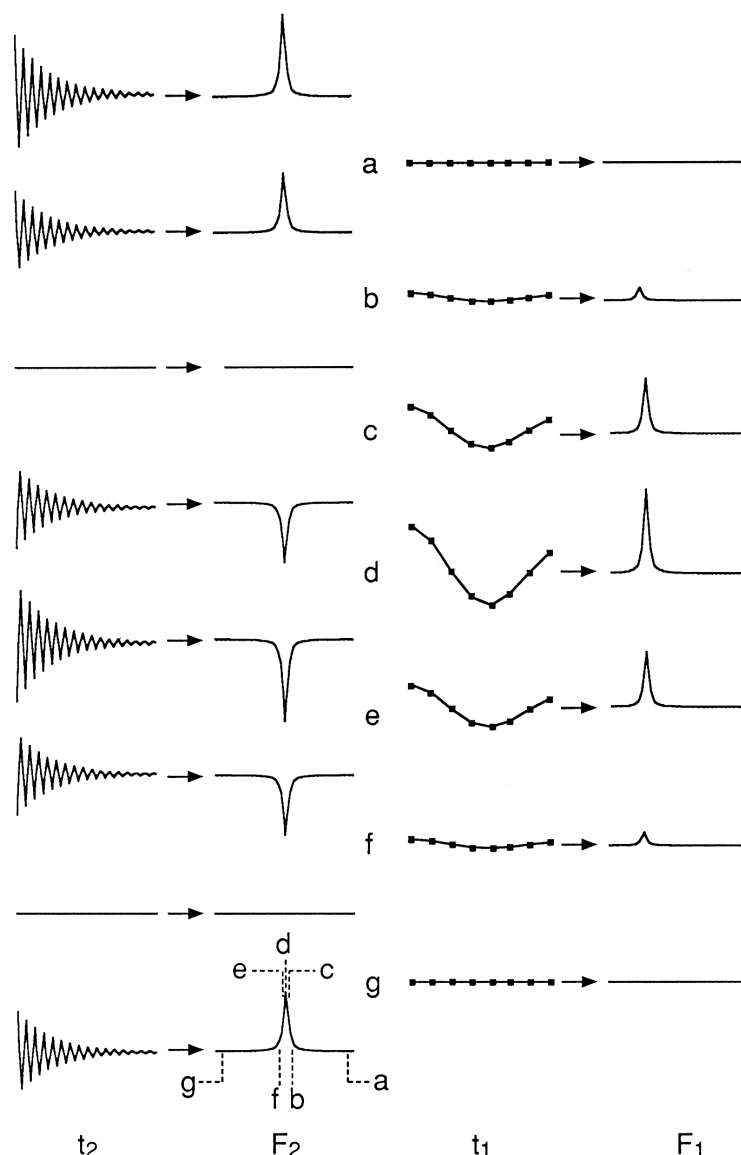


Illustration of how the modulation of a free induction decay by evolution during t_1 gives rise to a peak in the two-dimensional spectrum. In the left most column is shown a series of free induction decays that would be recorded for successive values of t_1 ; t_1 increases down the page. Note how the amplitude of these free induction decays varies with t_1 , something that becomes even plainer when the time domain signals are Fourier transformed, as shown in the second column. In practice, each of these F_2 spectra in column two consist of a series of data points. The data point at the same frequency in each of these spectra is extracted and assembled into an interferogram, in which the horizontal axis is the time t_1 . Several such interferograms, labelled *a* to *g*, are shown in the third column. Note that as there were eight F_2 spectra in column two corresponding to different t_1 values there are eight points in each interferogram. The F_2 frequencies at which the interferograms are taken are indicated on the lower spectrum of the second column. Finally, a second Fourier transformation of these interferograms gives a series of F_1 spectra shown in the right hand column. Note that in this column F_2 increases down the page, whereas in the first column t_1 increase down the page. The final result is a two-dimensional spectrum containing a single peak.

Several interferograms, labelled *a* to *g*, computed for different F_2 frequencies are shown in the third column of the figure. The particular F_2 frequency that each interferogram corresponds to is indicated in the bottom spectrum of the second column. The amplitude of the signal in each interferogram is different, but in this case the modulation frequency is the same. The final stage in the processing is to Fourier transform these interferograms to give the series of spectra which are shown in the right

most column of the figure. These spectra have F_1 running horizontally and F_2 running down the page. The modulation of the time domain signal has been transformed into a single two-dimensional peak. Note that the peak appears on several traces corresponding to different F_2 frequencies because of the width of the line in F_2 .

The time domain data in the t_1 dimension can be manipulated by multiplying by weighting functions or zero filling, just as with conventional free induction decays.

7.4 Two-dimensional experiments using coherence transfer through J-coupling

Perhaps the most important set of two-dimensional experiments are those which transfer magnetization from one spin to another via the scalar coupling between them. As was seen in section 6.3.3, this kind of transfer can be brought about by the action of a pulse on an anti-phase state. In outline the basic process is

$$I_{1x} \xrightarrow{\text{coupling}} 2I_{1y}I_{2z} \xrightarrow{90^\circ(x) \text{ to both spins}} 2I_{1z}I_{2y}$$

spin 1 spin 2

7.4.1 COSY

The pulse sequence for this experiment is shown opposite. It will be assumed in the analysis that all of the pulses are applied about the x-axis and for simplicity the calculation will start with equilibrium magnetization only on spin 1. The effect of the first pulse is to generate y-magnetization, as has been worked out previously many times

$$I_{1z} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -I_{1y}$$

This state then evolves for time t_1 , first under the influence of the offset of spin 1 (that of spin 2 has no effect on spin 1 operators):

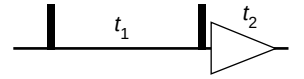
$$-I_{1y} \xrightarrow{\Omega_1 t_1 I_{1z}} -\cos \Omega_1 t_1 I_{1y} + \sin \Omega_1 t_1 I_{1x}$$

Both terms on the right then evolve under the coupling

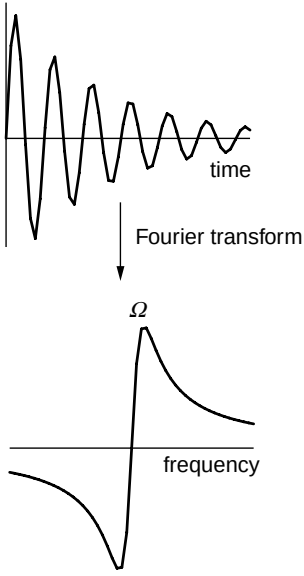
$$-\cos \Omega_1 t_1 I_{1y} \xrightarrow{2\pi J_{12} t_1 I_{1z} I_{2z}} -\cos \pi J_{12} t_1 \cos \Omega_1 t_1 I_{1y} + \sin \pi J_{12} t_1 \cos \Omega_1 t_1 2I_{1x} I_{2z}$$

$$\sin \Omega_1 t_1 I_{1x} \xrightarrow{2\pi J_{12} t_1 I_{1z} I_{2z}} \cos \pi J_{12} t_1 \sin \Omega_1 t_1 I_{1x} + \sin \pi J_{12} t_1 \sin \Omega_1 t_1 2I_{1y} I_{2z}$$

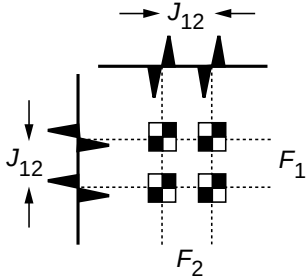
That completes the evolution under t_1 . Now all that remains is to consider the effect of the final pulse, remembering that the effect of the pulse on both spins needs to be computed. Taking the terms one by one:



Pulse sequence for the two-dimensional COSY experiment



The Fourier transform of a decaying sine function $\sin \Omega t \exp(-t/T_2)$ is a dispersion mode Lorentzian centred at frequency Ω .



Schematic view of the diagonal peak from a COSY spectrum. The squares are supposed to indicate the two-dimensional double dispersion lineshape illustrated below

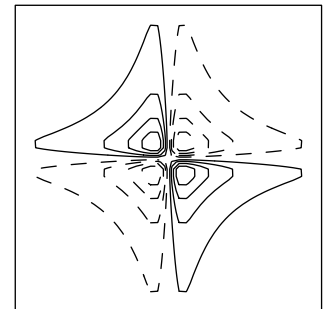
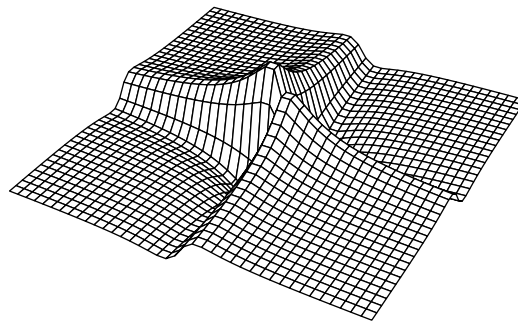
$$\begin{aligned}
 & -\cos \pi J_{12} t_1 \cos \Omega_1 t_1 I_{1y} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -\cos \pi J_{12} t_1 \cos \Omega_1 t_1 I_{1z} \quad \{1\} \\
 & \sin \pi J_{12} t_1 \cos \Omega_1 t_1 2I_{1x} I_{2z} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -\sin \pi J_{12} t_1 \cos \Omega_1 t_1 2I_{1x} I_{2y} \quad \{2\} \\
 & \cos \pi J_{12} t_1 \sin \Omega_1 t_1 I_{1x} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} \cos \pi J_{12} t_1 \sin \Omega_1 t_1 I_{1x} \quad \{3\} \\
 & \sin \pi J_{12} t_1 \sin \Omega_1 t_1 2I_{1y} I_{2z} \xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} -\sin \pi J_{12} t_1 \sin \Omega_1 t_1 2I_{1z} I_{2y} \quad \{4\}
 \end{aligned}$$

Terms {1} and {2} are unobservable. Term {3} corresponds to in-phase magnetization of spin 1, aligned along the x-axis. The t_1 modulation of this term depends on the offset of spin 1, so a diagonal peak centred at (Ω_1, Ω_1) is predicted. Term {4} is the really interesting one. It shows that anti-phase magnetization on spin 1, $2I_{1y}I_{2z}$, is transferred to anti-phase magnetization on spin 2, $2I_{1z}I_{2y}$; this is an example of coherence transfer. Term {4} appears as observable magnetization on spin 2, but it is modulated in t_1 with the offset of spin 1, thus it gives rise to a cross-peak centred at (Ω_1, Ω_2) . It has been shown, therefore, how cross- and diagonal-peaks arise in a COSY spectrum.

Some more consideration should be given to the form of the cross- and diagonal peaks. Consider again term {3}: it will give rise to an in-phase multiplet in F_2 , and as it is along the x-axis, the lineshape will be dispersive. The form of the modulation in t_1 can be expanded, using the formula, $\cos A \sin B = \frac{1}{2} \{ \sin(B+A) + \sin(B-A) \}$ to give

$$\cos \pi J_{12} t_1 \sin \Omega_1 t_1 = \frac{1}{2} \{ \sin(\Omega_1 t_1 + \pi J_{12} t_1) + \sin(\Omega_1 t_1 - \pi J_{12} t_1) \}$$

Two peaks in F_1 are expected at $\Omega_1 \pm \pi J_{12}$, these are just the two lines of the spin 1 doublet. In addition, since these are sine modulated they will have the dispersion lineshape. Note that both components in the spin 1 multiplet observed in F_2 are modulated in this way, so the appearance of the two-dimensional multiplet can best be found by "multiplying together" the multiplets in the two dimensions, as shown opposite. In addition, all four components of the diagonal-peak multiplet have the same sign, and have the double dispersion lineshape illustrated below

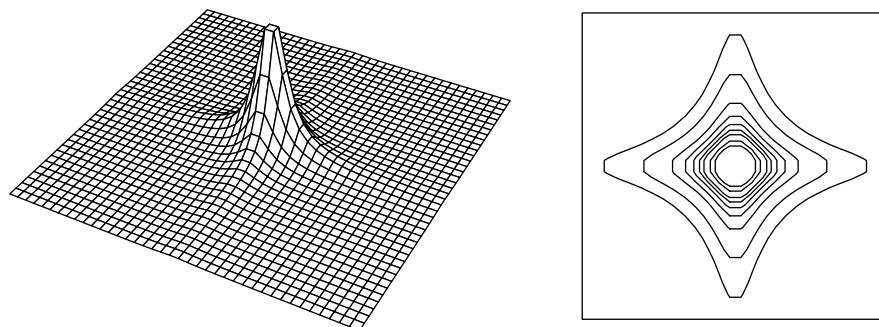


The double dispersion lineshape seen in pseudo 3D and as a contour plot; negative contours are indicated by dashed lines.

Term {4} can be treated in the same way. In F_2 we know that this term gives rise to an anti-phase absorption multiplet on spin 2. Using the relationship $\sin B \sin A = \frac{1}{2} \{ -\cos(B+A) + \cos(B-A) \}$ the modulation in t_1 can be expanded

$$\sin \pi J_{12} t_1 \sin \Omega_1 t_1 = \frac{1}{2} \{ -\cos(\Omega_1 t_1 + \pi J_{12} t_1) + \cos(\Omega_1 t_1 - \pi J_{12} t_1) \}$$

Two peaks in F_1 , at $\Omega_1 \pm \pi J_{12}$, are expected; these are just the two lines of the spin 1 doublet. Note that the two peaks have opposite signs – that is they are anti-phase in F_1 . In addition, since these are cosine modulated we expect the absorption lineshape (see section 7.2). The form of the cross-peak multiplet can be predicted by "multiplying together" the F_1 and F_2 multiplets, just as was done for the diagonal-peak multiplet. The result is shown opposite. This characteristic pattern of positive and negative peaks that constitutes the cross-peak is known as an *anti-phase square array*.



The double absorption lineshape seen in pseudo 3D and as a contour plot.

COSY spectra are sometimes plotted in the *absolute value mode*, where all the sign information is suppressed deliberately. Although such a display is convenient, especially for routine applications, it is generally much more desirable to retain the sign information. Spectra displayed in this way are said to be *phase sensitive*; more details of this are given in section 7.6.

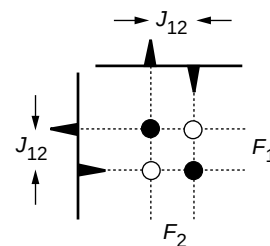
As the coupling constant becomes comparable with the linewidth, the positive and negative peaks in the cross-peak multiplet begin to overlap and cancel one another out. This leads to an overall reduction in the intensity of the cross-peak multiplet, and ultimately the cross-peak disappears into the noise in the spectrum. The smallest coupling which gives rise to a cross-peak is thus set by the linewidth and the signal-to-noise ratio of the spectrum.

7.1.2 Double-quantum filtered COSY (DQF COSY)

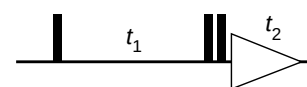
The conventional COSY experiment suffers from a disadvantage which arises from the different phase properties of the cross- and diagonal-peak multiplets. The components of a diagonal peak multiplet are all in-phase and so tend to reinforce one another. In addition, the dispersive tails of these peaks spread far into the spectrum. The result is a broad intense diagonal which can obscure nearby cross-peaks. This effect is particularly troublesome when the coupling is comparable with the linewidth as in such cases, as was described above, cancellation of anti-phase components in the cross-peak multiplet reduces the overall intensity of these multiplets.

This difficulty is neatly side-stepped by a modification called double quantum filtered COSY (DQF COSY). The pulse sequence is shown opposite.

Up to the second pulse the sequence is the same as COSY. However, it is arranged that only double-quantum coherence present during the (very short) delay between the second and third pulses is ultimately allowed to



Schematic view of the cross-peak multiplet from a COSY spectrum. The circles are supposed to indicate the two-dimensional double absorption lineshape illustrated below; filled circles represent positive intensity, open represent negative intensity.



The pulse sequence for DQF COSY; the delay between the last two pulses is usually just a few microseconds.

contribute to the spectrum. Hence the name, "double-quantum filtered", as all the observed signals are filtered through double-quantum coherence. The final pulse is needed to convert the double quantum coherence back into observable magnetization. This double-quantum derived signal is selected by the use of coherence pathway selection using phase cycling or field gradient pulses.

In the analysis of the COSY experiment, it is seen that after the second 90° pulse it is term {2} that contains double-quantum coherence; this can be demonstrated explicitly by expanding this term in the raising and lowering operators, as was done in section 6.5

$$\begin{aligned} 2I_{1x}I_{2y} &= 2 \times \frac{1}{2}(I_{1+} + I_{1-}) \times \frac{1}{2i}(I_{2+} - I_{2-}) \\ &= \frac{1}{2i}(I_{1+}I_{2+} - I_{1-}I_{2-}) + \frac{1}{2i}(-I_{1+}I_{2-} + I_{1-}I_{2+}) \end{aligned}$$

This term contains both double- and zero-quantum coherence. The pure double-quantum part is the term in the first bracket on the right; this term can be re-expressed in Cartesian operators:

$$\begin{aligned} \frac{1}{2i}(I_{1+}I_{2+} - I_{1-}I_{2-}) &= \frac{1}{2i}[(I_{1x} + iI_{1y})(I_{2x} + iI_{2y}) + (I_{1x} - iI_{1y})(I_{2x} - iI_{2y})] \\ &= \frac{1}{2}[2I_{1x}I_{2x} + 2I_{1y}I_{2y}] \end{aligned}$$

The effect of the last $90^\circ(x)$ pulse on the double quantum part of term {2} is thus

$$\begin{aligned} -\frac{1}{2} \sin \pi J_{12} t_1 \cos \Omega_1 t_1 (2I_{1x}I_{2y} + 2I_{1y}I_{2x}) &\xrightarrow{\pi/2 I_{1x}} \xrightarrow{\pi/2 I_{2x}} \\ &-\frac{1}{2} \sin \pi J_{12} t_1 \cos \Omega_1 t_1 (2I_{1x}I_{2z} + 2I_{1z}I_{2x}) \end{aligned}$$

The first term on the right is anti-phase magnetization of spin 1 aligned along the x-axis; this gives rise to a diagonal-peak multiplet. The second term is anti-phase magnetization of spin 2, again aligned along x; this will give rise to a cross-peak multiplet. Both of these terms have the same modulation in t_1 , which can be shown, by a similar analysis to that used above, to lead to an anti-phase multiplet in F_1 . As these peaks all have the same lineshape the overall phase of the spectrum can be adjusted so that they are all in absorption; see section 7.6 for further details. In contrast to the case of a simple COSY experiment *both* the diagonal- and cross-peak multiplets are in anti-phase in both dimensions, thus avoiding the strong in-phase diagonal peaks found in the simple experiment. The DQF COSY experiment is the method of choice for tracing out coupling networks in a molecule.

7.1.3 Heteronuclear correlation experiments

One particularly useful experiment is to record a two-dimensional spectrum in which the co-ordinate of a peak in one dimension is the chemical shift of one type of nucleus (*e.g.* proton) and the co-ordinate in the other dimension is the chemical shift of another nucleus (*e.g.* carbon-13) which is coupled to the first nucleus. Such spectra are often called *shift correlation maps* or *shift correlation spectra*.

The one-bond coupling between a carbon-13 and the proton directly

attached to it is relatively constant (around 150 Hz), and much larger than any of the long-range carbon-13 proton couplings. By utilizing this large difference experiments can be devised which give maps of carbon-13 shifts vs the shifts of directly attached protons. Such spectra are very useful as aids to assignment; for example, if the proton spectrum has already been assigned, simply recording a carbon-13 proton correlation experiment will give the assignment of all the protonated carbons.

Only one kind of nuclear species can be observed at a time, so there is a choice as to whether to observe carbon-13 or proton when recording a shift correlation spectrum. For two reasons, it is very advantageous from the sensitivity point of view to record protons. First, the proton magnetization is larger than that of carbon-13 because there is a larger separation between the spin energy levels giving, by the Boltzmann distribution, a greater population difference. Second, a given magnetization induces a larger voltage in the coil the higher the NMR frequency becomes.

Trying to record a carbon-13 proton shift correlation spectrum by proton observation has one serious difficulty. Carbon-13 has a natural abundance of only 1%, thus 99% of the molecules in the sample do not have any carbon-13 in them and so will not give signals that can be used to correlate carbon-13 and proton. The 1% of molecules with carbon-13 will give a perfectly satisfactory spectrum, but the signals from these resonances will be swamped by the much stronger signals from non-carbon-13 containing molecules. However, these unwanted signals can be suppressed using coherence selection in a way which will be described below.

7.1.3.1 Heteronuclear multiple-quantum correlation (HMQC)

The pulse sequence for this popular experiment is given opposite. The sequence will be analysed for a coupled carbon-13 proton pair, where spin 1 will be the proton and spin 2 the carbon-13.

The analysis will start with equilibrium magnetization on spin 1, I_{1z} . The whole analysis can be greatly simplified by noting that the 180° pulse is exactly midway between the first 90° pulse and the start of data acquisition. As has been shown in section 6.4, such a sequence forms a spin echo and so the evolution of the offset of spin 1 over the entire period ($t_1 + 2\Delta$) is refocused. Thus the evolution of the offset of spin 1 can simply be ignored for the purposes of the calculation.

At the end of the delay Δ the state of the system is simply due to evolution of the term $-I_{1y}$ under the influence of the scalar coupling:

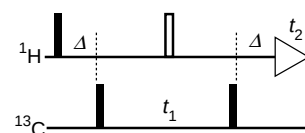
$$-\cos \pi J_{12} \Delta I_{1y} + \sin \pi J_{12} \Delta 2I_{1x} I_{2z}$$

It will be assumed that $\Delta = 1/(2J_{12})$, so only the anti-phase term is present.

The second 90° pulse is applied to carbon-13 (spin 2) only

$$2I_{1x} I_{2z} \xrightarrow{\pi/2 I_{2x}} -2I_{1x} I_{2y}$$

This pulse generates a mixture of heteronuclear double- and zero-quantum coherence, which then evolves during t_1 . In principle this term evolves under the influence of the offsets of spins 1 and 2 and the coupling between them. However, it has already been noted that the offset of spin 1 is



The pulse sequence for HMQC. Filled rectangles represent 90° pulses and open rectangles represent 180° pulses. The delay Δ is set to $1/(2J_{12})$.

refocused by the centrally placed 180° pulse, so it is not necessary to consider evolution due to this term. In addition, it can be shown that multiple-quantum coherence involving spins i and j does not evolve under the influence of the coupling, J_{ij} , between these two spins. As a result of these two simplifications, the only evolution that needs to be considered is that due to the offset of spin 2 (the carbon-13).

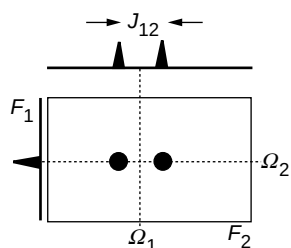
$$-2I_{1x}I_{2y} \xrightarrow{\Omega_2 t_1} -\cos \Omega_2 t_1 \ 2I_{1x}I_{2y} + \sin \Omega_2 t_1 \ 2I_{1x}I_{2x}$$

The second 90° pulse to spin 2 (carbon-13) regenerates the first term on the right into spin 1 (proton) observable magnetization; the other remains unobservable

$$-\cos \Omega_2 t_1 \ 2I_{1x}I_{2y} \xrightarrow{\pi/2 I_{2x}} -\cos \Omega_2 t_1 \ 2I_{1x}I_{2z}$$

This term then evolves under the coupling, again it is assumed that $\Delta = 1/(2J_{12})$

$$-\cos \Omega_2 t_1 \ 2I_{1x}I_{2z} \xrightarrow{2\pi J_{12} \Delta t_1 I_{2z}, \Delta=1/(2J_{12})} -\cos \Omega_2 t_1 \ I_{1y}$$



Schematic HMQC spectrum for two coupled spins.

This is a very nice result; in F_2 there will be an in-phase doublet centred at the offset of spin 1 (proton) and these two peaks will have an F_1 co-ordinate simply determined by the offset of spin 2 (carbon-13); the peaks will be in absorption. A schematic spectrum is shown opposite.

The problem of how to suppress the very strong signals from protons not coupled to any carbon-13 nuclei now has to be addressed. From the point of view of these protons the carbon-13 pulses might as well not even be there, and the pulse sequence looks like a simple spin echo. This insensitivity to the carbon-13 pulses is the key to suppressing the unwanted signals. Suppose that the phase of the first carbon-13 90° pulse is altered from x to $-x$. Working through the above calculation it is found that the wanted signal from the protons coupled to carbon-13 changes sign *i.e.* the observed spectrum will be inverted. In contrast the signal from a proton *not* coupled to carbon-13 will be unaffected by this change. Thus, for each t_1 increment the free induction decay is recorded twice: once with the first carbon-13 90° pulse set to phase x and once with it set to phase $-x$. The two free induction decays are then subtracted in the computer memory thus cancelling the unwanted signals. This is an example of a very simple phase cycle.

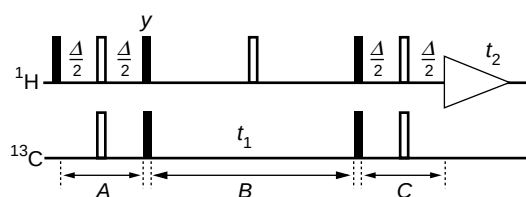
In the case of carbon-13 and proton the one bond coupling is so much larger than any of the long range couplings that a choice of $\Delta = 1/(2J_{\text{one bond}})$ does not give any correlations other than those through the one-bond coupling. There is simply insufficient time for the long-range couplings to become anti-phase. However, if Δ is set to a much longer value (30 to 60 ms), long-range correlations will be seen. Such spectra are very useful in assigning the resonances due to quaternary carbon-13 atoms. The experiment is often called HMBC (heteronuclear multiple-bond correlation).

Now that the analysis has been completed it can be seen what the function of various elements in the pulse sequence is. The first pulse and delay generate magnetization on proton which is anti-phase with respect to the coupling to carbon-13. The carbon-13 90° pulse turns this into multiple quantum coherence. This forms a filter through which magnetization not

bound to carbon-13 cannot pass and it is the basis of discrimination between signals from protons bound and not bound to carbon-13. The second carbon-13 pulse returns the multiple quantum coherence to observable anti-phase magnetization on proton. Finally, the second delay Δ turns the anti-phase state into an in-phase state. The centrally placed proton 180° pulse refocuses the proton shift evolution for both the delays Δ and t_1 .

7.1.3.2 Heteronuclear single-quantum correlation (HSQC)

This pulse sequence results in a spectrum identical to that found for HMQC. Despite the pulse sequence being a little more complex than that for HMQC, HSQC has certain advantages for recording the spectra of large molecules, such as proteins. The HSQC pulse sequence is often embedded in much more complex sequences which are used to record two- and three-dimensional spectra of carbon-13 and nitrogen-15 labelled proteins.



The pulse sequence for HSQC. Filled rectangles represent 90° pulses and open rectangles represent 180° pulses. The delay Δ is set to $1/(2J_{12})$; all pulses have phase x unless otherwise indicated.

If this sequence were to be analysed by considering each delay and pulse in turn the resulting calculation would be far too complex to be useful. A more intelligent approach is needed where simplifications are used, for example by recognizing the presence of spin echoes who refocus offsets or couplings. Also, it is often the case that attention can be focused a particular terms, as these are the ones which will ultimately lead to observable signals. This kind of "intelligent" analysis will be illustrated here.

Periods A and C are spin echoes in which 180° pulses are applied to both spins; it therefore follows that the offsets of spins 1 and 2 will be refocused, but the coupling between them will evolve throughout the entire period. As the total delay in the spin echo is $1/(2J_{12})$ the result will be the complete conversion of in-phase into anti-phase magnetization.

Period B is a spin echo in which a 180° pulse is applied only to spin 1. Thus, the offset of spin 1 is refocused, as is the coupling between spins 1 and 2; only the offset of spin 2 affects the evolution.

With these simplifications the analysis is easy. The first pulse generates $-I_{1y}$; during period A this then becomes $-2I_{1x}I_{2z}$. The $90^\circ(y)$ pulse to spin 1 turns this to $2I_{1z}I_{2z}$ and the $90^\circ(x)$ pulse to spin 2 turns it to $-2I_{1z}I_{2y}$. The evolution during period B is simply under the offset of spin 2

$$-2I_{1z}I_{2y} \xrightarrow{\Omega_2 t_1 I_{2z}} -\cos \Omega_2 t_1 2I_{1z}I_{2y} + \sin \Omega_2 t_1 2I_{1z}I_{2x}$$

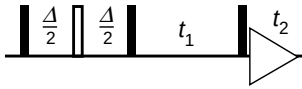
The next two 90° pulses transfer the first term to spin 1; the second term is rotated into multiple quantum and is not observed

$$-\cos \Omega_2 t_1 \ 2I_{1z}I_{2y} + \sin \Omega_2 t_1 \ 2I_{1z}I_{2x} \xrightarrow{\pi/2(I_{1x}+I_{2x})} \\ -\cos \Omega_2 t_1 \ 2I_{1y}I_{2z} - \sin \Omega_2 t_1 \ 2I_{1y}I_{2x}$$

The first term on the right evolves during period C into in-phase magnetization (the evolution of offsets is refocused). So the final observable term is $\cos \Omega_2 t_1 \ I_{1x}$. The resulting spectrum is therefore an in-phase doublet in F_2 , centred at the offset of spin 1, and these peaks will both have the same frequency in F_1 , namely the offset of spin 2. The spectrum looks just like the HMQC spectrum.

7.5 Advanced topic: Multiple-quantum spectroscopy

A key feature of two-dimensional NMR experiments is that no direct observations are made during t_1 , it is thus possible to detect, indirectly, the evolution of unobservable coherences. An example of the use of this feature is in the indirect detection of multiple-quantum spectra. A typical pulse sequence for such an experiment is shown opposite



Pulse sequence for multiple-quantum spectroscopy.

For a two-spin system the optimum value for Δ is $1/(2J_{12})$. The sequence can be dissected as follows. The initial $90^\circ - \Delta/2 - 180^\circ - \Delta/2 -$ sequence is a spin echo which, at time Δ , refocuses any evolution of offsets but allows the coupling to evolve and generate anti-phase magnetization. This anti-phase magnetization is turned into multiple-quantum coherence by the second 90° pulse. After evolving for time t_1 the multiple quantum is returned into observable (anti-phase) magnetization by the final 90° pulse. Thus the first three pulses form the preparation period and the last pulse is the mixing period.

7.5.1 Double-quantum spectrum for a three-spin system

The sequence will be analysed for a system of three spins. A complete analysis would be rather lengthy, so attention will be focused on certain terms as above, as many simplifying assumptions as possible will be made about the sequence.

The starting point will be equilibrium magnetization on spin 1, I_{1z} ; after the spin echo the magnetization has evolved due to the coupling between spin 1 and spin 2, and the coupling between spin 1 and spin 3 (the 180° pulse causes an overall sign change (see section 6.4.1) but this has no real effect here so it will be ignored)

$$-I_{1y} \xrightarrow{2\pi J_{12}\Delta I_{1z}I_{2z}} -\cos \pi J_{12}\Delta \ I_{1y} + \sin \pi J_{12}\Delta \ 2I_{1x}I_{2z} \\ \xrightarrow{2\pi J_{13}\Delta I_{1z}I_{3z}} -\cos \pi J_{13}\Delta \cos \pi J_{12}\Delta \ I_{1y} + \sin \pi J_{13}\Delta \cos \pi J_{12}\Delta \ 2I_{1x}I_{3z} \quad [3.1] \\ + \cos \pi J_{13}\Delta \sin \pi J_{12}\Delta \ 2I_{1x}I_{2z} + \sin \pi J_{13}\Delta \sin \pi J_{12}\Delta \ 4I_{1y}I_{2z}I_{3z}$$

Of these four terms, all but the first are turned into multiple-quantum by the second 90° pulse. For example, the second term becomes a mixture of double and zero quantum between spins 1 and 3

$$\sin \pi J_{13}\Delta \cos \pi J_{12}\Delta \ 2I_{1x}I_{3z} \xrightarrow{\pi/2(I_{1x}+I_{2x}+I_{3x})} -\sin \pi J_{13}\Delta \cos \pi J_{12}\Delta \ 2I_{1x}I_{3y}$$

It will be assumed that appropriate coherence pathway selection has been used so that ultimately only the double-quantum part contributes to the

spectrum. This part is

$$\left[-\sin \pi J_{13} \Delta \cos \pi J_{12} \Delta \right] \left\{ \frac{1}{2} (2I_{1x} I_{3y} + 2I_{1y} I_{3x}) \right\} \equiv B_{13} \text{DQ}_y^{(13)}$$

The term in square brackets just gives the overall intensity, but does not affect the frequencies of the peaks in the two-dimensional spectrum as it does not depend on t_1 or t_2 ; this intensity term is denoted B_{13} for brevity. The operators in the curly brackets represent a pure double quantum state which can be denoted $\text{DQ}_y^{(13)}$; the superscript (13) indicates that the double quantum is between spins 1 and 3 (see section 6.9).

As is shown in section 6.9, such a double-quantum term evolves under the offset according to

$$B_{13} \text{DQ}_y^{(13)} \xrightarrow{\Omega_1 t_1 I_{1z} + \Omega_2 t_1 I_{2z} + \Omega_3 t_1 I_{3z}} B_{13} \cos(\Omega_1 + \Omega_3) t_1 \text{DQ}_y^{(13)} - B_{13} \sin(\Omega_1 + \Omega_3) t_1 \text{DQ}_x^{(13)}$$

where $\text{DQ}_x^{(13)} \equiv \frac{1}{2} (2I_{1x} I_{3x} - 2I_{1y} I_{3y})$. This evolution is analogous to that of a single spin where y rotates towards $-x$.

As is also shown in section 6.9, $\text{DQ}_y^{(13)}$ and $\text{DQ}_x^{(13)}$ do not evolve under the coupling between spins 1 and 3, but they do evolve under the *sum* of the couplings between these two and all other spins; in this case this is simply $(J_{12} + J_{23})$. Taking each term in turn

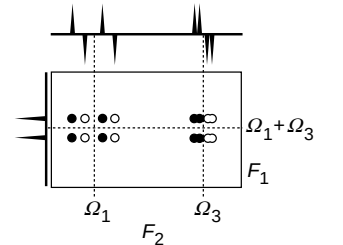
$$\begin{aligned} B_{13} \cos(\Omega_1 + \Omega_3) t_1 \text{DQ}_y^{(13)} &\xrightarrow{2\pi J_{12} t_1 I_{1z} I_{2z} + 2\pi J_{23} t_1 I_{2z} I_{3z}} \\ &B_{13} \cos(\Omega_1 + \Omega_3) t_1 \cos \pi (J_{12} + J_{23}) t_1 \text{DQ}_y^{(13)} \\ &- B_{13} \cos(\Omega_1 + \Omega_3) t_1 \sin \pi (J_{12} + J_{23}) t_1 2I_{2z} \text{DQ}_x^{(13)} \\ - B_{13} \sin(\Omega_1 + \Omega_3) t_1 \text{DQ}_x^{(13)} &\xrightarrow{2\pi J_{12} t_1 I_{1z} I_{2z} + 2\pi J_{23} t_1 I_{2z} I_{3z}} \\ &- B_{13} \sin(\Omega_1 + \Omega_3) t_1 \cos \pi (J_{12} + J_{23}) t_1 \text{DQ}_x^{(13)} \\ &- B_{13} \sin(\Omega_1 + \Omega_3) t_1 \sin \pi (J_{12} + J_{23}) t_1 2I_{2z} \text{DQ}_y^{(13)} \end{aligned}$$

Terms such as $2I_{2z} \text{DQ}_y^{(13)}$ and $2I_{2z} \text{DQ}_x^{(13)}$ can be thought of as double-quantum coherence which has become "anti-phase" with respect to the coupling to spin 2; such terms are directly analogous to single-quantum anti-phase magnetization.

Of all the terms present at the end of t_1 , only $\text{DQ}_y^{(13)}$ is rendered observable by the final pulse

$$\begin{aligned} \cos(\Omega_1 + \Omega_3) t_1 \cos \pi (J_{12} + J_{23}) t_1 B_{13} \text{DQ}_y^{(13)} &\xrightarrow{\pi/2(I_{1x} + I_{2x} + I_{3x})} \\ \cos(\Omega_1 + \Omega_3) t_1 \cos \pi (J_{12} + J_{23}) t_1 B_{13} [2I_{1x} I_{3x} + 2I_{1z} I_{3x}] \end{aligned}$$

The calculation predicts that two two-dimensional multiplets appear in the spectrum. Both have the same structure in F_1 , namely an in-phase doublet, split by $(J_{12} + J_{23})$ and centred at $(\Omega_1 + \Omega_3)$; this is analogous to a normal multiplet. In F_2 one two-dimensional multiplet is centred at the offset of spins 1, Ω_1 , and one at the offset of spin 3, Ω_3 ; both multiplets are anti-phase with respect to the coupling J_{13} . Finally, the overall amplitude, B_{13} , depends on the delay Δ and all the couplings in the system. The schematic spectrum is shown opposite. Similar multiplet structures are seen for the double-



Schematic two-dimensional double quantum spectrum showing the multiplets arising from evolution of double-quantum coherence between spins 1 and 3. It has been assumed that $J_{12} > J_{13} > J_{23}$.

quantum between spins 1 & 2 and spins 2 & 3.

7.5.2 Interpretation of double-quantum spectra

The double-quantum spectrum shows the relationship between the frequencies of the lines in the double quantum spectrum and those in the (conventional) single-quantum spectrum. If two two-dimensional multiplets appear at $(F_1, F_2) = (\Omega_A + \Omega_B, \Omega_A)$ and $(\Omega_A + \Omega_B, \Omega_B)$ the implication is that the two spins A and B are coupled, as it is only if there is a coupling present that double-quantum coherence between the two spins can be generated (e.g. in the previous section, if $J_{13} = 0$ the term B_{13} goes to zero). The fact that the two two-dimensional multiplets share a common F_1 frequency and that this frequency is the sum of the two F_2 frequencies constitute a double check as to whether or not the peaks indicate that the spins are coupled.

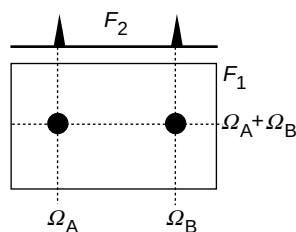
Double quantum spectra give very similar information to that obtained from COSY i.e. the identification of coupled spins. Each method has particular advantages and disadvantages:

- (1) In COSY the cross-peak multiplet is anti-phase in both dimensions, whereas in a double-quantum spectrum the multiplet is only anti-phase in F_2 . This may lead to stronger peaks in the double-quantum spectrum due to less cancellation. However, during the two delays Δ magnetization is lost by relaxation, resulting in reduced peak intensities in the double-quantum spectrum.
- (2) The value of the delay Δ in the double-quantum experiment affects the amount of multiple-quantum generated and hence the intensity in the spectrum. All of the couplings present in the spin system affect the intensity and as couplings cover a wide range, no single optimum value for Δ can be given. An unfortunate choice for Δ will result in low intensity, and it is then possible that correlations will be missed. No such problems occur with COSY.
- (3) There are no diagonal-peak multiplets in a double-quantum spectrum, so that correlations between spins with similar offsets are relatively easy to locate. In contrast, in a COSY the cross-peaks from such a pair of spins could be obscured by the diagonal.
- (4) In more complex spin systems the interpretation of a COSY remains unambiguous, but the double-quantum spectrum may show a peak with F_1 co-ordinate $(\Omega_A + \Omega_B)$ and F_2 co-ordinate Ω_A (or Ω_B) *even when spins A and B are not coupled*. Such *remote peaks*, as they are called, appear when spins A and B are *both* coupled to a third spin. There are various tests that can differentiate these remote from the more useful direct peaks, but these require additional experiments. The form of these remote peaks is considered in the next section.

On the whole, COSY is regarded as a more reliable and simple experiment, although double-quantum spectroscopy is used in some special circumstances.

7.5.3 Remote peaks in double-quantum spectra

The origin of remote peaks can be illustrated by returning to the calculation



Schematic spectrum showing the relationship between the single- and double-quantum frequencies for coupled spins.

of section 7.5.1. and focusing on the doubly anti-phase term which is present at the end of the spin echo (the fourth term in Eqn. [3.1])

$$\sin \pi J_{13} \Delta \sin \pi J_{12} \Delta 4 I_{1y} I_{2z} I_{3z}$$

The 90° pulse rotates this into multiple-quantum

$$\sin \pi J_{13} \Delta \sin \pi J_{12} \Delta 4 I_{1y} I_{2z} I_{3z} \xrightarrow{\pi/2(I_{1x}+I_{2x}+I_{3x})} \sin \pi J_{13} \Delta \sin \pi J_{12} \Delta 4 I_{1z} I_{2y} I_{3y}$$

The pure double-quantum part of this term is

$$-\frac{1}{2} \sin \pi J_{13} \Delta \sin \pi J_{12} \Delta (4 I_{1z} I_{2x} I_{3x} - 4 I_{1z} I_{2y} I_{3y}) \equiv B_{23,1} 2 I_{1z} DQ_x^{(23)}$$

In words, what has been generated is double-quantum between spins 2 and 3, anti-phase with respect to spin 1. The key thing is that no coupling between spins 2 and 3 is required for the generation of this term – the intensity just depends on J_{12} and J_{13} ; all that is required is that both spins 2 and 3 have a coupling to the third spin, spin 1.

During t_1 this term evolves under the influence of the offsets and the couplings. Only two terms ultimately lead to observable signals; at the end of t_1 these two terms are

$$B_{23,1} \cos(\Omega_2 + \Omega_3) t_1 \cos \pi (J_{12} + J_{13}) t_1 2 I_{1z} DQ_x^{(23)}$$

$$B_{23,1} \cos(\Omega_2 + \Omega_3) t_1 \sin \pi (J_{12} + J_{13}) t_1 DQ_y^{(23)}$$

and after the final 90° pulse the observable parts are

$$B_{23,1} \cos(\Omega_2 + \Omega_3) t_1 \cos \pi (J_{12} + J_{13}) t_1 4 I_{1y} I_{2z} I_{3z}$$

$$B_{23,1} \cos(\Omega_2 + \Omega_3) t_1 \sin \pi (J_{12} + J_{13}) t_1 (2 I_{2x} I_{3z} + 2 I_{2z} I_{3x})$$

The first term results in a multiplet appearing at Ω_1 in F_2 and at $(\Omega_2 + \Omega_3)$ in F_1 . The multiplet is doubly anti-phase (with respect to the couplings to spins 2 and 3) in F_2 ; in F_1 it is in-phase with respect to the sum of the couplings J_{12} and J_{13} . This multiplet is a remote peak, as its frequency coordinates do not conform to the simple pattern described in section 7.5.2. It is distinguished from direct peaks not only by its frequency coordinates, but also by having a different lineshape in F_2 to direct peaks and by being doubly anti-phase in that dimension.

The second and third terms are anti-phase with respect to the coupling between spins 2 and 3, and if this coupling is zero there will be cancellation within the multiplet and no signals will be observed. This is despite the fact that multiple-quantum coherence between these two spins has been generated.

7.6 Advanced topic:

Lineshapes and frequency discrimination

This is a somewhat involved topic which will only be possible to cover in outline here.

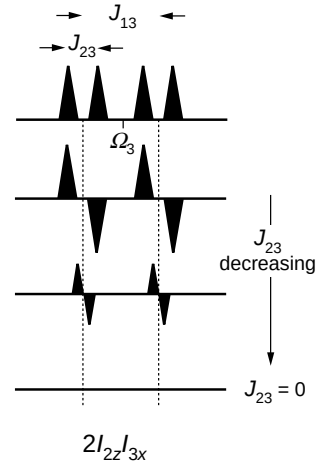


Illustration of how the intensity of an anti-phase multiplet decreases as the coupling which it is in anti-phase with respect to decreases. The in-phase multiplet is shown at the top, and below are three versions of the anti-phase multiplet for successively decreasing values of J_{23} .

All modern spectrometers use a method known as *quadrature detection*, which in effect means that both the x- and y-components of the magnetization are detected simultaneously.

7.6.1 One-dimensional spectra

Suppose that a $90^\circ(y)$ pulse is applied to equilibrium magnetization resulting in the generation of pure x-magnetization which then precesses in the transverse plane with frequency Ω . NMR spectrometers are set up to detect the x- and y-components of this magnetization. If it is assumed (arbitrarily) that these components decay exponentially with time constant T_2 the resulting signals, $S_x(t)$ and $S_y(t)$, from the two channels of the detector can be written

$$S_x(t) = \gamma \cos \Omega t \exp(-t/T_2) \quad S_y(t) = \gamma \sin \Omega t \exp(-t/T_2)$$

where γ is a factor which gives the absolute intensity of the signal.

Usually, these two components are combined in the computer to give a complex time-domain signal, $S(t)$

$$\begin{aligned} S(t) &= S_x(t) + iS_y(t) \\ &= \gamma(\cos \Omega t + i \sin \Omega t) \exp(-t/T_2) \\ &= \gamma \exp(i\Omega t) \exp(-t/T_2) \end{aligned} \quad [7.2]$$

The Fourier transform of $S(t)$ is also a complex function, $S(\omega)$:

$$\begin{aligned} S(\omega) &= FT[S(t)] \\ &= \gamma\{A(\omega) + iD(\omega)\} \end{aligned}$$

where $A(\omega)$ and $D(\omega)$ are the absorption and dispersion Lorentzian lineshapes:

$$A(\omega) = \frac{1}{(\omega - \Omega)^2 T_2^2 + 1} \quad D(\omega) = \frac{(\omega - \Omega)T_2}{(\omega - \Omega)^2 T_2^2 + 1}$$

These lineshapes are illustrated opposite. For NMR it is usual to display the spectrum with the absorption mode lineshape and in this case this corresponds to displaying the real part of $S(\omega)$.

7.6.1.1 Phase

Due to instrumental factors it is almost never the case that the real and imaginary parts of $S(t)$ correspond exactly to the x- and y-components of the magnetization. Mathematically, this is expressed by multiplying the ideal function by an instrumental phase factor, ϕ_{instr}

$$S(t) = \gamma \exp(i\phi_{\text{instr}}) \exp(i\Omega t) \exp(-t/T_2)$$

The real and imaginary parts of $S(t)$ are

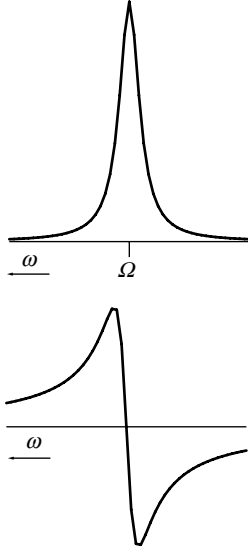
$$\begin{aligned} \text{Re}[S(t)] &= \gamma(\cos \phi_{\text{instr}} \cos \Omega t - \sin \phi_{\text{instr}} \sin \Omega t) \exp(-t/T_2) \\ \text{Im}[S(t)] &= \gamma(\cos \phi_{\text{instr}} \sin \Omega t + \sin \phi_{\text{instr}} \cos \Omega t) \exp(-t/T_2) \end{aligned}$$

Clearly, these do not correspond to the x- and y-components of the ideal time-domain function.

The Fourier transform of $S(t)$ carries forward the phase term

$$S(\omega) = \gamma \exp(i\phi_{\text{instr}}) \{A(\omega) + iD(\omega)\}$$

The real and imaginary parts of $S(\omega)$ are no longer the absorption and



Absorption (above) and dispersion (below) Lorentzian lineshapes, centred at frequency Ω .

dispersion signals:

$$\operatorname{Re}[S(\omega)] = \gamma(\cos \phi_{\text{instr}} A(\omega) - \sin \phi_{\text{instr}} D(\omega))$$

$$\operatorname{Im}[S(\omega)] = \gamma(\cos \phi_{\text{instr}} D(\omega) + \sin \phi_{\text{instr}} A(\omega))$$

Thus, displaying the real part of $S(\omega)$ will not give the required absorption mode spectrum; rather, the spectrum will show lines which have a mixture of absorption and dispersion lineshapes.

Restoring the pure absorption lineshape is simple. $S(\omega)$ is multiplied, in the computer, by a phase correction factor, ϕ_{corr} :

$$\begin{aligned} S(\omega) \exp(i\phi_{\text{corr}}) &= \gamma \exp(i\phi_{\text{corr}}) \exp(i\phi_{\text{instr}}) \{A(\omega) + iD(\omega)\} \\ &= \gamma \exp(i(\phi_{\text{corr}} + \phi_{\text{instr}})) \{A(\omega) + iD(\omega)\} \end{aligned}$$

By choosing ϕ_{corr} such that $(\phi_{\text{corr}} + \phi_{\text{instr}}) = 0$ (i.e. $\phi_{\text{corr}} = -\phi_{\text{instr}}$) the phase terms disappear and the real part of the spectrum will have the required absorption lineshape. In practice, the value of the phase correction is set "by eye" until the spectrum "looks phased". NMR processing software also allows for an additional phase correction which depends on frequency; such a correction is needed to compensate for, amongst other things, imperfections in radiofrequency pulses.

7.6.1.2 Phase is arbitrary

Suppose that the phase of the 90° pulse is changed from y to x . The magnetization now starts along $-y$ and precesses towards x ; assuming that the instrumental phase is zero, the output of the two channels of the detector are

$$S_x(t) = \gamma \sin \Omega t \exp(-t/T_2) \quad S_y(t) = -\gamma \cos \Omega t \exp(-t/T_2)$$

The complex time-domain signal can then be written

$$\begin{aligned} S(t) &= S_x(t) + iS_y(t) \\ &= \gamma(\sin \Omega t - i \cos \Omega t) \exp(-t/T_2) \\ &= \gamma(-i)(\cos \Omega t + i \sin \Omega t) \exp(-t/T_2) \\ &= \gamma(-i) \exp(i\Omega t) \exp(-t/T_2) \\ &= \gamma \exp(i\phi_{\text{exp}}) \exp(i\Omega t) \exp(-t/T_2) \end{aligned}$$

Where ϕ_{exp} , the "experimental" phase, is $-\pi/2$ (recall that $\exp(i\phi) = \cos \phi + i \sin \phi$, so that $\exp(-i\pi/2) = -i$).

It is clear from the form of $S(t)$ that this phase introduced by altering the experiment (in this case, by altering the phase of the pulse) takes exactly the same form as the instrumental phase error. It can, therefore, be corrected by applying a phase correction so as to return the real part of the spectrum to the absorption mode lineshape. In this case the phase correction would be $\pi/2$.

The Fourier transform of the original signal is

$$\begin{aligned} S(\omega) &= \gamma(-i)\{A(\omega) + iD(\omega)\} \\ \operatorname{Re}[S(\omega)] &= \gamma D(\omega) \quad \operatorname{Im}[S(\omega)] = -\gamma A(\omega) \end{aligned}$$

Thus the real part shows the dispersion mode lineshape, and the imaginary part shows the absorption lineshape. The 90° phase shift simply swaps over the real and imaginary parts.

7.6.1.3 *Relative phase is important*

The conclusion from the previous two sections is that the lineshape seen in the spectrum is under the control of the spectroscopist. It does not matter, for example, whether the pulse sequence results in magnetization appearing along the x - or y - axis (or anywhere in between, for that matter). It is always possible to phase correct the spectrum afterwards to achieve the desired lineshape.

However, if an experiment leads to magnetization from different processes or spins appearing along different axes, there is no single phase correction which will put the whole spectrum in the absorption mode. This is the case in the COSY spectrum (section 7.4.1). The terms leading to diagonal-peaks appear along the x -axis, whereas those leading to cross-peaks appear along y . Either can be phased to absorption, but if one is in absorption, one will be in dispersion; the two signals are fundamentally 90° out of phase with one another.

7.6.1.4 *Frequency discrimination*

Suppose that a particular spectrometer is only capable of recording one, say the x -, component of the precessing magnetization. The time domain signal will then just have a real part (compare Eqn. [7.2] in section 7.6.1)

$$S(t) = \gamma \cos \Omega t \exp(-t/T_2)$$

Using the identity $\cos \theta = \frac{1}{2}(\exp(i\theta) + \exp(-i\theta))$ this can be written

$$\begin{aligned} S(t) &= \frac{1}{2} \gamma [\exp(i\Omega t) + \exp(-i\Omega t)] \exp(-t/T_2) \\ &= \frac{1}{2} \gamma \exp(i\Omega t) \exp(-t/T_2) + \frac{1}{2} \gamma \exp(-i\Omega t) \exp(-t/T_2) \end{aligned}$$

The Fourier transform of the first term gives, in the real part, an absorption mode peak at $\omega = +\Omega$; the transform of the second term gives the same but at $\omega = -\Omega$.

$$\text{Re}[S(\omega)] = \frac{1}{2} \gamma A_+ + \frac{1}{2} \gamma A_-$$

where A_+ represents an absorption mode Lorentzian line at $\omega = +\Omega$ and A_- represents the same at $\omega = -\Omega$; likewise, D_+ and D_- represent dispersion mode peaks at $+\Omega$ and $-\Omega$, respectively.

This spectrum is said to lack *frequency discrimination*, in the sense that it does not matter if the magnetization went round at $+\Omega$ or $-\Omega$, the spectrum still shows peaks at both $+\Omega$ and $-\Omega$. This is in contrast to the case where both the x - and y -components are measured where *one* peak appears at either positive or negative ω depending on the sign of Ω .

The lack of frequency discrimination is associated with the signal being modulated by a cosine wave, which has the property that $\cos(\Omega t) = \cos(-\Omega t)$, as opposed to a complex exponential, $\exp(i\Omega t)$ which is sensitive to the sign of Ω . In one-dimensional spectroscopy it is virtually always possible to arrange for the signal to have this desirable complex phase modulation, but in the case of two-dimensional spectra it is almost always the case that the signal modulation in the t_1 dimension is of the form $\cos(\Omega t_1)$ and so such spectra are not naturally frequency discriminated in the F_1 dimension.

Suppose now that only the y -component of the precessing magnetization could be detected. The time domain signal will then be (compare Eqn. [3.2] in section 7.6.1)

$$S(t) = i\gamma \sin \Omega t \exp(-t/T_2)$$

Using the identity $\sin \theta = \frac{1}{2i} (\exp(i\theta) - \exp(-i\theta))$ this can be written

$$\begin{aligned} S(t) &= \frac{1}{2} \gamma [\exp(i\Omega t) - \exp(-i\Omega t)] \exp(-t/T_2) \\ &= \frac{1}{2} \gamma \exp(i\Omega t) \exp(-t/T_2) - \frac{1}{2} \gamma \exp(-i\Omega t) \exp(-t/T_2) \end{aligned}$$

and so

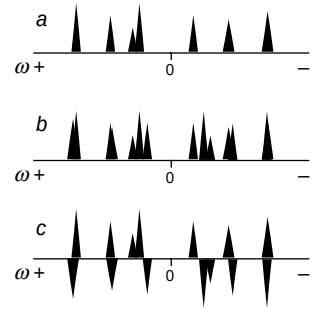
$$\text{Re}[S(\omega)] = \frac{1}{2} \gamma A_+ - \frac{1}{2} \gamma A_-$$

This spectrum again shows two peaks, at $\pm\Omega$, but the two peaks have opposite signs; this is associated with the signal being modulated by a sine wave, which has the property that $\sin(-\Omega t) = -\sin(\Omega t)$. If the sign of Ω changes the two peaks swap over, but there are still two peaks. In a sense the spectrum is frequency discriminated, as positive and negative frequencies can be distinguished, but in practice in a spectrum with many lines with a range of positive and negative offsets the resulting set of possibly cancelling peaks would be impossible to sort out satisfactorily.

7.6.2 Two-dimensional spectra

7.6.2.1 Phase and amplitude modulation

There are two basic types of time-domain signal that are found in two-dimensional experiments. The first is *phase modulation*, in which the evolution in t_1 is encoded as a phase, *i.e.* mathematically as a complex



Spectrum a has peaks at positive and negative frequencies and is frequency discriminated. Spectrum b results from a cosine modulated time-domain data set; each peak appears at both positive and negative frequency, regardless of whether its real offset is positive or negative. Spectrum c results from a sine modulated data set; like b each peak appears twice, but with the added complication that one peak is inverted. Spectra b and c lack frequency discrimination and are quite uninterpretable as a result.

exponential

$$S(t_1, t_2)_{\text{phase}} = \gamma \exp(i\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)})$$

where Ω_1 and Ω_2 are the modulation frequencies in t_1 and t_2 respectively, and $T_2^{(1)}$ and $T_2^{(2)}$ are the decay time constants in t_1 and t_2 respectively.

The second type is *amplitude modulation*, in which the evolution in t_1 is encoded as an amplitude, i.e. mathematically as sine or cosine

$$S(t)_c = \gamma \cos(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)})$$

$$S(t)_s = \gamma \sin(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)})$$

Generally, two-dimensional experiments produce amplitude modulation, indeed all of the experiments analysed in this chapter have produced either sine or cosine modulated data. Therefore most two-dimensional spectra are fundamentally not frequency discriminated in the F_1 dimension. As explained above for one-dimensional spectra, the resulting confusion in the spectrum is not acceptable and steps have to be taken to introduce frequency discrimination.

It will turn out that the key to obtaining frequency discrimination is the ability to record, in separate experiments, both sine and cosine modulated data sets. This can be achieved by simply altering the phase of the pulses in the sequence.

For example, consider the EXSY sequence analysed in section 7.2 . The observable signal, at time $t_2 = 0$, can be written

$$(1 - f) \cos \Omega_1 t_1 I_{1y} + f \cos \Omega_1 t_1 I_{2y}$$

If, however, the first pulse in the sequence is changed in phase from x to y the corresponding signal will be

$$-(1 - f) \sin \Omega_1 t_1 I_{1y} - f \sin \Omega_1 t_1 I_{2y}$$

i.e. the modulation has changed from the form of a cosine to sine. In COSY and DQF COSY a similar change can be brought about by altering the phase of the first 90° pulse. In fact there is a general procedure for effecting this change, the details of which are given in a later chapter.

7.6.2.2 Two-dimensional lineshapes

The spectra resulting from two-dimensional Fourier transformation of phase and amplitude modulated data sets can be determined by using the following Fourier pair

$$FT[\exp(i\Omega t) \exp(-t/T_2)] = \{A(\omega) + iD(\omega)\}$$

where A and D are the dispersion Lorentzian lineshapes described in section 7.6.1

Phase modulation

For the phase modulated data set the transform with respect to t_2 gives

$$S(t_1, \omega_2)_{\text{phase}} = \gamma \exp(i\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) [A_+^{(2)} + iD_+^{(2)}]$$

where $A_+^{(2)}$ indicates an absorption mode line in the F_2 dimension at $\omega_2 =$

$+\Omega_2$ and with linewidth set by $T_2^{(2)}$; similarly $D_+^{(2)}$ is the corresponding dispersion line.

The second transform with respect to t_1 gives

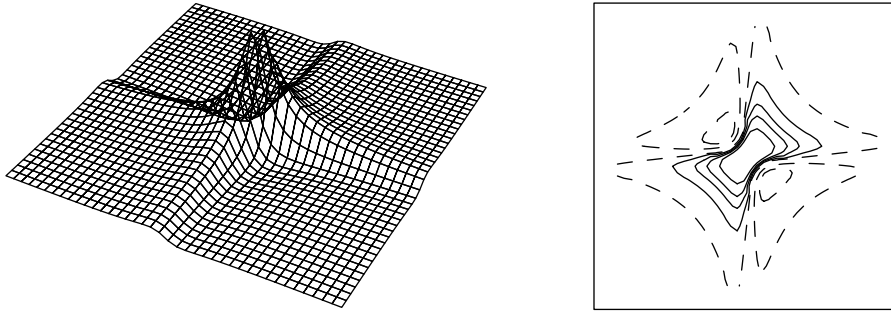
$$S(\omega_1, \omega_2)_{\text{phase}} = \gamma [A_+^{(1)} + iD_+^{(1)}] [A_+^{(2)} + iD_+^{(2)}]$$

where $A_+^{(1)}$ indicates an absorption mode line in the F_1 dimension at $\omega_1 = +\Omega_1$ and with linewidth set by $T_2^{(1)}$; similarly $D_+^{(1)}$ is the corresponding dispersion line.

The real part of the resulting two-dimensional spectrum is

$$\text{Re}[S(\omega_1, \omega_2)_{\text{phase}}] = \gamma (A_+^{(1)} A_+^{(2)} - D_+^{(1)} D_+^{(2)})$$

This is a single line at $(\omega_1, \omega_2) = (+\Omega_1, +\Omega_2)$ with the *phase-twist lineshape*, illustrated below.



Pseudo 3D view and contour plot of the phase-twist lineshape.

The phase-twist lineshape is an inextricable mixture of absorption and dispersion; it is a superposition of the double absorption and double dispersion lineshape (illustrated in section 7.4.1). No phase correction will restore it to pure absorption mode. Generally the phase twist is not a very desirable lineshape as it has both positive and negative parts, and the dispersion component only dies off slowly.

Cosine amplitude modulation

For the cosine modulated data set the transform with respect to t_2 gives

$$S(t_1, \omega_2)_c = \gamma \cos(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) [A_+^{(2)} + iD_+^{(2)}]$$

The cosine is then rewritten in terms of complex exponentials to give

$$S(t_1, \omega_2)_c = \frac{1}{2} \gamma [\exp(i\Omega_1 t_1) + \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) [A_+^{(2)} + iD_+^{(2)}]$$

The second transform with respect to t_1 gives

$$S(\omega_1, \omega_2)_c = \frac{1}{2} \gamma [\{A_+^{(1)} + iD_+^{(1)}\} + \{A_-^{(1)} + iD_-^{(1)}\}] [A_+^{(2)} + iD_+^{(2)}]$$

where $A_-^{(1)}$ indicates an absorption mode line in the F_1 dimension at $\omega_1 = -\Omega_1$ and with linewidth set by $T_2^{(1)}$; similarly $D_-^{(1)}$ is the corresponding dispersion line.

The real part of the resulting two-dimensional spectrum is

$$\text{Re}[S(\omega_1, \omega_2)_c] = \frac{1}{2} \gamma (A_+^{(1)} A_+^{(2)} - D_+^{(1)} D_+^{(2)}) + \frac{1}{2} \gamma (A_-^{(1)} A_+^{(2)} - D_-^{(1)} D_+^{(2)})$$

This is a two lines, both with the phase-twist lineshape; one is located at

$(+\Omega_1, +\Omega_2)$ and the other is at $(-\Omega_1, +\Omega_2)$. As expected for a data set which is cosine modulated in t_1 the spectrum is symmetrical about $\omega_1 = 0$.

A spectrum with a pure absorption mode lineshape can be obtained by discarding the imaginary part of the time domain data immediately after the transform with respect to t_2 ; i.e. taking the real part of $S(t_1, \omega_2)_c$

$$\begin{aligned} S(t_1, \omega_2)_c^{\text{Re}} &= \text{Re}[S(t_1, \omega_2)_c] \\ &= \gamma \cos(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) A_+^{(2)} \end{aligned}$$

Following through the same procedure as above:

$$\begin{aligned} S(t_1, \omega_2)_c^{\text{Re}} &= \frac{1}{2} \gamma [\exp(i\Omega_1 t_1) + \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) A_+^{(2)} \\ S(\omega_1, \omega_2)_c^{\text{Re}} &= \frac{1}{2} \gamma [\{A_+^{(1)} + iD_+^{(1)}\} + \{A_-^{(1)} + iD_-^{(1)}\}] A_+^{(2)} \end{aligned}$$

The real part of the resulting two-dimensional spectrum is

$$\text{Re}[S(\omega_1, \omega_2)_c^{\text{Re}}] = \frac{1}{2} \gamma A_+^{(1)} A_+^{(2)} + \frac{1}{2} \gamma A_-^{(1)} A_+^{(2)}$$

This is two lines, located at $(+\Omega_1, +\Omega_2)$ and $(-\Omega_1, +\Omega_2)$, but in contrast to the above both have the double absorption lineshape. There is still lack of frequency discrimination, but the undesirable phase-twist lineshape has been avoided.

Sine amplitude modulation

For the sine modulated data set the transform with respect to t_2 gives

$$S(t_1, \omega_2)_s = \gamma \sin(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) [A_+^{(2)} + iD_+^{(2)}]$$

The cosine is then rewritten in terms of complex exponentials to give

$$S(t_1, \omega_2)_s = \frac{1}{2i} \gamma [\exp(i\Omega_1 t_1) - \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) [A_+^{(2)} + iD_+^{(2)}]$$

The second transform with respect to t_1 gives

$$S(\omega_1, \omega_2)_s = \frac{1}{2i} \gamma [\{A_+^{(1)} + iD_+^{(1)}\} - \{A_-^{(1)} + iD_-^{(1)}\}] [A_+^{(2)} + iD_+^{(2)}]$$

The imaginary part of the resulting two-dimensional spectrum is

$$\text{Im}[S(\omega_1, \omega_2)_s] = -\frac{1}{2} \gamma (A_+^{(1)} A_+^{(2)} - D_+^{(1)} D_+^{(2)}) + \frac{1}{2} \gamma (A_-^{(1)} A_+^{(2)} - D_-^{(1)} D_+^{(2)})$$

This is two lines, both with the phase-twist lineshape but with opposite signs; one is located at $(+\Omega_1, +\Omega_2)$ and the other is at $(-\Omega_1, +\Omega_2)$. As expected for a data set which is sine modulated in t_1 the spectrum is anti-symmetric about $\omega_1 = 0$.

As before, a spectrum with a pure absorption mode lineshape can be obtained by discarding the imaginary part of the time domain data immediately after the transform with respect to t_2 ; i.e. taking the real part of $S(t_1, \omega_2)_s$

$$\begin{aligned} S(t_1, \omega_2)_s^{\text{Re}} &= \text{Re}[S(t_1, \omega_2)_s] \\ &= \gamma \sin(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) A_+^{(2)} \end{aligned}$$

Following through the same procedure as above:

$$S(t_1, \omega_2)_s^{\text{Re}} = \frac{1}{2i} \gamma [\exp(i\Omega_1 t_1) - \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) A_+^{(2)}$$

$$S(\omega_1, \omega_2)_s^{\text{Re}} = \frac{1}{2i} \gamma [\{A_+^{(1)} + iD_+^{(1)}\} - \{A_-^{(1)} + iD_-^{(1)}\}] A_+^{(2)}$$

The imaginary part of the resulting two-dimensional spectrum is

$$\text{Im}[S(\omega_1, \omega_2)_s^{\text{Re}}] = -\frac{1}{2} \gamma A_+^{(1)} A_+^{(2)} + \frac{1}{2} \gamma A_-^{(1)} A_+^{(2)}$$

The two lines now have the pure absorption lineshape.

7.6.2.3 Frequency discrimination with retention of absorption lineshapes

It is essential to be able to combine frequency discrimination in the F_1 dimension with retention of pure absorption lineshapes. Three different ways of achieving this are commonly used; each will be analysed here.

States-Haberkorn-Ruben method

The essence of the States-Haberkorn-Ruben (SHR) method is the observation that the cosine modulated data set, processed as described in section 7.6.1.2, gives two positive absorption mode peaks at $(+\Omega_1, +\Omega_2)$ and $(-\Omega_1, +\Omega_2)$, whereas the sine modulated data set processed in the same way gives a spectrum in which one peak is negative and one positive. Subtracting these spectra from one another gives the required absorption mode frequency discriminated spectrum (see the diagram below):

$$\begin{aligned} \text{Re}[S(\omega_1, \omega_2)_c^{\text{Re}}] - \text{Im}[S(\omega_1, \omega_2)_s^{\text{Re}}] \\ = \left[\frac{1}{2} \gamma A_+^{(1)} A_+^{(2)} + \frac{1}{2} \gamma A_-^{(1)} A_+^{(2)} \right] - \left[-\frac{1}{2} \gamma A_+^{(1)} A_+^{(2)} + \frac{1}{2} \gamma A_-^{(1)} A_+^{(2)} \right] \\ = \gamma A_+^{(1)} A_+^{(2)} \end{aligned}$$

In practice it is usually more convenient to achieve this result in the following way, which is mathematically identical.

The cosine and sine data sets are transformed with respect to t_2 and the real parts of each are taken. Then a new complex data set is formed using the cosine data for the real part and the sine data for the imaginary part:

$$\begin{aligned} S(t_1, \omega_2)_{\text{SHR}} &= S(t_1, \omega_2)_c^{\text{Re}} + i S(t_1, \omega_2)_s^{\text{Re}} \\ &= \gamma \cos(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) A_+^{(2)} + i \gamma \sin(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) A_+^{(2)} \\ &= \gamma \exp(i\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) A_+^{(2)} \end{aligned}$$

Fourier transformation with respect to t_1 gives a spectrum whose real part contains the required frequency discriminated absorption mode spectrum

$$\begin{aligned} S(\omega_1, \omega_2)_{\text{SHR}} &= \gamma [A_+^{(1)} + iD_+^{(1)}] A_+^{(2)} \\ &= \gamma A_+^{(1)} A_+^{(2)} + iD_+^{(1)} A_+^{(2)} \end{aligned}$$

Marion-Wüthrich or TPPI method

The idea behind the TPPI (time proportional phase incrementation) or Marion-Wüthrich (MW) method is to arrange things so that all of the peaks have positive offsets. Then, frequency discrimination would not be required as there would be no ambiguity.

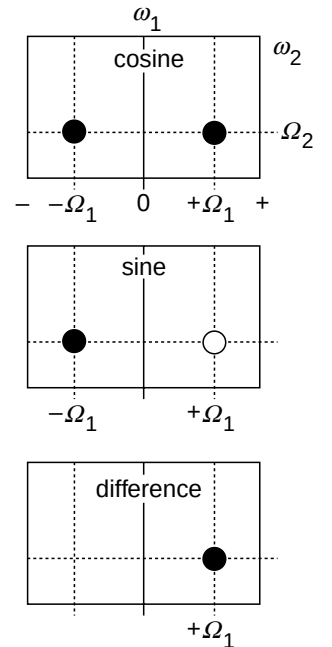


Illustration of the way in which the SHR method achieves frequency discrimination by combining cosine and sine modulated spectra.

One simple way to make all offsets positive is to set the receiver carrier frequency deliberately at the edge of the spectrum. Simple though this is, it is not really a very practical method as the resulting spectrum would be very inefficient in its use of data space and in addition off-resonance effects associated with the pulses in the sequence will be accentuated.

In the TPPI method the carrier can still be set in the middle of the spectrum, but it is made to *appear* that all the frequencies are positive by phase shifting systematically some of the pulses in the sequence in concert with the incrementation of t_1 .

In section 7.2 it was shown that in the EXSY sequence the cosine modulation in t_1 , $\cos(\Omega_1 t_1)$, could be turned into sine modulation, $-\sin(\Omega_1 t_1)$, by shifting the phase of the first pulse by 90° . The effect of such a phase shift can be represented mathematically in the following way.

Recall that Ω is in units of radians s^{-1} , and so if t is in seconds Ωt is in radians; Ωt can therefore be described as a phase which depends on time. It is also possible to consider phases which do not depend on time, as was the case for the phase errors considered in section 7.6.1.1

The change from cosine to sine modulation in the EXSY experiment can be thought of as a phase shift of the signal in t_1 . Mathematically, such a phase shifted cosine wave is written as $\cos(\Omega_1 t_1 + \phi)$, where ϕ is the phase shift in radians. This expression can be expanded using the well known formula $\cos(A + B) = \cos A \cos B - \sin A \sin B$ to give

$$\cos(\Omega_1 t_1 + \phi) = \cos \Omega_1 t_1 \cos \phi - \sin \Omega_1 t_1 \sin \phi$$

If the phase shift, ϕ , is $\pi/2$ radians the result is

$$\begin{aligned} \cos(\Omega_1 t_1 + \pi/2) &= \cos \Omega_1 t_1 \cos \pi/2 - \sin \Omega_1 t_1 \sin \pi/2 \\ &= -\sin \Omega_1 t_1 \end{aligned}$$

In words, a cosine wave, phase shifted by $\pi/2$ radians (90°) is the same thing as a sine wave. Thus, in the EXSY experiment the effect of changing the phase of the first pulse by 90° can be described as a phase shift of the signal by 90° .

Suppose that instead of a fixed phase shift, the phase shift is made proportional to t_1 ; what this means is that each time t_1 is incremented the phase is also incremented in concert. The constant of proportion between the time dependent phase, $\phi(t_1)$, and t_1 will be written $\omega_{\text{additional}}$

$$\phi(t_1) = \omega_{\text{additional}} t_1$$

Clearly the units of $\omega_{\text{additional}}$ are radians s^{-1} , that is $\omega_{\text{additional}}$ is a frequency. The new time-domain function with the inclusion of this incrementing phase is thus

$$\begin{aligned} \cos(\Omega_1 t_1 + \phi(t_1)) &= \cos(\Omega_1 t_1 + \omega_{\text{additional}} t_1) \\ &= \cos(\Omega_1 + \omega_{\text{additional}}) t_1 \end{aligned}$$

In words, the effect of incrementing the phase in concert with t_1 is to add a frequency $\omega_{\text{additional}}$ to *all* of the offsets in the spectrum. The TPPI method utilizes this option of shifting all the frequencies in the following way.

In one-dimensional pulse-Fourier transform NMR the free induction signal is sampled at regular intervals Δ . After transformation the resulting spectrum displays correctly peaks with offsets in the range $-(SW/2)$ to $+(SW/2)$ where SW is the spectral width which is given by $1/\Delta$ (this comes about from the Nyquist theorem of data sampling). Frequencies outside this range are not represented correctly.

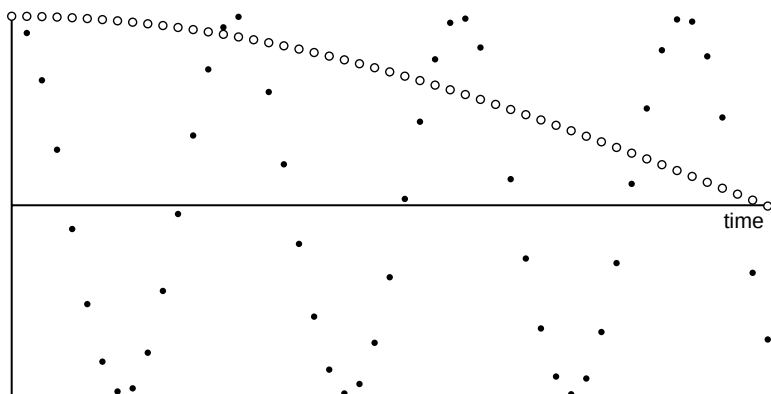
Suppose that the required frequency range in the F_1 dimension is from $-(SW_1/2)$ to $+(SW_1/2)$ (in COSY and EXSY this will be the same as the range in F_2). To make it appear that all the peaks have a positive offset, it will be necessary to add $(SW_1/2)$ to all the frequencies. Then the peaks will be in the range 0 to (SW_1) .

As the maximum frequency is now (SW_1) rather than $(SW_1/2)$ the sampling interval, Δ_1 , will have to be halved *i.e.* $\Delta_1 = 1/(2SW_1)$ in order that the range of frequencies present are represented properly.

The phase increment is $\omega_{\text{additional}} t_1$, but t_1 can be written as $n\Delta_1$ for the n th increment of t_1 . The required value for $\omega_{\text{additional}}$ is $2\pi(SW_1/2)$, where the 2π is to convert from frequency (the units of SW_1) to rad s^{-1} , the units of $\omega_{\text{additional}}$. Putting all of this together $\omega_{\text{additional}} t_1$ can be expressed, for the n th increment as

$$\begin{aligned}\omega_{\text{additional}} t_1 &= 2\pi \left(\frac{SW_1}{2} \right) (n\Delta_1) \\ &= 2\pi \left(\frac{SW_1}{2} \right) \left(n \frac{1}{2SW_1} \right) \\ &= n \frac{\pi}{2}\end{aligned}$$

The way in which the phase incrementation increases the frequency of the cosine wave is shown below:



The open circles lie on a cosine wave, $\cos(\Omega \times n\Delta)$, where Δ is the sampling interval and n runs 0, 1, 2 ... The closed circles lie on a cosine wave in which an additional phase is incremented on each point *i.e.* the function is $\cos(\Omega \times n\Delta + n\phi)$; here $\phi = \pi/8$. The way in which this phase increment increases the frequency of the cosine wave is apparent.

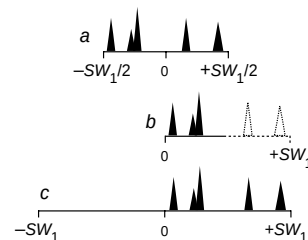
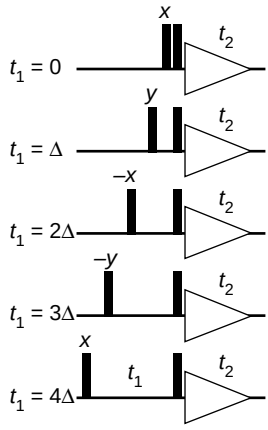


Illustration of the TPPI method. The normal spectrum is shown in a, with peaks in the range $-SW/2$ to $+SW/2$. Adding a frequency of $SW/2$ to all the peaks gives them all positive offsets, but some, shown dotted) will then fall outside the spectral window – spectrum b. If the spectral width is doubled all peaks are represented correctly – spectrum c.



TPPI phase incrementation applied to a COSY sequence. The phase of the first pulse is incremented by 90° each time t_1 is incremented.

In words this means that each time t_1 is incremented, the phase of the signal should also be incremented by 90°, for example by incrementing the phase of one of the pulses. The way in which it can be decided which pulse to increment will be described in a later chapter.

A data set from an experiment to which TPPI has been applied is simply amplitude modulated in t_1 and so can be processed according to the method described for cosine modulated data so as to obtain absorption mode lineshapes. As the spectrum is symmetrical about $F_1 = 0$ it is usual to use a modified Fourier transform routine which saves effort and space by only calculating the positive frequency part of the spectrum.

Echo anti-echo method

Few two-dimensional experiments naturally produce phase modulated data sets, but if gradient pulses are used for coherence pathway selection it is then quite often found that the data are phase modulated. In one way this is an advantage, as it means that no special steps are required to obtain frequency discrimination. However, phase modulated data sets give rise to spectra with phase-twist lineshapes, which are very undesirable. So, it is usual to attempt to use some method to eliminate the phase-twist lineshape, while at the same time retaining frequency discrimination.

The key to how this can be done lies in the fact that two kinds of phase modulated data sets can usually be recorded. The first is called the *P*-type or anti-echo spectrum

$$S(t_1, t_2)_P = \gamma \exp(i\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)})$$

the "*P*" indicates positive, meaning here that the sign of the frequencies in F_1 and F_2 are the same.

The second data set is called the echo or *N*-type

$$S(t_1, t_2)_N = \gamma \exp(-i\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)})$$

the "*N*" indicates negative, meaning here that the sign of the frequencies in F_1 and F_2 are opposite. As will be explained in a later chapter in gradient experiments it is easy to arrange to record either the *P*- or *N*-type spectrum.

The simplest way to proceed is to compute two new data sets which are the sum and difference of the *P*- and *N*-type data sets:

$$\begin{aligned} \frac{1}{2} [S(t_1, t_2)_P + S(t_1, t_2)_N] &= \\ \frac{1}{2} \gamma [\exp(i\Omega_1 t_1) + \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)}) &= \\ = \gamma \cos(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)}) \end{aligned}$$

$$\begin{aligned} \frac{1}{2i} [S(t_1, t_2)_P - S(t_1, t_2)_N] &= \\ \frac{1}{2i} \gamma [\exp(i\Omega_1 t_1) - \exp(-i\Omega_1 t_1)] \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)}) &= \\ = \gamma \sin(\Omega_1 t_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2) \exp(-t_2/T_2^{(2)}) \end{aligned}$$

These two combinations are just the cosine and sine modulated data sets that

are the inputs needed for the SHR method. The pure absorption spectrum can therefore be calculated in the same way starting with these combinations.

7.6.2.4 Phase in two-dimensional spectra

In practice there will be instrumental and other phase shifts, possibly in both dimensions, which mean that the time-domain functions are not the idealised ones treated above. For example, the cosine modulated data set might be

$$S(t)_c = \gamma \cos(\Omega_1 t_1 + \phi_1) \exp(-t_1/T_2^{(1)}) \exp(i\Omega_2 t_2 + i\phi_2) \exp(-t_2/T_2^{(2)})$$

where ϕ_1 and ϕ_2 are the phase errors in F_1 and F_2 , respectively. Processing this data set in the manner described above will not give a pure absorption spectrum. However, it is possible to recover the pure absorption spectrum by software manipulations of the spectrum, just as was described for the case of one-dimensional spectra. Usually, NMR data processing software provides options for making such phase corrections to two-dimensional data sets.

8 Relaxation[†]

Relaxation is the process by which the spins in the sample come to equilibrium with the surroundings. At a practical level, the rate of relaxation determines how fast an experiment can be repeated, so it is important to understand how relaxation rates can be measured and the factors that influence their values. The rate of relaxation is influenced by the physical properties of the molecule and the sample, so a study of relaxation phenomena can lead to information on these properties. Perhaps the most often used and important of these phenomena in the nuclear Overhauser effect (NOE) which can be used to probe internuclear distances in a molecule. Another example is the use of data on relaxation rates to probe the internal motions of macromolecules.

In this chapter the language and concepts used to describe relaxation will be introduced and illustrated. To begin with it will simply be taken for granted that there are processes which give rise to relaxation and we will not concern ourselves with the source of relaxation. Having described the experiments which can be used to probe relaxation we will then go on to see what the source of relaxation is and how it depends on molecular parameters and molecular motion.

8.1 What is relaxation?

Relaxation is the process by which the spins return to *equilibrium*. Equilibrium is the state in which (a) the populations of the energy levels are those predicted by the Boltzmann distribution and (b) there is no transverse magnetization and, more generally, no coherences present in the system.

In Chapter 3 we saw that when an NMR sample is placed in a static magnetic field and allowed to come to equilibrium it is found that a net magnetization of the sample along the direction of the applied field (traditionally the z-axis) is developed. Magnetization parallel to the applied field is termed longitudinal.

This equilibrium magnetization arises from the unequal population of the two energy levels that correspond to the α and β spin states. In fact, the z-magnetization, M_z , is proportional to the population difference

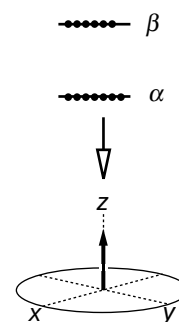
$$M_z \propto (n_\alpha - n_\beta)$$

where n_α and n_β are the populations of the two corresponding energy levels. Ultimately, the constant of proportion just determines the absolute size of the signal we will observe. As we are generally interested in the *relative* size of magnetizations and signals we may just as well write

$$M_z = (n_\alpha - n_\beta) \quad [2]$$

8.2 Rate equations and rate constants

The populations of energy levels are in many ways analogous to



A population difference gives rise to net magnetization along the z-axis.

[†] Chapter 8 "Relaxation" © James Keeler 1999, 2002 and 2004

concentrations in chemical kinetics, and many of the same techniques that are used to describe the rates of chemical reactions can also be used to describe the dynamics of populations. This will lead to a description of the dynamics of the z-magnetization.

Suppose that the populations of the α and β states at time t are n_α and n_β respectively. If these are not the equilibrium values, then for the system to reach equilibrium the population of one level must increase and that of the other must decrease. This implies that there must be *transitions* between the two levels *i.e.* something must happen which causes a spin to move from the α state to the β state or *vice versa*. It is this process which results in relaxation.

The simplest assumption that we can make about the rate of transitions from α to β is that it is proportional to the population of state α , n_α , and is a first order process with rate constant W . With these assumptions the rate of loss of population from state α is Wn_α . In the same way, the rate of loss of population of state β is Wn_β .

However, the key thing to realize is that whereas a transition from α to β causes a *loss* of population of level α , a transition from β to α causes the population of state α to *increase*. So we can write

$$\text{rate of change of population of state } \alpha = -Wn_\alpha + Wn_\beta$$

The first term is negative as it represents a loss of population of state α and in contrast the second term is positive as it represents a gain in the population of state α . The rate of change of the population can be written using the language of calculus as dn_α/dt so we have

$$\frac{dn_\alpha}{dt} = -Wn_\alpha + Wn_\beta .$$

Similarly we can write for the population of state β :

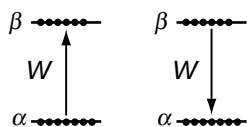
$$\frac{dn_\beta}{dt} = -Wn_\beta + Wn_\alpha .$$

These equations are almost correct, but we need to make one modification. At equilibrium the populations are not changing so $dn_\alpha/dt = 0$; this immediately implies that at equilibrium $n_\alpha = n_\beta$, which simply is not correct. We know that at equilibrium the population of state α exceeds that of state β . This defect is easily remedied by replacing the population n_α with the *deviation* of the population from its equilibrium value $(n_\alpha - n_\alpha^0)$, where n_α^0 is the population of state α at equilibrium. Doing the same with state β gives us the final, correct, equations:

$$\frac{dn_\alpha}{dt} = W(n_\beta - n_\beta^0) - W(n_\alpha - n_\alpha^0) \quad \frac{dn_\beta}{dt} = W(n_\alpha - n_\alpha^0) - W(n_\beta - n_\beta^0) .$$

You will recognise here that this kind of approach is exactly the same as that used to analyse the kinetics of a reversible chemical reaction.

Using Eqn. [2], we can use these two equations to work out how the z-magnetization varies with time:



A transition from state α to state β decreases the population of state α , but a transition from state β to state α increases the population of state α .

$$\begin{aligned}
\frac{dM_z}{dt} &= \frac{d(n_\alpha - n_\beta)}{dt} \\
&= \frac{dn_\alpha}{dt} - \frac{dn_\beta}{dt} \\
&= W(n_\beta - n_\beta^0) - W(n_\alpha - n_\alpha^0) - W(n_\alpha - n_\alpha^0) + W(n_\beta - n_\beta^0) \\
&= -2W(n_\alpha - n_\beta) + 2W(n_\alpha^0 - n_\beta^0) \\
&= -2W(M_z - M_z^0)
\end{aligned}$$

where $M_z^0 = (n_\alpha^0 - n_\beta^0)$, the equilibrium z-magnetization.

8.2.1 Consequences of the rate equation

The discussion in the previous section led to a (differential) equation describing the motion of the z-magnetization

$$\frac{dM_z(t)}{dt} = -R_z(M_z(t) - M_z^0) \quad [7]$$

where the rate constant, $R_z = 2W$ and M_z has been written as a function of time, $M_z(t)$, to remind us that it may change.

What this equations says is that the rate of change of M_z is proportional to the deviation of M_z from its equilibrium value, M_z^0 . If $M_z = M_z^0$, that is the system is at equilibrium, the right-hand side of Eqn. [7] is zero and hence so is the rate of change of M_z ; nothing happens. On the other hand, if M_z deviates from M_z^0 there will be a rate of change of M_z , and this rate will be proportional to the deviation of M_z from M_z^0 . The change will also be such as to return M_z to its equilibrium value, M_z^0 . In summary, Eqn. [7] predicts that over time M_z will return to M_z^0 ; this is exactly what we expect. The rate at which this happens will depend on R_z .

This equation can easily be integrated:

$$\begin{aligned}
\int \frac{dM_z(t)}{(M_z(t) - M_z^0)} &= \int -R_z dt \\
\ln(M_z(t) - M_z^0) &= -R_z t + \text{const.}
\end{aligned}$$

If, at time zero, the magnetization is $M_z(0)$, the constant of integration can be determined as $\ln(M_z(0) - M_z^0)$. Hence, with some rearrangement:

$$\ln \left[\frac{M_z(t) - M_z^0}{M_z(0) - M_z^0} \right] = -R_z t \quad [7A]$$

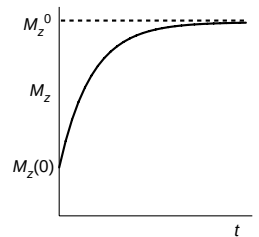
or

$$M_z(t) = [M_z(0) - M_z^0] \exp(-R_z t) + M_z^0 \quad [7B]$$

In words, this says that the z-magnetization returns from $M_z(0)$ to M_z^0 following an exponential law. The time constant of the exponential is $1/R_z$, and this is often called T_1 , the *longitudinal* or *spin-lattice relaxation time*.

8.2.2 The inversion recovery experiment

We described this experiment in section 3.10. First, a 180° pulse is applied,



thereby inverting the magnetization. Then a delay t is left for the magnetization to relax. Finally, a 90° pulse is applied so that the size of the z-magnetization can be measured.

We can now analyse this experiment fully. The starting condition for M_z is $M_z(0) = -M_z^0$, i.e. inversion. With this condition the predicted time evolution can be found from Eqn. [7A] to be:

$$\ln \left[\frac{M_z(t) - M_z^0}{-2M_z^0} \right] = -R_z t$$

Recall that the amplitude of the signal we record in this experiment is proportional to the z-magnetization. So, if this signal is $S(t)$ it follows that

$$\ln \left[\frac{S(t) - S^0}{-2S^0} \right] = -R_z t \quad [7C]$$

where S^0 is the signal intensity from equilibrium magnetization; we would find this from a simple 90° – acquire experiment.

Equation [7C] implies that a plot of $\ln[(S(t) - S^0)/-2S^0]$ against t should be a straight line of slope $-R_z$. This, then, is the basis of a method of determining the relaxation rate constant.

8.2.3 A quick estimate for R_z (or T_1)

Often we want to obtain a quick estimate for the relaxation rate constant (or, equivalently, the relaxation time). One way to do this is to do an inversion recovery experiment but rather than varying t systematically we look for the value of t which results in no signal i.e. a null. If the time when $S(t)$ is zero is t_{null} it follows immediately from Eqn. [7C] that:

$$\ln \left[\frac{1}{2} \right] = -R_z t_{\text{null}} \quad \text{or} \quad R_z = \frac{\ln 2}{t_{\text{null}}} \quad \text{or} \quad T_1 = \frac{t_{\text{null}}}{\ln 2}$$

Probably the most useful relationship is the last, which is $T_1 \approx 1.4 t_{\text{null}}$.

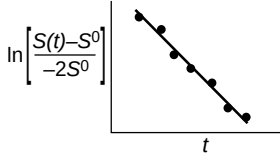
This method is rather crude, but it good enough for estimating T_1 . Armed with this estimate we can then, for example, decide on the time to leave between transients (typically three to five times T_1).

8.2.4 Writing relaxation in terms of operators

As we saw in Chapter 6, in quantum mechanics z-magnetization is represented by the operator I_z . It is therefore common to write Eqn. [7] in terms of operators rather than magnetizations, to give:

$$\frac{dI_z(t)}{dt} = -R_z (I_z(t) - I_z^0) \quad [8]$$

where $I_z(t)$ represents the z-magnetization at time t and I_z^0 represents the equilibrium z-magnetization. As it stands this last equation seems to imply that the operators change with time, which is not what is meant. What are changing are the populations of the energy levels and these in turn lead to changes in the z-magnetization represented by the operator. We will use this notation from now on.



Plot used to extract a value of R_z from the data from an inversion recovery experiment.

8.3 Solomon equations

The idea of writing differential equations for the populations, and then transcribing these into magnetizations, is a particularly convenient way of describing relaxation, especially in more complex system. This will be illustrated in this section.

Consider a sample consisting of molecules which contain two spins, I and S; the spins are not coupled. As was seen in section 2.4, the two spins have between them four energy levels, which can be labelled according to the spin states of the two spins.

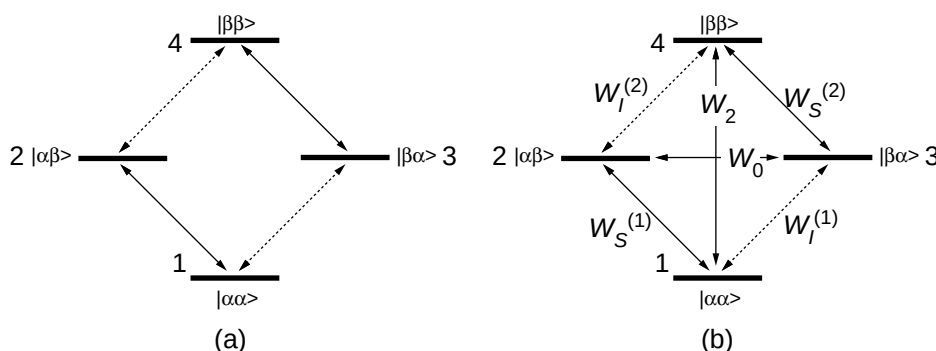


Diagram (a) shows the energy levels of a two spin system; the levels are labelled with the spin of I first and the spin of S second. The dashed arrows indicate allowed transitions of the I spin, and the solid arrows indicate allowed transitions of the S spin. Diagram (b) shows the relaxation induced transitions which are possible amongst the same set of levels.

It turns out that in such a system it is possible to have relaxation induced transitions between *all* possible pairs of energy levels, even those transitions which are forbidden in normal spectroscopy; why this is so will be seen in detail below. The rate constants for the two allowed I spin transitions will be denoted $W_I^{(1)}$ and $W_I^{(2)}$, and likewise for the spin S transitions. The rate constant for the transition between the $\alpha\alpha$ and $\beta\beta$ states is denoted W_2 , the "2" indicating that it is a double quantum transition. Finally, the rate constant for the transition between the $\alpha\beta$ and $\beta\alpha$ states is denoted W_0 , the "0" indicating that it is a zero quantum transition.

Just in the same way as was done in Section 8.2, rate equations can be written for the flow of population from any of the levels. For example, for level 1

$$\frac{dn_1}{dt} = -W_S^{(1)}n_1 - W_I^{(1)}n_1 - W_2n_1 + W_S^{(1)}n_2 + W_I^{(1)}n_3 + W_2n_4$$

The negative terms are rates which lead to a loss of population of level 1 and the positive terms are ones that lead to a gain in its population. As was discussed in section 8.2 the populations ought to be written as deviations from their equilibrium values, $(n_i - n_i^0)$. However, to do this results in unnecessary complexity; rather, the calculation will be carried forward as written and then at the last stage the populations will be replaced by their deviations from equilibrium.

The corresponding equations for the other populations are

$$\begin{aligned}\frac{dn_2}{dt} &= -W_S^{(1)}n_2 - W_I^{(2)}n_2 - W_0n_2 + W_S^{(1)}n_1 + W_I^{(2)}n_4 + W_0n_3 \\ \frac{dn_3}{dt} &= -W_I^{(1)}n_3 - W_S^{(2)}n_3 - W_0n_3 + W_I^{(1)}n_1 + W_S^{(2)}n_4 + W_0n_2 \\ \frac{dn_4}{dt} &= -W_S^{(2)}n_4 - W_I^{(2)}n_4 - W_2n_4 + W_S^{(2)}n_3 + W_I^{(2)}n_2 + W_2n_1\end{aligned}$$

All of this can be expressed in a more compact way if we introduce the I and S spin z-magnetizations. The I spin magnetization is equal to the population difference across the two I spin transitions, 1–3 and 2–4

$$I_z = n_1 - n_3 + n_2 - n_4 \quad [9]$$

As discussed above, the magnetization has been represented as the corresponding operator, I_z . Likewise for the S-spin magnetization

$$S_z = n_1 - n_2 + n_3 - n_4 \quad [10]$$

A third combination of populations will be needed, which is represented by the operator $2I_zS_z$

$$2I_zS_z = n_1 - n_3 - n_2 + n_4 \quad [11]$$

Comparing this with Eq. [9] reveals that $2I_zS_z$ represents the *difference* in population differences across the two I-spin transitions; likewise, comparison with Eq. [10] shows that the same operator also represents the difference in population differences across the two S-spin transitions.

Taking the derivative of Eq. [9] and then substituting for the derivatives of the populations gives

$$\begin{aligned}\frac{dI_z}{dt} &= \frac{dn_1}{dt} - \frac{dn_3}{dt} + \frac{dn_2}{dt} - \frac{dn_4}{dt} \\ &= -W_S^{(1)}n_1 - W_I^{(1)}n_1 - W_2n_1 + W_S^{(1)}n_2 + W_I^{(1)}n_3 + W_2n_4 \\ &\quad + W_I^{(1)}n_3 + W_S^{(2)}n_3 + W_0n_3 - W_I^{(1)}n_1 - W_S^{(2)}n_4 - W_0n_2 \\ &\quad - W_S^{(1)}n_2 - W_I^{(2)}n_2 - W_0n_2 + W_S^{(1)}n_1 + W_I^{(2)}n_4 + W_0n_3 \\ &\quad + W_S^{(2)}n_4 + W_I^{(2)}n_4 + W_2n_4 - W_S^{(2)}n_3 - W_I^{(2)}n_2 - W_2n_1\end{aligned} \quad [12]$$

This unpromising looking equation can be expressed in terms of I_z , S_z etc. by first introducing one more operator E , which is essentially the identity or unit operator

$$E = n_1 + n_2 + n_3 + n_4 \quad [13]$$

and then realizing that the populations, n_i , can be written in terms of E , I_z , S_z , and $2I_zS_z$:

$$\begin{aligned}n_1 &= \frac{1}{4}(E + I_z + S_z + 2I_zS_z) \\ n_2 &= \frac{1}{4}(E + I_z - S_z - 2I_zS_z) \\ n_3 &= \frac{1}{4}(E - I_z + S_z - 2I_zS_z) \\ n_4 &= \frac{1}{4}(E - I_z - S_z + 2I_zS_z)\end{aligned}$$

where these relationships can easily be verified by substituting back in the definitions of the operators in terms of populations, Eqs. [9] – [13].

After some tedious algebra, the following differential equation is found

for I_z

$$\begin{aligned} \frac{dI_z}{dt} = & -\left(W_I^{(1)} + W_I^{(2)} + W_2 + W_0\right)I_z \\ & -\left(W_2 - W_0\right)S_z - \left(W_I^{(1)} - W_I^{(2)}\right)2I_zS_z \end{aligned} \quad [14]$$

Similar algebra gives the following differential equations for the other operators

$$\begin{aligned} \frac{dS_z}{dt} = & -\left(W_2 - W_0\right)I_z - \left(W_S^{(1)} + W_S^{(2)} + W_2 + W_0\right)S_z - \left(W_S^{(1)} - W_S^{(2)}\right)2I_zS_z \\ \frac{d2I_zS_z}{dt} = & -\left(W_I^{(1)} - W_I^{(2)}\right)I_z - \left(W_S^{(1)} - W_S^{(2)}\right)S_z \\ & -\left(W_I^{(1)} + W_I^{(2)} + W_S^{(1)} + W_S^{(2)}\right)2I_zS_z \end{aligned}$$

As expected, the total population, represented by E , does not change with time. These three differential equations are known as the *Solomon equations*.

It must be remembered that the populations used to derive these equations are really the deviation of the populations from their equilibrium values. As a result, the I and S spin magnetizations should properly be their deviations from their equilibrium values, I_z^0 and S_z^0 ; the equilibrium value of $2I_zS_z$ is easily shown, from its definition, to be zero. For example, Eq. [14] becomes

$$\begin{aligned} \frac{d(I_z - I_z^0)}{dt} = & -\left(W_I^{(1)} + W_I^{(2)} + W_2 + W_0\right)(I_z - I_z^0) \\ & -\left(W_2 - W_0\right)(S_z - S_z^0) - \left(W_I^{(1)} - W_I^{(2)}\right)2I_zS_z \end{aligned}$$

8.3.1 Interpreting the Solomon equations

What the Solomon equations predict is, for example, that the rate of change of I_z depends not only on $I_z - I_z^0$, but also on $S_z - S_z^0$ and $2I_zS_z$. In other words the way in which the magnetization on the I spin varies with time depends on what is happening to the S spin – the two magnetizations are connected. This phenomena, by which the magnetizations of the two different spins are connected, is called *cross relaxation*.

The rate at which S magnetization is transferred to I magnetization is given by the term

$$(W_2 - W_0)(S_z - S_z^0)$$

in Eq. [14]; $(W_2 - W_0)$ is called the cross-relaxation rate constant, and is sometimes given the symbol σ_{IS} . It is clear that in the absence of the relaxation pathways between the $\alpha\alpha$ and $\beta\beta$ states (W_2), or between the $\alpha\beta$ and $\beta\alpha$ states (W_0), there will be no cross relaxation. This term is described as giving rise to transfer from S to I as it says that the rate of change of the I spin magnetization is proportional to the deviation of the S spin magnetization from its equilibrium value. Thus, if the S spin is not at equilibrium the I spin magnetization is perturbed.

In Eq. [14] the term

$$(W_I^{(1)} + W_I^{(2)} + W_2 + W_0)(I_z - I_z^0)$$

describes the relaxation of I spin magnetization on its own; this is sometimes called the self relaxation. Even if W_2 and W_0 are absent, self relaxation still occurs. The self relaxation rate constant, given in the previous equation as a sum of W values, is sometimes given the symbol R_I or ρ_I .

Finally, the term

$$(W_I^{(1)} - W_I^{(2)})2I_zS_z$$

in Eq. [14] describes the transfer of I_zS_z into I spin magnetization. Recall that $W_I^{(1)}$ and $W_I^{(2)}$ are the relaxation induced rate constants for the two allowed transitions of the I spin (1–3 and 2–4). Only if these two rate constants are different will there be transfer from $2I_zS_z$ into I spin magnetization. This situation arises when there is *cross-correlation* between different relaxation mechanisms; a further discussion of this is beyond the scope of these lectures. The rate constants for this transfer will be written

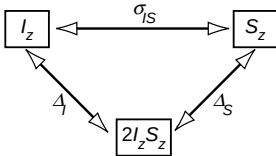
$$\Delta_I = (W_I^{(1)} - W_I^{(2)}) \quad \Delta_S = (W_S^{(1)} - W_S^{(2)})$$

According to the final Solomon equation, the operator $2I_zS_z$ shows self relaxation with a rate constant

$$R_{IS} = (W_I^{(1)} + W_I^{(2)} + W_S^{(1)} + W_S^{(2)})$$

Note that the W_2 and W_0 pathways do not contribute to this. This rate combined constant will be denoted R_{IS} .

Using these combined rate constants, the Solomon equations can be written



$$\begin{aligned} \frac{d(I_z - I_z^0)}{dt} &= -R_I(I_z - I_z^0) - \sigma_{IS}(S_z - S_z^0) - \Delta_I 2I_zS_z \\ \frac{d(S_z - S_z^0)}{dt} &= -\sigma_{IS}(I_z - I_z^0) - R_S(S_z - S_z^0) - \Delta_S 2I_zS_z \\ \frac{d2I_zS_z}{dt} &= -\Delta_I(I_z - I_z^0) - \Delta_S(S_z - S_z^0) - R_{IS} 2I_zS_z \end{aligned} \quad [15]$$

The pathways between the different magnetization are visualized in the diagram opposite. Note that as $dI_z^0/dt = 0$ (the equilibrium magnetization is a constant), the derivatives on the left-hand side of these equations can equally well be written dI_z/dt and dS_z/dt .

It is important to realize that in such a system I_z and S_z do not relax with a simple exponentials. They only do this if the differential equation is of the form

$$\frac{dI_z}{dt} = -R_I(I_z - I_z^0)$$

which is plainly not the case here. For such a two-spin system, therefore, it is not proper to talk of a " T_1 " relaxation time constant.

8.4 Nuclear Overhauser effect

The Solomon equations are an excellent way of understanding and analysing experiments used to measure the nuclear Overhauser effect. Before embarking on this discussion it is important to realize that although the states represented by operators such as I_z and S_z cannot be observed directly, they can be made observable by the application of a radiofrequency pulse, ideally a 90° pulse

$$aI_z \xrightarrow{(\pi/2)I_x} -aI_y$$

The subsequent recording of the free induction signal due to the evolution of the operator I_y will give, after Fourier transformation, a spectrum with a peak of size $-a$ at frequency Ω_I . In effect, by computing the value of the coefficient a , the appearance of the subsequently observed spectrum is predicted.

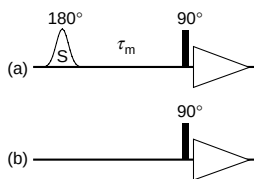
The basis of the nuclear Overhauser effect can readily be seen from the Solomon equation (for simplicity, it is assumed in this section that $\Delta_I = \Delta_S = 0$)

$$\frac{d(I_z - I_z^0)}{dt} = -R_I(I_z - I_z^0) - \sigma_{IS}(S_z - S_z^0)$$

What this says is that if the S spin magnetization deviates from equilibrium there will be a change in the I spin magnetization at a rate proportional to (a) the cross-relaxation rate, σ_{IS} and (b) the extent of the deviation of the S spin from equilibrium. This change in the I spin magnetization will manifest itself as a change in the intensity in the corresponding spectrum, and it is this change in intensity of the I spin when the S spin is perturbed which is termed the nuclear Overhauser effect.

Plainly, there will be no such effect unless σ_{IS} is non-zero, which requires the presence of the W_2 and W_0 relaxation pathways. It will be seen later on that such pathways are only present when there is dipolar relaxation between the two spins and that the resulting cross-relaxation rate constants have a strong dependence on the distance between the two spins. The observation of a nuclear Overhauser effect is therefore diagnostic of dipolar relaxation and hence the proximity of pairs of spins. The effect is of enormous value, therefore, in structure determination by NMR.

8.4.1 Transient experiments



Pulse sequence for recording transient NOE enhancements. Sequence (a) involves selective inversion of the S spin – shown here using a shaped pulse. Sequence (b) is used to record the reference spectrum in which the intensities are unperturbed.

A simple experiment which reveals the NOE is to invert just the S spin by applying a selective 180° pulse to its resonance. The S spin is then not at equilibrium so magnetization is transferred to the I spin by cross-relaxation. After a suitable period, called the mixing time, τ_m , a non-selective 90° pulse is applied and the spectrum recorded.

After the selective pulse the situation is

$$I_z(0) = I_z^0 \quad S_z(0) = -S_z^0 \quad [16]$$

where I_z has been written as $I_z(t)$ to emphasize that it depends on time and likewise for S. To work out what will happen during the mixing time the differential equations

$$\begin{aligned} \frac{dI_z(t)}{dt} &= -R_I(I_z(t) - I_z^0) - \sigma_{IS}(S_z(t) - S_z^0) \\ \frac{dS_z(t)}{dt} &= -\sigma_{IS}(I_z(t) - I_z^0) - R_S(S_z(t) - S_z^0) \end{aligned}$$

need to be solved (integrated) with this initial condition. One simple way to do this is to use the *initial rate approximation*. This involves assuming that the mixing time is sufficiently short that, on the *right-hand side* of the equations, it can be assumed that the initial conditions set out in Eq. [16] apply, so, for the first equation

$$\begin{aligned} \frac{dI_z(t)}{dt} \Big|_{\text{init}} &= -R_I(I_z^0 - I_z^0) - \sigma_{IS}(-S_z^0 - S_z^0) \\ &= 2\sigma_{IS}S_z^0 \end{aligned}$$

This is now easy to integrate as the right-hand side has no dependence on $I_z(t)$

$$\begin{aligned} \int_0^{\tau_m} dI_z(t) &= \int_0^{\tau_m} 2\sigma_{IS}S_z^0 dt \\ I_z(\tau_m) - I_z(0) &= 2\sigma_{IS}\tau_m S_z^0 \\ I_z(\tau_m) &= 2\sigma_{IS}\tau_m S_z^0 + I_z^0 \end{aligned}$$

This says that for zero mixing time the I magnetization is equal to its equilibrium value, but that as the mixing time increases the I magnetization has an additional contribution which is proportional to the mixing time and the cross-relaxation rate, σ_{IS} . This latter term results in a change in the intensity of the I spin signal, and this change is called an *NOE enhancement*.

The normal procedure for visualizing these enhancements is to record a reference spectrum in which the intensities are unperturbed. In terms of z-magnetizations this means that $I_{z,\text{ref}} = I_z^0$. The difference spectrum, defined as (perturbed spectrum – unperturbed spectrum) corresponds to the difference

$$\begin{aligned} I_z(\tau_m) - I_{z,\text{ref}} &= 2\sigma_{IS}\tau_m S_z^0 + I_z^0 - I_z^0 \\ &= 2\sigma_{IS}\tau_m S_z^0 \end{aligned}$$

The NOE enhancement factor, η , is defined as

$$\eta = \frac{\text{intensity in enhanced spectrum} - \text{intensity in reference spectrum}}{\text{intensity in reference spectrum}}$$

so in this case η is

$$\eta(\tau_m) = \frac{I_z(\tau_m) - I_{z,\text{ref}}}{I_{z,\text{ref}}} = \frac{2\sigma_{IS}\tau_m S_z^0}{I_z^0}$$

and if I and S are of the same nuclear species (*e.g.* both proton), their equilibrium magnetizations are equal so that

$$\eta(\tau_m) = 2\sigma_{IS}\tau_m$$

Hence a plot of η against mixing time will give a straight line of slope σ_{IS} ; this is a method used for measuring the cross-relaxation rate constant. A single experiment for one value of the mixing time will reveal the presence of NOE enhancements.

This initial rate approximation is valid provided that

$$\sigma_{IS}\tau_m \ll 1 \text{ and } R_S\tau_m \ll 1$$

the first condition means that there is little transfer of magnetization from S to I, and the second means that the S spin remains very close to complete inversion.

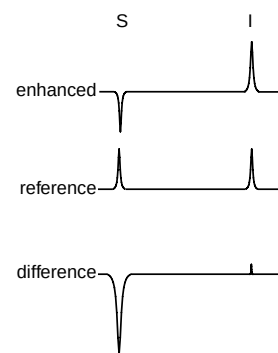
8.4.1.1 Advanced topic: longer mixing times

At longer mixing times the differential equations are a little more difficult to solve, but they can be integrated using standard methods (symbolic mathematical programmes such as *Mathematica* are particularly useful for this). Using the initial conditions given in Eq. [16] and, assuming for simplicity that $I_z^0 = S_z^0$ the following solutions are found

$$\begin{aligned} \frac{I_z(\tau_m)}{I_z^0} &= \frac{2\sigma_{IS}}{R} [\exp(-\lambda_2\tau_m) - \exp(-\lambda_1\tau_m)] + 1 \\ \frac{S_z(\tau_m)}{I_z^0} &= \left[\frac{R_I - R_S}{R} \right] [\exp(-\lambda_1\tau_m) - \exp(-\lambda_2\tau_m)] \\ &\quad + 1 - \exp(-\lambda_1\tau_m) + \exp(-\lambda_2\tau_m) \end{aligned}$$

where

$$\begin{aligned} R &= \sqrt{R_I^2 - 2R_IR_S + R_S^2 + 4\sigma_{IS}^2} \\ \lambda_1 &= \frac{1}{2}[R_I + R_S + R] \quad \lambda_2 = \frac{1}{2}[R_I + R_S - R] \end{aligned}$$

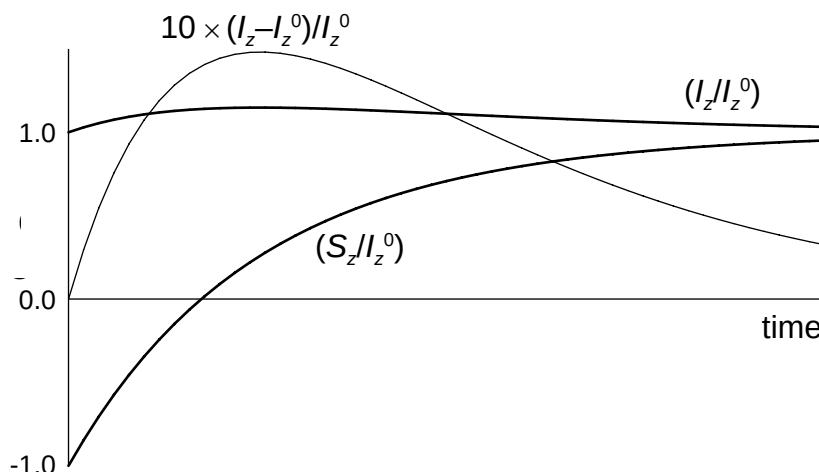


Visualization of how an NOE difference spectrum is recorded. The enhancement is assumed to be positive.

These definitions ensure that $\lambda_1 > \lambda_2$. If R_I and R_S are not too dissimilar, R is of the order of σ_{IS} , and so the two rate constants λ_1 and λ_2 differ by a quantity of the order of σ_{IS} .

As expected for these two coupled differential equations, integration gives a time dependence which is the sum of two exponentials with different time constants.

The figure below shows the typical behaviour predicted by these equations (the parameters are $R_I = R_S = 5\sigma_{IS}$)



The S spin magnetization returns to its equilibrium value with what appears to be an exponential curve; in fact it is the sum of two exponentials but their time constants are not sufficiently different for this to be discerned. The I spin magnetization grows towards a maximum and then drops off back towards the equilibrium value. The NOE enhancement is more easily visualized by plotting the difference magnetization, $(I_z - I_z^0)/I_z^0$, on an expanded scale; the plot now shows the positive NOE enhancement reaching a maximum of about 15%.

Differentiation of the expression for I_z as a function of τ_m shows that the maximum enhancement is reached at time

$$\tau_{m,\max} = \frac{1}{\lambda_1 - \lambda_2} \ln \frac{\lambda_1}{\lambda_2}$$

and that the maximum enhancement is

$$\frac{I_z(\tau_{m,\max}) - I_z^0}{I_z^0} = \frac{2\sigma_{IS}}{R} \left[\left(\frac{\lambda_1}{\lambda_2} \right)^{\frac{-\lambda_1}{R}} - \left(\frac{\lambda_1}{\lambda_2} \right)^{\frac{-\lambda_2}{R}} \right]$$

8.4.2 The DPFGE NOE experiment

From the point of view of the relaxation behaviour the DPFGE experiment is essentially identical to the transient NOE experiment. The only difference is that the I spin starts out *saturated* rather than at equilibrium. This does not influence the build up of the NOE enhancement on I. It does, however, have the advantage of reducing the size of the I spin signal which has to be removed in the difference experiment. Further discussion of this experiment is deferred to Chapter 9.

8.4.3 Steady state experiments

The steady-state NOE experiment involves irradiating the S spin with a radiofrequency field which is sufficiently weak that the I spin is not affected. The irradiation is applied for long enough that the S spin is saturated, meaning $S_z = 0$, and that the steady state has been reached, which means that none of the magnetizations are changing, *i.e.* $(dI_z/dt) = 0$.

Under these conditions the first of Eqs. [15] can be written

$$\left. \frac{d(I_z - I_z^0)}{dt} \right|_{ss} = -R_I(I_{z,ss} - I_z^0) - \sigma_{IS}(0 - S_z^0) = 0$$

therefore

$$I_{z,ss} = \frac{\sigma_{IS}}{R_I} S_z^0 + I_z^0$$

As in the transient experiment, the NOE enhancement is revealed by subtracting a reference spectrum which has equilibrium intensities. The NOE enhancement, as defined above, will be

$$\eta_{ss} = \frac{I_{z,ss} - I_{z,ref}}{I_{z,ref}} = \frac{\sigma_{IS}}{R_I} \frac{S_z^0}{I_z^0}$$

In contrast to the transient experiment, the steady state enhancement only depends on the relaxation of the receiving spin (here I); the relaxation rate of the S spin does not enter into the relationship simply because this spin is held saturated during the experiment.

It is important to realise that the value of the steady-state NOE enhancement depends on the ratio of cross-relaxation rate constant to the self relaxation rate constant for the spin which is receiving the enhancement. If this spin is relaxing quickly, for example as a result of interaction with many other spins, the size of the NOE enhancement will be reduced. So, although the size of the enhancement does depend on the cross-relaxation rate constant the size of the enhancement cannot be interpreted in terms of this rate constant alone. This point is illustrated by the example in the margin.

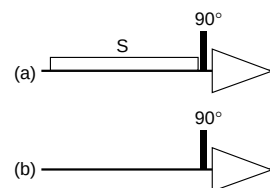
8.4.4 Advanced topic: NOESY

The dynamics of the NOE in NOESY are very similar to those for the transient NOE experiment. The key difference is that instead of the magnetization of the S spin being inverted at the start of the mixing time, the magnetization has an amplitude label which depends on the evolution during t_1 .

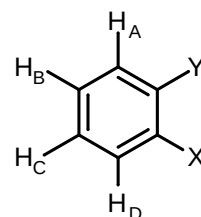
Starting with equilibrium magnetization on the I and S spins, the z-magnetizations present at the start of the mixing time are (other magnetization will be rejected by appropriate phase cycling)

$$S_z(0) = -\cos \Omega_S t_1 S_z^0 \quad I_z(0) = -\cos \Omega_I t_1 I_z^0$$

The equation of motion for S_z is



Pulse sequence for recording steady state NOE enhancements. Sequence (a) involves selective irradiation of the S spin leading to saturation. Sequence (b) is used to record the reference spectrum in which the intensities are unperturbed.



Irradiation of proton B gives a much larger enhancement on proton A than on C despite the fact that the distances to the two spins are equal. The smaller enhancement on C is due to the fact that it is relaxing more quickly than A, due to the interaction with proton D.



Pulse sequence for NOESY.

$$\frac{dS_z(t)}{dt} = -\sigma_{IS}(I_z(t) - I_z^0) - R_S(S_z(t) - S_z^0)$$

As before, the initial rate approximation will be used:

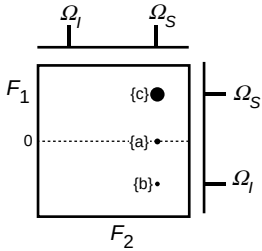
$$\begin{aligned} \left. \frac{dS_z(\tau_m)}{dt} \right|_{\text{init}} &= -\sigma_{IS}(-\cos \Omega_I t_1 I_z^0 - I_z^0) - R_S(-\cos \Omega_S t_1 S_z^0 - S_z^0) \\ &= \sigma_{IS}(\cos \Omega_I t_1 + 1)I_z^0 + R_S(\cos \Omega_S t_1 + 1)S_z^0 \end{aligned}$$

Integrating gives

$$\begin{aligned} \int_0^{\tau_m} dS_z(t) &= \int_0^{\tau_m} [\sigma_{IS}(\cos \Omega_I t_1 + 1)I_z^0 + R_S(\cos \Omega_S t_1 + 1)S_z^0] dt \\ S_z(\tau_m) - S_z(0) &= \sigma_{IS}\tau_m(\cos \Omega_I t_1 + 1)I_z^0 + R_S\tau_m(\cos \Omega_S t_1 + 1)S_z^0 \\ S_z(\tau_m) &= \sigma_{IS}\tau_m(\cos \Omega_I t_1 + 1)I_z^0 + R_S\tau_m(\cos \Omega_S t_1 + 1)S_z^0 - \cos \Omega_S t_1 S_z^0 \\ &= \sigma_{IS}\tau_m I_z^0 + R_S\tau_m S_z^0 \quad \{a\} \\ &\quad + \cos \Omega_I t_1 [\sigma_{IS}\tau_m] I_z^0 \quad \{b\} \\ &\quad + \cos \Omega_S t_1 [R_S\tau_m - 1] S_z^0 \quad \{c\} \end{aligned}$$

After the end of the mixing time, this z-magnetization on spin S is rendered observable by the final 90° pulse; the magnetization is on spin S, and so will precess at Ω_S during t_2 .

The three terms {a}, {b} and {c} all represent different peaks in the NOESY spectrum.



Term {a} has no evolution as a function of t_1 and so will appear at $F_1 = 0$; in t_2 it evolves at Ω_S . This is therefore an *axial peak* at $\{F_1, F_2\} = \{0, \Omega_S\}$. This peak arises from z-magnetization which has recovered during the mixing time. In this initial rate limit, it is seen that the axial peak is zero for zero mixing time and then grows linearly depending on R_S and σ_{IS} .

Term {b} evolves at Ω_I during t_1 and Ω_S during t_2 ; it is therefore a cross peak at $\{\Omega_I, \Omega_S\}$. The intensity of the cross peak grows linearly with the mixing time and also depends on σ_{IS} ; this is analogous to the transient NOE experiment.

Term {c} evolves at Ω_S during t_1 and Ω_S during t_2 ; it is therefore a diagonal peak at $\{\Omega_S, \Omega_S\}$ and as $R_S\tau_m \ll 1$ in the initial rate, this peak is negative. The intensity of the peak grows back towards zero linearly with the mixing time and at a rate depending on R_S . This peak arises from S spin magnetization which remains on S during the mixing time, decaying during that time at a rate determined by R_S .

If the calculation is repeated using the differential equation for I_z a complimentary set of peaks at $\{0, \Omega_I\}$, $\{\Omega_S, \Omega_I\}$ and $\{\Omega_I, \Omega_I\}$ are found.

It will be seen later that whereas R_I and R_S are positive, σ_{IS} can be either positive or negative. If σ_{IS} is positive, the diagonal and cross peaks will be of opposite sign, whereas if σ_{IS} is negative all the peaks will have the same sign.

8.4.5 Sign of the NOE enhancement

We see that the time dependence and size of the NOE enhancement depends on the relative sizes of the cross-relaxation rate constant σ_{IS} and the self relaxation rate constants R_I and R_S . It turns out that these self-rates are always positive, but the cross-relaxation rate constant can be positive or negative. The reason for this is that $\sigma_{IS} = (W_2 - W_0)$ and it is quite possible for W_0 to be greater or less than W_2 .

A positive cross-relaxation rate constant means that if spin S deviates from equilibrium cross-relaxation will *increase* the magnetization on spin I. This leads to an *increase* in the signal from I, and hence a *positive* NOE enhancement. This situation is typical for small molecules in non-viscous solvents.

A negative cross-relaxation rate constant means that if spin S deviates from equilibrium cross-relaxation will *decrease* the magnetization on spin I. This leads to a *negative* NOE enhancement, a situation typical for large molecules in viscous solvents. Under some conditions W_0 and W_2 can become equal and then the NOE enhancement goes to zero.

8.5 Origins of relaxation

We now turn to the question as to what causes relaxation. Recall from section 8.1 that relaxation involves transitions between energy levels, so what we seek is the origin of these transitions. We already know from Chapter 3 that transitions are caused by transverse magnetic fields (*i.e.* in the *xy*-plane) which are oscillating close to the Larmor frequency. An RF pulse gives rise to just such a field.

However, there is an important distinction between the kind of transitions caused by RF pulses and those which lead to relaxation. When an RF pulse is applied all of the spins experience the same oscillating field. The kind of transitions which lead to relaxation are different in that the transverse fields are *local*, meaning that they only affect a few spins and not the whole sample. In addition, these fields vary randomly in direction and amplitude. In fact, it is precisely their random nature which drives the sample to equilibrium.

The fields which are responsible for relaxation are generated within the sample, often due to interactions of spins with one another or with their environment in some way. They are made time varying by the random motions (rotations, in particular) which result from the thermal agitation of the molecules and the collisions between them. Thus we will see that NMR relaxation rate constants are particularly sensitive to molecular motion.

If the spins need to lose energy to return to equilibrium they give this up to the motion of the molecules. Of course, the amounts of energy given up by the spins are tiny compared to the kinetic energies that molecules have, so they are hardly affected. Likewise, if the spins need to increase their energy to go to equilibrium, for example if the population of the β state has to be increased, this energy comes from the motion of the molecules.

Relaxation is essentially the process by which energy is allowed to flow

between the spins and molecular motion. This is the origin of the original name for longitudinal relaxation: *spin-lattice relaxation*. The lattice does not refer to a solid, but to the motion of the molecules with which energy can be exchanged.

8.5.1 Factors influencing the relaxation rate constant

The detailed theory of the calculation of relaxation rate constants is beyond the scope of this course. However, we are in a position to discuss the kinds of factors which influence these rate constants.

Let us consider the rate constant W_{ij} for transitions between levels i and j ; this turns out to depend on three factors:

$$W_{ij} = A_{ij} \times Y \times J(\omega_{ij})$$

We will consider each in turn.

The spin factor, A_{ij}

This factor depends on the quantum mechanical details of the interaction. For example, not all oscillating fields can cause transitions between all levels. In a two spin system the transition between the $\alpha\alpha$ and $\beta\beta$ cannot be brought about by a simple oscillating field in the transverse plane; in fact it needs a more complex interaction that is only present in the dipolar mechanism (section 8.6.2). We can think of A_{ij} as representing a kind of selection rule for the process – like a selection rule it may be zero for some transitions.

The size factor, Y

This is just a measure of how large the interaction causing the relaxation is. Its size depends on the detailed origin of the random fields and often it is related to molecular geometry.

The spectral density, $J(\omega_{ij})$

This is a measure of the amount of molecular motion which is at the correct frequency, ω_{ij} , to cause the transitions. Recall that molecular motion is the effect which makes the random fields vary with time. However, as we saw with RF pulses, the field will only have an effect on the spins if it is oscillating at the correct frequency. The spectral density is a measure of how much of the motion is present at the correct frequency.

8.5.2 Spectral densities and correlation functions

The value of the spectral density, $J(\omega)$, has a large effect on relaxation rate constants, so it is well worthwhile spending some time in understanding the form that this function takes.

Correlation functions

To make the discussion concrete, suppose that a spin in a sample experiences a magnetic field due to a dissolved paramagnetic species. The size of the magnetic field will depend on the relative orientation of the spin and the paramagnetic species, and as both are subject to random thermal motion, this orientation will vary randomly with time (it is said to be a random function of time), and so the magnetic field will be a random function of time. Let the field experienced by this first spin be $F_1(t)$.

Now consider a second spin in the sample. This also experiences a random magnetic field, $F_2(t)$, due to the interaction with the paramagnetic species. At any instant, this random field will not be the same as that experienced by the first spin.

For a macroscopic sample, each spin experiences a different random field, $F_i(t)$. There is no way that a detailed knowledge of each of these random fields can be obtained, but in some cases it is possible to characterise the *overall* behaviour of the system quite simply.

The average field experienced by the spins is found by taking the ensemble average – that is adding up the fields for all members of the ensemble (*i.e.* all spins in the system)

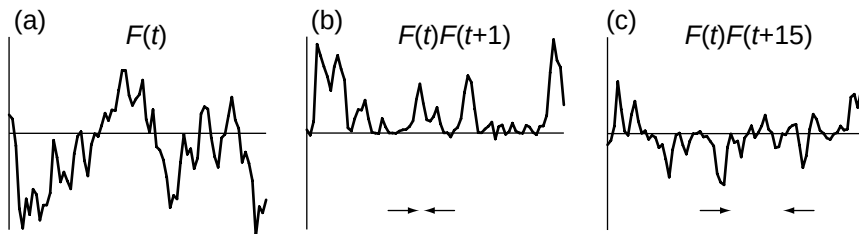
$$\overline{F(t)} = F_1(t) + F_2(t) + F_3(t) + \dots$$

For random thermal motion, this ensemble average turns out to be independent of the time; this is a property of *stationary* random functions. Typically, the $F_i(t)$ are signed quantities, randomly distributed about zero, so this ensemble average will be zero.

An important property of random functions is the *correlation function*, $G(t, \tau)$, defined as

$$\begin{aligned} G(t, \tau) &= F_1(t)F_1^*(t + \tau) + F_2(t)F_2^*(t + \tau) + F_3(t)F_3^*(t + \tau) + \dots \\ &= \overline{F(t)F^*(t + \tau)} \end{aligned}$$

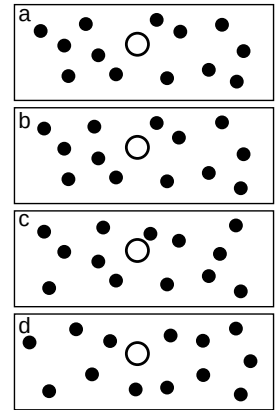
$F_1(t)$ is the field experienced by spin 1 at time t , and $F_1(t + \tau)$ is the field experienced at a time τ later; the star indicates the complex conjugate, which allows for the possibility that $F(t)$ may be complex. If the time τ is short the spins will not have moved very much and so $F_1(t + \tau)$ will be very little different from $F_1(t)$. As a result, the product $F_1(t)F_1^*(t + \tau)$ will be positive. This is illustrated in the figure below, plot (b).



(a) is a plot of the random function $F(t)$ against time; there are about 100 separate time points. (b) is a plot of the value of F multiplied by its value one data point later – *i.e.* one data point to the right; all possible pairs are plotted. (c) is the same as (b) but for a time interval of 15 data points. The two arrows indicate the spacing over which the correlation is calculated.

The same is true for all of the other members of the ensemble, so when the $F_i(t)F_i^*(t + \tau)$ are added together for a particular time, t , – that is, the

Paramagnetic species have unpaired electrons. These generate magnetic fields which can interact with nearby nuclei. On account of the large gyromagnetic ratio of the electron (when compared to the nucleus) such paramagnetic species are often a significant source of relaxation.

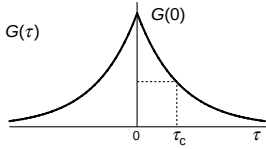


Visualization of the different timescales for random motion. (a) is the starting position: the black dots are spins and the open circle represents a paramagnetic species. (b) is a snapshot a very short time after (a); hardly any of the spins have moved. (c) is a snapshot at a longer time; more spins have moved, but part of the original pattern is still discernible. (d) is after a long time, all the spins have moved and the original pattern is lost.

ensemble average is taken – the result will be for them to reinforce one another and hence give a finite value for $G(t, \tau)$.

As τ gets longer, the spin will have had more chance of moving and so $F_1(t+\tau)$ will differ more and more from $F_1(t)$; the product $F_1(t)F_1^*(t+\tau)$ need not necessarily be positive. This is illustrated in plot (c) above. The ensemble average of all these $F_i(t)F_i^*(t+\tau)$ is thus less than it was when τ was shorter. In the limit, once τ becomes sufficiently long, the $F_i(t)F_i^*(t+\tau)$ are randomly distributed and their ensemble average, $G(t, \tau)$, goes to zero. $G(t, \tau)$ thus has its maximum value at $\tau = 0$ and then decays to zero at long times. For stationary random functions, the correlation function is independent of the time t ; it will therefore be written $G(\tau)$.

The correlation function, $G(\tau)$, is thus a function which characterises the memory that the system has of a particular arrangement of spins in the sample. For times τ which are much less than the time it takes for the system to rearrange itself $G(\tau)$ will be close to its maximum value. As time proceeds, the initial arrangement becomes more and more disturbed, and $G(\tau)$ falls. For sufficiently long times, $G(\tau)$ tends to zero.



The simplest form for $G(\tau)$ is

$$G(\tau) = G(0) \exp(-|\tau|/\tau_c) \quad [18]$$

the variable τ appears as the modulus, resulting in the same value of $G(\tau)$ for positive and negative values of τ . This means that the correlation is the same with time τ before and time τ after the present time.

τ_c is called the *correlation time*. For times much less than the correlation time the spins have not moved much and the correlation function is close to its original value; when the time is of the order of τ_c significant rearrangements have taken place and the correlation function has fallen to about half its initial value. For times much longer than τ_c the spins have moved to completely new positions and the correlation function has fallen close to zero.

Spectral densities

The correlation function is a function of time, just like a free induction decay. So, it can be Fourier transformed to give a function of frequency. The resulting frequency domain function is called the *spectral density*; as the name implies, the spectral density gives a measure of the amount of motion present at different frequencies. The spectral density is usually denoted $J(\omega)$

$$G(\tau) \xrightarrow{\text{Fourier Transform}} J(\omega)$$

If the spins were executing a well ordered motion, such as oscillating back and forth about a mean position, the spectral density would show a peak at that frequency. However, the spins are subject to random motions with a range of different periods, so the spectral density shows a range of frequencies rather than having peaks at discrete frequencies.

Generally, for random motion characterised by a correlation time τ_c , frequencies from zero up to about $1/\tau_c$ are present. The amount at frequencies higher than $1/\tau_c$ tails off quite rapidly as the frequency increases.

For a simple exponential correlation function, given in Eq. [18], the corresponding spectral density is a Lorentzian

$$\exp(-|\tau|/\tau_c) \xrightarrow{\text{Fourier Transform}} \frac{2\tau_c}{1 + \omega^2 \tau_c^2}$$

This function is plotted in the margin; note how it drops off significantly once the product $\omega\tau_c$ begins to exceed ~ 1 .

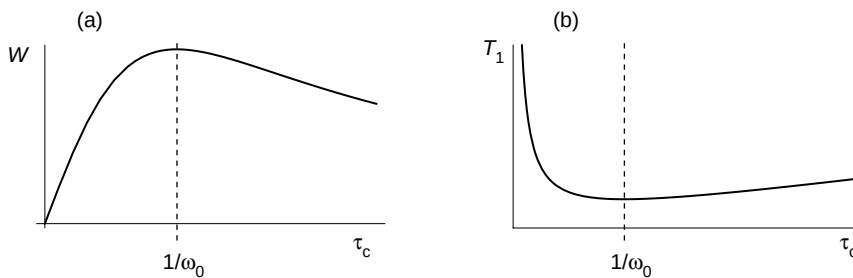
The plot opposite compares the spectral densities for three different correlation times; curve *a* is the longest, *b* an intermediate value and *c* the shortest. Note that as the correlation time decreases the spectral density moves out to higher frequencies. However, the area under the plot remains the same, so the contribution at lower frequencies is decreased. In particular, at the frequency indicated by the dashed line the contribution at correlation time *b* is greater than that for either correlation times *a* or *c*.

For this spectral density function, the maximum contribution at frequency ω is found when τ_c is $1/\omega$; this has important consequences which are described in the next section.

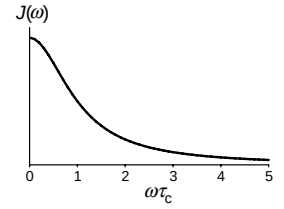
8.5.3 The " T_1 minimum"

In the case of relaxation of a single spin by a random field (such as that generated by a paramagnetic species), the only relevant spectral density is that at the Larmor frequency, ω_0 . This is hardly surprising as to cause relaxation – that is to cause transitions – the field needs to have components oscillating at the Larmor frequency.

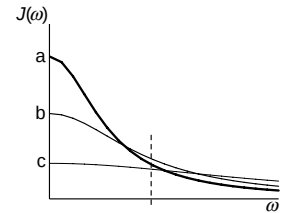
We have just seen that for a given frequency, ω_0 , the spectral density is a maximum when τ_c is $1/\omega_0$, so to have the most rapid relaxation the correlation time should be $1/\omega_0$. This is illustrated in the plots below which show the relaxation rate constant, W , and the corresponding relaxation time ($T_1 = 1/W$) plotted as a function of the correlation time.



Plot (a) shows how the relaxation rate constant, W , varies with the correlation time, τ_c , for a given Larmor frequency; there is a maximum in the rate constant when $\tau_c = 1/\omega_0$. Plot (b) shows the same effect, but here we have plotted the relaxation time constant, T_1 ; this shows a minimum.



Plot of the spectral density as a function of the dimensionless variable $\omega\tau_c$. The curve is a Lorentzian



At very short correlation times ($\tau_c \ll 1/\omega_0$) there is some spectral density at the Larmor frequency, but not that much as the energy of the motion is spread over a very wide frequency range. As the correlation time increases the amount of spectral density at the Larmor frequency increases and so the relaxation rate constant increases, reaching a maximum when $\tau_c = 1/\omega_0$. After this point, the spectral density at the Larmor frequency, and hence the rate constant, falls.

In terms of the relaxation time, T_1 , there is a minimum in T_1 which corresponds to the maximum in W . We see that, like Goldilocks and the Three Bears, efficient relaxation requires a correlation time which is neither too fast nor too slow.

Motion which gives rise to correlation times which are much shorter than $1/\omega_0$ is described as being in the *fast motion* (or *extreme narrowing*) limit. Put mathematically, fast motion means $\omega_0\tau_c \ll 1$. Motion which gives rise to correlation times which are much slower than $1/\omega_0$ is described as being in the *slow motion* (or *spin diffusion*) limit; mathematically this limit is $\omega_0\tau_c \gg 1$. Clearly, which limit we are in depends on the Larmor frequency, which in turn depends on the nucleus and the magnetic field.

For a Larmor frequency of 400 MHz we would expect the fastest relaxation when the correlation time is 0.4 ns. Small molecules have correlation times significantly shorter than this (say tens of ps), so such molecules are clearly in the fast motion limit. Large molecules, such as proteins, can easily have correlation times of the order of a few ns, and these clearly fall in the slow motion limit.

Somewhat strangely, therefore, both very small and very large molecules tend to relax more slowly than medium-sized molecules.

8.6 Relaxation mechanisms

So far, the source of the magnetic fields which give rise to relaxation and the origin of their time dependence have not been considered. Each such source is referred to as a *relaxation mechanism*. There are quite a range of different mechanisms that can act, but of these only a few are really important for spin half nuclei.

8.6.1 Paramagnetic species

We have already mentioned this source of varying fields several times. The large magnetic moment of the electron means that paramagnetic species in solution are particularly effective at promoting relaxation. Such species include dissolved oxygen and certain transition metal compounds.

8.6.2 The dipolar mechanism

Each spin has associated with it a magnetic moment, and this in turn gives rise to a magnetic field which can interact with other spins. Two spins are thus required for this interaction, one to "create" the field and one to "experience" it. However, their roles are reversible, in the sense that the second spin creates a field which is experienced by the first. So, the overall interaction is a property of the pair of nuclei.

The size of the interaction depends on the inverse cube of the distance between the two nuclei and the direction of the vector joining the two nuclei, measured relative to that of the applied magnetic field. As a molecule tumbles in solution the direction of this vector changes and so the magnetic field changes. Changes in the distance between the nuclei also result in a change in the magnetic field. However, molecular vibrations, which do give such changes, are generally at far too high frequencies to give significant spectral density at the Larmor frequency. As a result, it is generally changes in orientation which are responsible for relaxation.

The pair of interacting nuclei can be in the same or different molecules, leading to intra- and inter-molecular relaxation. Generally, however, nuclei in the same molecule can approach much more closely than those in different molecules so that intra-molecular relaxation is dominant.

The relaxation induced by the dipolar coupling is proportional to the *square* of the coupling. Thus it goes as

$$\gamma_1^2 \gamma_2^2 \frac{1}{r_{12}^6}$$

where γ_1 and γ_2 are the gyromagnetic ratios of the two nuclei involved and r_{12} is the distance between them.

As the size of the dipolar interaction depends on the product of the gyromagnetic ratios of the two nuclei involved, and the resulting relaxation rate constants depends on the square of this. Thus, pairs of nuclei with high gyromagnetic ratios are most efficient at promoting relaxation. For example, every thing else being equal, a proton-proton pair will relax 16 times faster than a carbon-13 proton pair.

It is important to realize that in dipolar relaxation the effect is not primarily to distribute the energy from one of the spins to the other. This would not, on its own, bring the spins to equilibrium. Rather, the dipolar interaction provides a path by which energy can be transferred between the lattice and the spins. In this case, the lattice is the molecular motion. Essentially, the dipole-dipole interaction turns molecular motion into an oscillating magnetic field which can cause transitions of the spins.

Relation to the NOE

The dipolar mechanism is the only common relaxation mechanism which can cause transitions in which more than one spin flips. Specifically, with reference to section 8.3, the dipolar mechanism gives rise to transitions between the $\alpha\alpha$ and $\beta\beta$ states (W_2) and between the $\alpha\beta$ and $\beta\alpha$ states (W_0).

The rate constant W_2 corresponds to transitions which are at the sum of the Larmor frequencies of the two spins, ($\omega_{0,I} + \omega_{0,S}$) and so it is the spectral density at this sum frequency which is relevant. In contrast, W_0 corresponds to transitions at ($\omega_{0,I} - \omega_{0,S}$) and so for these it is the spectral density at this difference frequency which is relevant.

In the case where the two spins are the same (e.g. two protons) the two relevant spectral densities are $J(2\omega_0)$ and $J(0)$. In the fast motion limit ($\omega_0\tau_c \ll 1$) $J(2\omega_0)$ is somewhat less than $J(0)$, but not by very much. A detailed

calculation shows that $W_2 > W_0$ and so we expect to see positive NOE enhancements (section 8.4.5). In contrast, in the slow motion limit ($\omega_0 \tau_c \gg 1$) $J(2\omega_0)$ is all but zero and so $J(0) \gg J(2\omega_0)$; not surprisingly it follows that $W_0 > W_2$ and a negative NOE enhancement is seen.

8.6.3 The chemical shift anisotropy mechanism

The chemical shift arises because, due to the effect of the electrons in a molecule, the magnetic field experienced by a nucleus is different to that applied to the sample. In liquids, all that is observable is the average chemical shift, which results from the molecule rapidly experiencing all possible orientations by rapid molecular tumbling.

At a more detailed level, the magnetic field experienced by the nucleus depends on the orientation of the molecule relative to the applied magnetic field. This is called chemical shift anisotropy (CSA). In addition, it is not only the magnitude of the field which is altered but also its *direction*. The changes are very small, but sufficient to be detectable in the spectrum and to give rise to relaxation.

One convenient way of imagining the effect of CSA is to say that due to it there are small additional fields created at the nucleus – in general in all three directions. These fields vary in size as the molecule reorients, and so they have the necessary time variation to cause relaxation. As has already been discussed, it is the transverse fields which will give rise to changes in population.

The size of the CSA is specified by a tensor, which is a mathematical object represented by a three by three matrix.

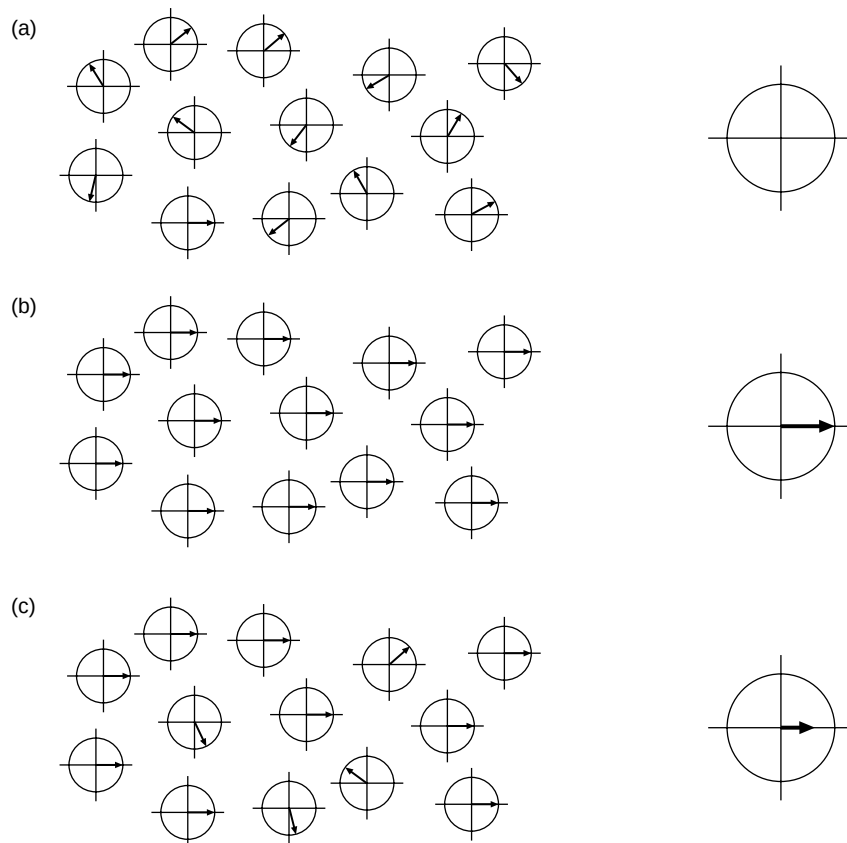
$$\sigma = \begin{pmatrix} \sigma_{xx} & \sigma_{xy} & \sigma_{xz} \\ \sigma_{yx} & \sigma_{yy} & \sigma_{yz} \\ \sigma_{zx} & \sigma_{zy} & \sigma_{zz} \end{pmatrix}$$

The element σ_{xz} gives the size of the extra field in the x-direction which results from a field being applied in the z-direction; likewise, σ_{yz} gives the extra field in the y-direction and σ_{zz} that in the z-direction. These elements depend on the electronic properties of the molecule and the orientation of the molecule with respect to the magnetic field.

Detailed calculations show that the relaxation induced by CSA goes as the square of the field strength and is also proportional to the shift anisotropy. A rough estimate of the size of this anisotropy is that it is equal to the typical shift range. So, CSA relaxation is expected to be significant for nuclei with large shift ranges observed at high fields. It is usually insignificant for protons.

8.7 Transverse relaxation

Right at the start of this section we mentioned that relaxation involved two processes: the populations returning to equilibrium and the transverse magnetization decaying to zero. So far, we have only discussed the first of these two. The second, in which the transverse magnetization decays, is called *transverse* (or *spin-spin*) relaxation.



Depiction of how the individual contributions from different spins (shown on the left) add up to give the net transverse magnetization (on the right). See text for details.

Each spin in the sample can be thought of as giving rise to a small contribution to the magnetization; these contributions can be in any direction, and in general have a component along x , y and z . The individual contributions along z add up to give the net z -magnetization of the sample.

The transverse contributions behave in a more complex way as, just like the net transverse magnetization, these contributions are precessing at the Larmor frequency in the transverse plane. We can represent each of these contributions by a vector precessing in the transverse plane.

The direction in which these vectors point can be specified by giving each a phase – arbitrarily the angle measured around from the x -axis. It is immediately clear that if these phases are random the net transverse magnetization of the sample will be zero as all the individual contributions will cancel. This is the situation that pertains at equilibrium and is shown in (a) in the figure above.

For there to be net magnetization, the phases must not be random, rather there has to be a preference for one direction; this is shown in (b) in the figure above. In quantum mechanics this is described as a *coherence*. An RF pulse applied to equilibrium magnetization generates transverse magnetization, or in other words the pulse generates a coherence. Transverse relaxation *destroys* this coherence by destroying the alignment of the individual contributions, as shown in (c) above.

Our picture indicates that there are two ways in which the coherence could be destroyed. The first is to make the vectors jump to new positions,

at random. Drawing on our analogy between these vectors and the behaviour of the bulk magnetization, we can see that these jumps could be brought about by local oscillating fields which have the same effect as pulses.

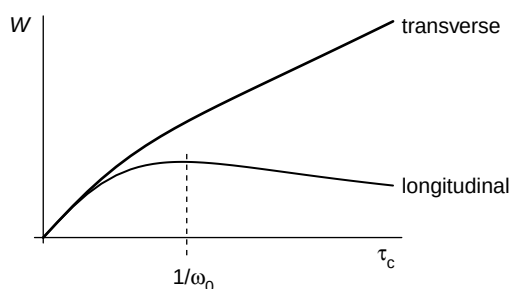
This is exactly what causes longitudinal relaxation, in which we imagine the local fields causing the spins to flip. So, *anything* that causes longitudinal relaxation will *also cause* transverse relaxation.

The second way of destroying the coherence is to make the vectors get out of step with one another as a result of them precessing at *different* Larmor frequencies. Again, a local field plays the part we need but this time we do not need it to oscillate; rather, all we need for it to do is to be different at different locations in the sample.

This latter contribution is called the *secular* part of transverse relaxation; the part which has the same origin as longitudinal relaxation is called the *non-secular* part.

It turns out that the secular part depends on the spectral density at zero frequency, $J(0)$. We can see that this makes sense as this part of transverse relaxation requires no transitions, just a field to cause a local variation in the magnetic field. Looking at the result from section 8.5.2 we see that $J(0) = 2\tau_c$, and so as the correlation time gets longer and longer, so too does the relaxation rate constant. Thus large molecules in the slow motion limit are characterised by very rapid transverse relaxation; this is in contrast to longitudinal relaxation is most rapid for a particular value of the correlation time.

The plot below compares the behaviour of the longitudinal and transverse relaxation rate constants. As the correlation time increases the longitudinal rate constant goes through a maximum. However, the transverse rate constant carries on increasing and shows no such maximum. We can attribute this to the secular part of transverse relaxation which depends on $J(0)$ and which simply goes on increasing as the correlation time increases. Detailed calculations show that in the fast motion limit the two relaxation rate constants are equal.



Comparison of the longitudinal and transverse relaxation rate constants as a function of the correlation time for the fixed Larmor frequency. The longitudinal rate constant shows a maximum, but the transverse rate constant simply goes on increasing.

9 Coherence Selection: Phase Cycling and Gradient Pulses[†]

9.1 Introduction

The pulse sequence used in an NMR experiment is carefully designed to produce a particular outcome. For example, we may wish to pass the spins through a state of multiple quantum coherence at a particular point, or plan for the magnetization to be aligned along the z -axis during a mixing period. However, it is usually the case that the particular series of events we designed the pulse sequence to cause is only one out of many possibilities. For example, a pulse whose role is to generate double-quantum coherence from anti-phase magnetization may also generate zero-quantum coherence or transfer the magnetization to another spin. Each time a radiofrequency pulse is applied there is this possibility of branching into many different pathways. If no steps are taken to suppress these unwanted pathways the resulting spectrum will be hopelessly confused and uninterpretable.

There are two general ways in which one pathway can be isolated from the many possible. The first is *phase cycling*. In this method the phases of the pulses and the receiver are varied in a systematic way so that the signal from the desired pathways adds and signal from all other pathways cancels. Phase cycling requires that the experiment is repeated several times, something which is probably required in any case in order to achieve the required signal-to-noise ratio.

The second method of selection is to use *field gradient pulses*. These are short periods during which the applied magnetic field is made inhomogeneous. As a result, any coherences present dephase and are apparently lost. However, this dephasing can be undone, and the coherence restored, by application of a subsequent gradient. We shall see that this dephasing and rephasing approach can be used to select particular coherences. Unlike phase cycling, the use of field gradient pulses does not require repetition of the experiment.

Both of these selection methods can be described in a unified framework which classifies the coherences present at any particular point according to a *coherence order* and then uses *coherence transfer pathways* to specify the desired outcome of the experiment.

9.2 Phase in NMR

In NMR we have control over both the phase of the pulses and the *receiver phase*. The effect of changing the phase of a pulse is easy to visualise in the usual rotating frame. So, for example, a 90° pulse about the x -axis rotates magnetization from z onto $-y$, whereas the same pulse applied about the y -axis rotates the magnetization onto the x -axis. The idea of the receiver phase is slightly more complex and will be explored in this section.

[†] Chapter 9 "Coherence Selection: Phase Cycling and Gradient Pulses" © James Keeler 2001, 2003 and 2004.

The NMR signal – that is the free induction decay – which emerges from the probe is a radiofrequency signal oscillating at close to the Larmor frequency (usually hundreds of MHz). Within the spectrometer this signal is shifted down to a much lower frequency in order that it can be digitized and then stored in computer memory. The details of this down-shifting process can be quite complex, but the overall result is simply that a fixed frequency, called the *receiver reference* or *carrier*, is subtracted from the frequency of the incoming NMR signal. Frequently, this receiver reference frequency is the same as the transmitter frequency used to generate the pulses applied to the observed nucleus. We shall assume that this is the case from now on.

The rotating frame which we use to visualise the effect of pulses is set at the transmitter frequency, ω_{rf} , so that the field due to the radiofrequency pulse is static. In this frame, a spin whose Larmor frequency is ω_0 precesses at $(\omega_0 - \omega_{\text{rf}})$, called the offset Ω . In the spectrometer the incoming signal at ω_0 is down-shifted by subtracting the receiver reference which, as we have already decided, will be equal to the frequency of the radiofrequency pulses. So, in effect, the frequencies of the signals which are digitized in the spectrometer are the offset frequencies at which the spins evolve in the rotating frame. Often this whole process is summarised by saying that the "signal is detected in the rotating frame".

9.1.1 Detector phase

The quantity which is actually detected in an NMR experiment is the transverse magnetization. Ultimately, this appears at the probe as an oscillating voltage, which we will write as

$$S_{\text{FID}} = \cos \omega_0 t$$

where ω_0 is the Larmor frequency. The down-shifting process in the spectrometer is achieved by an electronic device called a *mixer*; this effectively multiplies the incoming signal by a locally generated reference signal, S_{ref} , which we will assume is oscillating at ω_{rf}

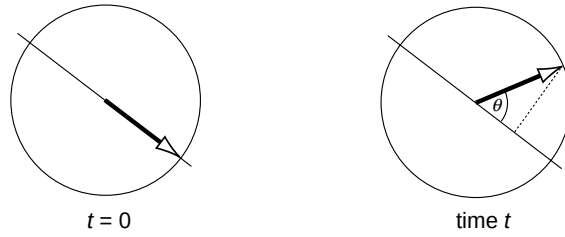
$$S_{\text{ref}} = \cos \omega_{\text{rf}} t$$

The output of the mixer is the product $S_{\text{FID}} S_{\text{ref}}$

$$\begin{aligned} S_{\text{FID}} S_{\text{ref}} &= A \cos \omega_{\text{rf}} t \cos \omega_0 t \\ &= \frac{1}{2} A [\cos(\omega_{\text{rf}} + \omega_0) t + \cos(\omega_{\text{rf}} - \omega_0) t] \end{aligned}$$

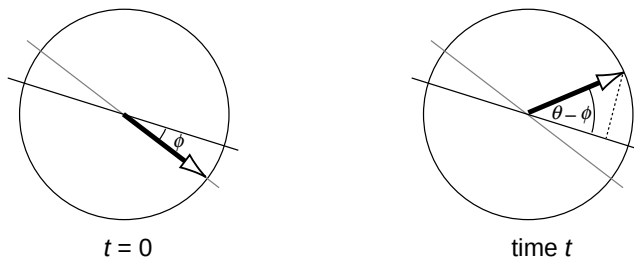
The first term is an oscillation at a high frequency (of the order of twice the Larmor frequency as $\omega_0 \approx \omega_{\text{rf}}$) and is easily removed by a low-pass filter. The second term is an oscillation at the offset frequency, Ω . This is in line with the previous comment that this down-shifting process is equivalent to detecting the precession in the rotating frame.

We can go further with this interpretation and say that the second term represents the component of the magnetization along a particular axis (called the *reference axis*) in the rotating frame. Such a component varies as $\cos \Omega t$, assuming that at time zero the magnetization is aligned along the chosen axis; this is illustrated below



At time zero the magnetization is assumed to be aligned along the reference axis. After time t the magnetization has precessed through an angle $\theta = \Omega t$. The projection of the magnetization onto the reference axis is proportional to $\cos \Omega t$.

Suppose now that the phase of the reference signal is shifted by ϕ , something which is easily achieved in the spectrometer. Effectively, this shifts the reference axis by ϕ , as shown below



Shifting the phase of the receiver reference by ϕ is equivalent to detecting the component along an axis rotated by ϕ from its original position (the previous axis is shown in grey). Now the apparent angle of precession is $\theta = \Omega t - \phi$ and the projection of the magnetization onto the reference axis is proportional to $\cos(\Omega t - \phi)$.

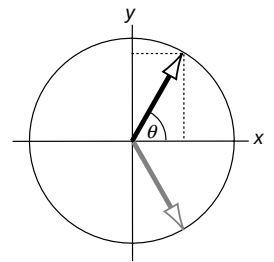
The component along the new reference axis is proportional to $\cos(\Omega t - \phi)$. How this is put to good effect is described in the next section.

9.1.2 Quadrature detection

We normally want to place the transmitter frequency in the centre of the resonances of interest so as to minimise off-resonance effects. If the receiver reference frequency is the same as the transmitter frequency, it immediately follows that the offset frequencies, Ω , may be both positive and negative. However, as we have seen in the previous section the effect of the down-shifting scheme used is to generate a signal of the form $\cos \Omega t$. Since $\cos(\theta) = \cos(-\theta)$ such a signal does not discriminate between positive and negative offset frequencies. The resulting spectrum, obtained by Fourier transformation, would be confusing as each peak would appear at both $+\Omega$ and $-\Omega$.

The way out of this problem is to detect the signal along two *perpendicular* axes. As is illustrated opposite, the projection along one axis is proportional to $\cos(\Omega t)$ and to $\sin(\Omega t)$ along the other. Knowledge of both of these projections enables us to work out the sense of rotation of the vector *i.e.* the sign of the offset.

The sin modulated component is detected by having a second mixer fed with a reference whose phase is shifted by $\pi/2$. Following the above discussion the output of the mixer is



The x and y projections of the black vector are both positive. If the vector had precessed in the opposite direction (shown shaded), and at the same frequency, the projection along x would be the same, but along y it would be minus that of the black vector.

$$\begin{aligned}\cos(\Omega t - \pi/2) &= \cos \Omega t \cos \pi/2 + \sin \Omega t \sin \pi/2 \\ &= \sin \Omega t\end{aligned}$$

The output of these two mixers can be regarded as being the components of the magnetization along perpendicular axes in the rotating frame.

The final step in this whole process is regard the outputs of the two mixers as being the real and imaginary parts of a complex number:

$$\cos \Omega t + i \sin \Omega t = \exp(i\Omega t)$$

The overall result is the generation of a complex signal whose phase varies according to the offset frequency Ω .

9.1.3 Control of phase

In the previous section we supposed that the signal coming from the probe was of the form $\cos \omega_0 t$ but it is more realistic to write the signal as $\cos(\omega_0 t + \phi_{\text{sig}})$ in recognition of the fact that in addition to a frequency the signal has a phase, ϕ_{sig} . This phase is a combination of factors that are not under our control (such as phase shifts produced in the amplifiers and filters through which the signal passes) and a phase due to the pulse sequence, which certainly is under our control.

The overall result of this phase is simply to multiply the final complex signal by a phase factor, $\exp(i\phi_{\text{sig}})$:

$$\exp(i\Omega t) \exp(i\phi_{\text{sig}})$$

As we saw in the previous section, we can also introduce another phase shift by altering the phase of the reference signal fed to the mixer, and indeed we saw that the cosine and sine modulated signals are generated by using two mixers fed with reference signals which differ in phase by $\pi/2$. If the phase of each of these reference signals is advanced by ϕ_{rx} , usually called the receiver phase, the output of the two mixers becomes $\cos(\Omega t - \phi_{\text{rx}})$ and $\sin(\Omega t - \phi_{\text{rx}})$. In the complex notation, the overall signal thus acquires another phase factor

$$\exp(i\Omega t) \exp(i\phi_{\text{sig}}) \exp(-i\phi_{\text{rx}})$$

Overall, then, the phase of the final signal depends on the *difference* between the phase introduced by the pulse sequence and the phase introduced by the receiver reference.

9.1.4 Lineshapes

Let us suppose that the signal can be written

$$S(t) = B \exp(i\Omega t) \exp(i\Phi) \exp(-t/T_2)$$

where Φ is the overall phase ($= \phi_{\text{sig}} - \phi_{\text{rx}}$) and B is the amplitude. The term, $\exp(-t/T_2)$ has been added to impose a decay on the signal. Fourier transformation of $S(t)$ gives the spectrum $S(\omega)$:

$$S(\omega) = B[A(\omega) + iD(\omega)] \exp(i\Phi) \quad [1]$$

where $A(\omega)$ is an absorption mode lorentzian lineshape centred at $\omega = \Omega$ and $D(\omega)$ is the corresponding dispersion mode lorentzian:

$$A(\omega) = \frac{T_2}{1 + (\omega - \Omega)^2 T_2^2} \quad D(\omega) = \frac{(\omega - \Omega) T_2^2}{1 + (\omega - \Omega)^2 T_2^2}$$

Normally we display just the real part of $S(\omega)$ which is, in this case,

$$\text{Re}[S(\omega)] = B[\cos \Phi A(\omega) - \sin \Phi D(\omega)]$$

In general this is a mixture of the absorption and dispersion lineshape. If we want just the absorption lineshape we need to somehow set Φ to zero, which is easily done by multiplying $S(\omega)$ by a phase factor $\exp(i\Theta)$.

$$\begin{aligned} S(\omega) \exp(i\Theta) &= B[A(\omega) + iD(\omega)] \exp(i\Phi) \exp(i\Theta) \\ &= B[A(\omega) + iD(\omega)] \exp(i[\Phi + \Theta]) \end{aligned}$$

As this is a numerical operation which can be carried out in the computer we are free to choose Θ to be the required value (here $-\Phi$) in order to remove the phase factor entirely and hence give an absorption mode spectrum in the real part. This is what we do when we "phase the spectrum".

9.1.5 Relative phase and lineshape

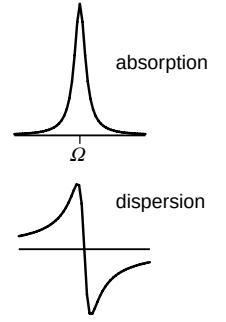
We have seen that we can alter the phase of the spectrum by altering the phase of the pulse or of the receiver, but that what really counts is the difference in these two phases.

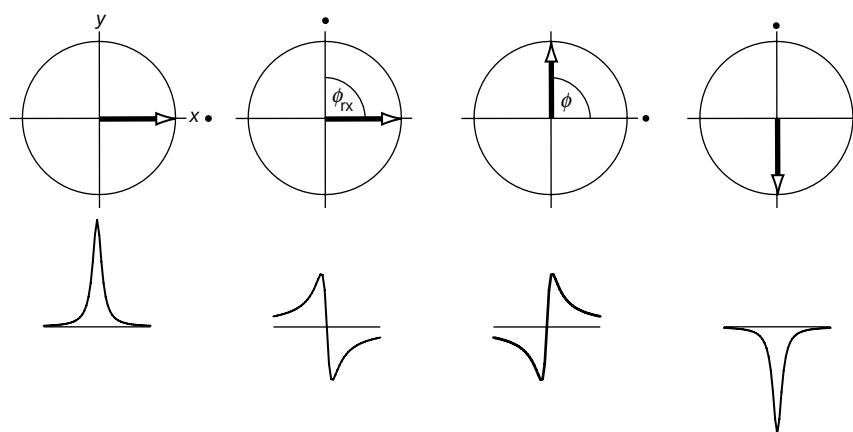
We will illustrate this with the simple vector diagrams shown below. Here, the vector shows the position of the magnetization at time zero and its phase, ϕ_{sig} , is measured anti-clockwise from the x -axis. The dot shows the axis along which the receiver is aligned; this phase, ϕ_{rx} , is also measured anti-clockwise from the x -axis.

If the vector and receiver are aligned along the same axis, $\Phi = 0$, and the real part of the spectrum shows the absorption mode lineshape. If the receiver phase is advanced by $\pi/2$, $\Phi = 0 - \pi/2$ and, from Eq. [1]

$$\begin{aligned} S(\omega) &= B[A(\omega) + iD(\omega)] \exp(-i\pi/2) \\ &= B[-iA(\omega) + D(\omega)] \end{aligned}$$

This means that the real part of the spectrum shows a dispersion lineshape. On the other hand, if the magnetization is advanced by $\pi/2$, $\Phi = \phi_{\text{sig}} - \phi_{\text{rx}} = \pi/2 - 0 = \pi/2$ and it can be shown from Eq. [1] that the real part of the spectrum shows a negative dispersion lineshape. Finally, if either phase is advanced by π , the result is a negative absorption lineshape.

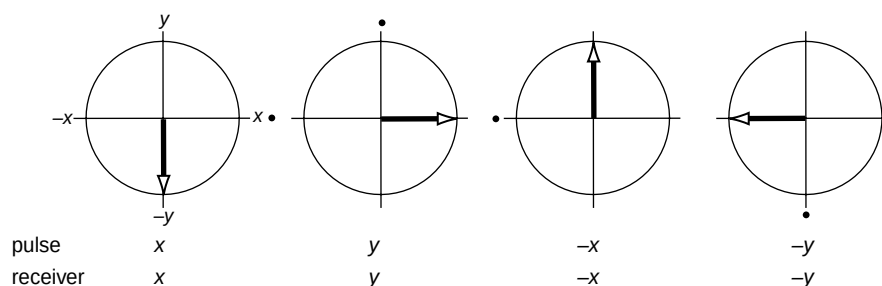




9.1.6 CYCLOPS

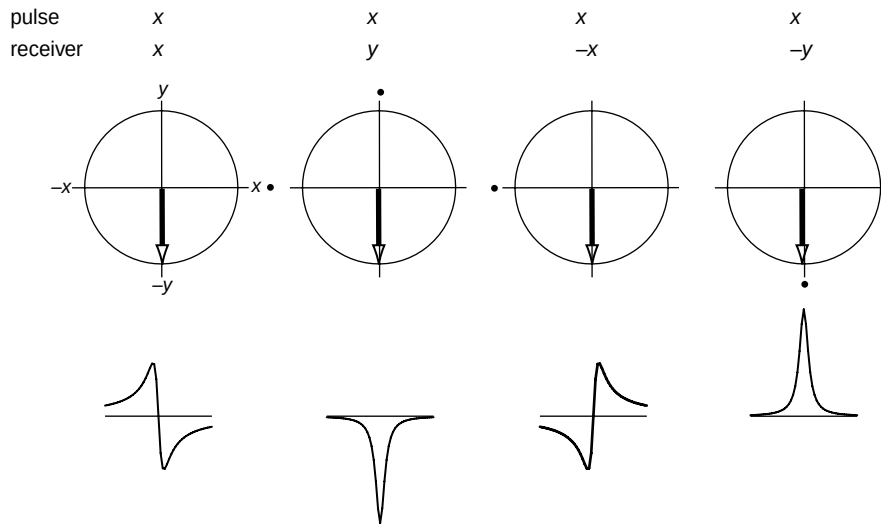
The CYCLOPS phase cycling scheme is commonly used in even the simplest pulse-acquire experiments. The sequence is designed to cancel some imperfections associated with errors in the two phase detectors mentioned above; a description of how this is achieved is beyond the scope of this discussion. However, the cycle itself illustrates very well the points made in the previous section.

There are four steps in the cycle, the pulse phase goes $x, y, -x, -y$ *i.e.* it advances by 90° on each step; likewise the receiver advances by 90° on each step. The figure below shows how the magnetization and receiver phases are related for the four steps of this cycle



Although both the receiver and the magnetization shift phase on each step, the phase difference between them remains constant. Each step in the cycle thus gives the same lineshape and so the signal adds on all four steps, which is just what is required.

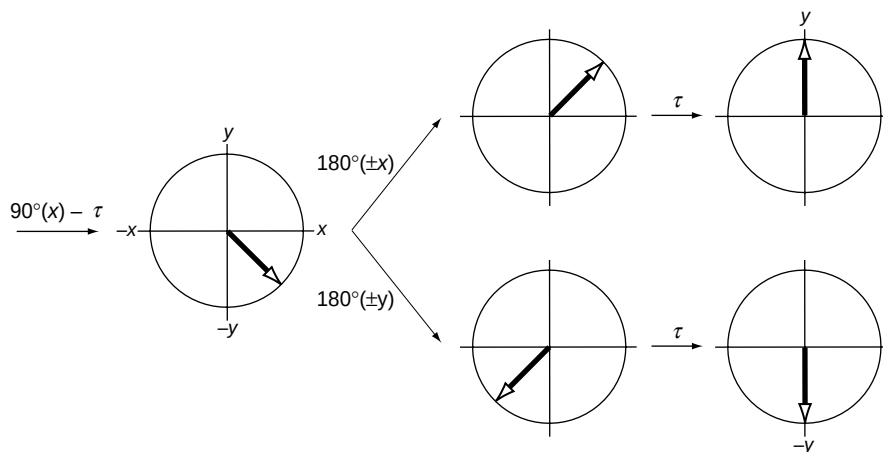
Suppose that we forget to advance the pulse phase; the outcome is quite different



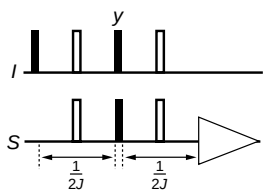
Now the phase difference between the receiver and the magnetization is no longer constant. A different lineshape thus results from each step and it is clear that adding all four together will lead to complete cancellation (steps 2 and 4 cancel, as do steps 1 and 3). For the signal to add up it is clearly essential for the receiver to follow the magnetization.

9.1.7 EXORCYCLE

EXORCYCLE is perhaps the original phase cycle. It is a cycle used for 180° pulses when they form part of a spin echo sequence. The 180° pulse cycles through the phases $x, y, -x, -y$ and the receiver phase goes $x, -x, x, -x$. The diagram below illustrates the outcome of this sequence



If the phase of the 180° pulse is $+x$ or $-x$ the echo forms along the y -axis, whereas if the phase is $\pm y$ the echo forms on the $-y$ axis. Therefore, as the 180° pulse is advanced by 90° (e.g. from x to y) the receiver must be advanced by 180° (e.g. from x to $-x$). Of course, we could just as well cycle the receiver phases $y, -y, y, -y$; all that matters is that they advance in steps of 180° . We will see later on how it is that this phase cycle cancels out the results of imperfections in the 180° pulse.



Pulse sequence for INEPT. Filled rectangles represent 90° pulses and open rectangles represent 180° pulses. Unless otherwise indicated, all pulses are of phase x .

9.1.8 Difference spectroscopy

Often a simple two step sequence suffices to cancel unwanted magnetization; essentially this is a form of difference spectroscopy. The idea is well illustrated by the INEPT sequence, shown opposite. The aim of the sequence is to transfer magnetization from spin I to a coupled spin S .

With the phases and delays shown equilibrium magnetization of spin I , I_z , is transferred to spin S , appearing as the operator S_x . Equilibrium magnetization of S , S_z , appears as S_y . We wish to preserve only the signal that has been transferred from I .

The procedure to achieve this is very simple. If we change the phase of the second I spin 90° pulse from y to $-y$ the magnetization arising from transfer of the I spin magnetization to S becomes $-S_x$ i.e. it changes sign. In contrast, the signal arising from equilibrium S spin magnetization is unaffected simply because the S_z operator is unaffected by the I spin pulses. By repeating the experiment twice, once with the phase of the second I spin 90° pulse set to y and once with it set to $-y$, and then subtracting the two resulting signals, the undesired signal is cancelled and the desired signal adds. It is easily confirmed that shifting the phase of the S spin 90° pulse does not achieve the desired separation of the two signals as both are affected in the same way.

In practice the subtraction would be carried out by shifting the receiver by 180° , so the I spin pulse would go $y, -y$ and the receiver phase go $x, -x$. This is a two step phase cycle which is probably best viewed as difference spectroscopy.

This simple two step cycle is the basic element used in constructing the phase cycling of many two- and three-dimensional heteronuclear experiments.

9.2 Coherence transfer pathways

Although we can make some progress in writing simple phase cycles by considering the vector picture, a more general framework is needed in order to cope with experiments which involve multiple-quantum coherence and related phenomena. We also need a theory which enables us to predict the degree to which a phase cycle can discriminate against different classes of unwanted signals. A convenient and powerful way of doing both these things is to use the coherence transfer pathway approach.

9.2.1 Coherence order

Coherences, of which transverse magnetization is one example, can be classified according to a coherence order, p , which is an integer taking values $0, \pm 1, \pm 2 \dots$. Single quantum coherence has $p = \pm 1$, double has $p = \pm 2$ and so on; z -magnetization, " zz " terms and zero-quantum coherence have $p = 0$. This classification comes about by considering the phase which different coherences acquire in response to a rotation about the z -axis.

A coherence of order p , represented by the density operator $\sigma^{(p)}$, evolves under a z -rotation of angle ϕ according to

$$\exp(-i\phi F_z) \sigma^{(p)} \exp(i\phi F_z) = \exp(-i p \phi) \sigma^{(p)} \quad [2]$$

where F_z is the operator for the total z-component of the spin angular momentum. In words, a coherence of order p experiences a phase shift of $-p\phi$. Equation [2] is the definition of coherence order.

To see how this definition can be applied, consider the effect of a z-rotation on transverse magnetization aligned along the x-axis. Such a rotation is identical in nature to that due to evolution under an offset, and using product operators it can be written

$$\exp(-i\phi I_z) I_x \exp(i\phi I_z) = \cos \phi I_x + \sin \phi I_y \quad [3]$$

The right hand sides of Eqs. [2] and [3] are not immediately comparable, but by writing the sine and cosine terms as complex exponentials the comparison becomes clearer. Using

$$\cos \phi = \frac{1}{2} [\exp(i\phi) + \exp(-i\phi)] \quad \sin \phi = \frac{1}{2i} [\exp(i\phi) - \exp(-i\phi)]$$

Eq. [3] becomes

$$\begin{aligned} & \exp(-i\phi I_z) I_x \exp(i\phi I_z) \\ &= \frac{1}{2} [\exp(i\phi) + \exp(-i\phi)] I_x + \frac{1}{2i} [\exp(i\phi) - \exp(-i\phi)] I_y \\ &= \frac{1}{2} [I_x + \frac{1}{i} I_y] \exp(i\phi) + \frac{1}{2} [I_x - \frac{1}{i} I_y] \exp(-i\phi) \end{aligned}$$

It is now clear that the first term corresponds to coherence order -1 and the second to $+1$; in other words, I_x is an equal mixture of coherence orders ± 1 .

The cartesian product operators do not correspond to a single coherence order so it is more convenient to rewrite them in terms of the raising and lowering operators, I_+ and I_- , defined as

$$I_+ = I_x + iI_y \quad I_- = I_x - iI_y$$

from which it follows that

$$I_x = \frac{1}{2} [I_+ + I_-] \quad I_y = \frac{1}{2i} [I_+ - I_-] \quad [4]$$

Under z-rotations the raising and lowering operators transform simply

$$\exp(-i\phi I_z) I_{\pm} \exp(i\phi I_z) = \exp(\mp i\phi) I_{\pm}$$

which, by comparison with Eq. [2] shows that I_+ corresponds to coherence order $+1$ and I_- to -1 . So, from Eq. [4] we can see that I_x and I_y correspond to mixtures of coherence orders $+1$ and -1 .

As a second example consider the pure double quantum operator for two coupled spins,

$$2I_{1x}I_{2y} + 2I_{1y}I_{2x}$$

Rewriting this in terms of the raising and lowering operators gives

$$\frac{1}{i} (I_1^+ I_2^+ - I_1^- I_2^-)$$

The effect of a z-rotation on the term $I_1^+ I_2^+$ is found as follows:

$$\begin{aligned} & \exp(-i\phi I_{1z}) \exp(-i\phi I_{2z}) I_{1+} I_{2+} \exp(i\phi I_{2z}) \exp(i\phi I_{1z}) \\ &= \exp(-i\phi I_{1z}) \exp(-i\phi) I_{1+} I_{2+} \exp(i\phi I_{1z}) \\ &= \exp(-i\phi) \exp(-i\phi) I_{1+} I_{2+} = \exp(-2i\phi) I_{1+} I_{2+} \end{aligned}$$

Thus, as the coherence experiences a phase shift of -2ϕ the coherence is classified according to Eq. [2] as having $p = 2$. It is easy to confirm that the term $I_{1-}I_{2-}$ has $p = -2$. Thus the pure double quantum term, $2I_{1x}I_{2y} + 2I_{1y}I_{2x}$, is an equal mixture of coherence orders $+2$ and -2 .

As this example indicates, it is possible to determine the order or orders of any state by writing it in terms of raising and lowering operators and then simply inspecting the number of such operators in each term. A raising operator contributes $+1$ to the coherence order whereas a lowering operator contributes -1 . A z-operator, I_{iz} , has coherence order 0 as it is invariant to z-rotations.

Coherences involving heteronuclei can be assigned both an overall order and an order with respect to each nuclear species. For example the term $I_{1+}S_{1-}$ has an overall order of 0, is order $+1$ for the I spins and -1 for the S spins. The term $I_{1+}I_{2+}S_{1z}$ is overall of order 2, is order 2 for the I spins and is order 0 for the S spins.

9.2.2 Evolution under offsets

The evolution under an offset, Ω , is simply a z-rotation, so the raising and lowering operators simply acquire a phase Ωt

$$\exp(-i\Omega t I_z) I_{\pm} \exp(i\Omega t I_z) = \exp(\mp i\Omega t) I_{\pm}$$

For products of these operators, the overall phase is the sum of the phases acquired by each term

$$\begin{aligned} & \exp(-i\Omega_j t I_{jz}) \exp(-i\Omega_i t I_{iz}) I_{i-} I_{j+} \exp(i\Omega_i t I_{iz}) \exp(i\Omega_j t I_{jz}) \\ &= \exp(i(\Omega_i - \Omega_j)t) I_{i-} I_{j+} \end{aligned}$$

It also follows that coherences of opposite sign acquire phases of opposite signs under free evolution. So the operator $I_{1+}I_{2+}$ (with $p = 2$) acquires a phase $-(\Omega_1 + \Omega_2)t$ *i.e.* it evolves at a frequency $-(\Omega_1 + \Omega_2)$ whereas the operator $I_{1-}I_{2-}$ (with $p = -2$) acquires a phase $(\Omega_1 + \Omega_2)t$ *i.e.* it evolves at a frequency $(\Omega_1 + \Omega_2)$. We will see later on that this observation has important consequences for the lineshapes in two-dimensional NMR.

The observation that coherences of different orders respond differently to evolution under a z-rotation (*e.g.* an offset) lies at the heart of the way in which gradient pulses can be used to separate different coherence orders.

9.2.3 Phase shifted pulses

In general, a radiofrequency pulse causes coherences to be transferred from one order to one or more different orders; it is this spreading out of the coherence which makes it necessary to select one transfer among many possibilities. An example of this spreading between coherence orders is the effect of a non-selective pulse on antiphase magnetization, such as $2I_{1x}I_{2z}$, which corresponds to coherence orders ± 1 . Some of the coherence may be transferred into double- and zero-quantum coherence, some may be transferred into two-spin order and some will remain unaffected. The precise outcome depends on the phase and flip angle of the pulse, but in general we can see that there are many possibilities.

If we consider just one coherence, of order p , being transferred to a coherence of order p' by a radiofrequency pulse we can derive a very general result for the way in which the phase of the pulse affects the phase of the coherence. It is on this relationship that the phase cycling method is based.

We will write the initial state of order p as $\sigma^{(p)}$, and the final state of order p' as $\sigma^{(p')}$. The effect of the radiofrequency pulse causing the transfer is represented by the (unitary) transformation U_ϕ where ϕ is the phase of the pulse. The initial and final states are related by the usual transformation

$$U_0 \sigma^{(p)} U_0^{-1} = \sigma^{(p')} + \text{terms of other orders} \quad [5]$$

which has been written for phase 0; the other terms will be dropped as we are only interested in the transfer from p to p' . The transformation brought about by a radiofrequency pulse phase shifted by ϕ , U_ϕ , is related to that with the phase set to zero, U_0 , in the following way

$$U_\phi = \exp(-i\phi F_z) U_0 \exp(i\phi F_z) \quad [6]$$

Using this, the effect of the phase shifted pulse on the initial state $\sigma^{(p)}$ can be written

$$\begin{aligned} U_\phi \sigma^{(p)} U_\phi^{-1} \\ = \exp(-i\phi F_z) U_0 \exp(i\phi F_z) \sigma^{(p)} \exp(-i\phi F_z) U_0^{-1} \exp(i\phi F_z) \end{aligned} \quad [7]$$

The central three terms can be simplified by application of Eq. [2]

$$\exp(i\phi F_z) \sigma^{(p)} \exp(-i\phi F_z) U_0^{-1} = \exp(ip\phi) \sigma^{(p)}$$

giving

$$U_\phi \sigma^{(p)} U_\phi^{-1} = \exp(ip\phi) \exp(-i\phi F_z) U_0 \sigma^{(p)} U_0^{-1} \exp(i\phi F_z)$$

The central three terms can, from Eq. [5], be replaced by $\sigma^{(p')}$ to give

$$U_\phi \sigma^{(p)} U_\phi^{-1} = \exp(ip\phi) \exp(-i\phi F_z) \sigma^{(p')} \exp(i\phi F_z)$$

Finally, Eq. [5] is applied again to give

$$U_\phi \sigma^{(p)} U_\phi^{-1} = \exp(ip\phi) \exp(-ip'\phi) \sigma^{(p')}$$

Defining $\Delta p = (p' - p)$ as the change in coherence order, this simplifies to

$$U_\phi \sigma^{(p)} U_\phi^{-1} = \exp(-i\Delta p \phi) \sigma^{(p')} \quad [8]$$

Equation [8] says that if the phase of a pulse which is causing a change in coherence order of Δp is shifted by ϕ the coherence will acquire a phase label $(-\Delta p \phi)$. It is this property which enables us to separate different changes in coherence order from one another by altering the phase of the pulse.

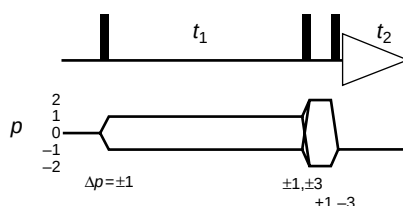
In the discussion so far it has been assumed that U_ϕ represents a single pulse. However, any sequence of pulses and delays can be represented by a single unitary transformation, so Eq. [8] applies equally well to the effect of phase shifting all of the pulses in such a sequence. We will see that this property is often of use in writing phase cycles.

If a series of phase shifted pulses (or pulse sandwiches) are applied a

phase ($-\Delta p \phi$) is acquired from each. The total phase is found by adding up these individual contributions. In an NMR experiment this total phase affects the signal which is recorded at the end of the sequence, even though the phase shift may have been acquired earlier in the pulse sequence. These phase shifts are, so to speak, carried forward.

9.2.4 Coherence transfer pathways diagrams

In designing a multiple-pulse NMR experiment the intention is to have specific orders of coherence present at various points in the sequence. One way of indicating this is to use a *coherence transfer pathway* (CTP) diagram along with the timing diagram for the pulse sequence. An example of shown below, which gives the pulse sequence and CTP for the DQF COSY experiment.



The solid lines under the sequence represent the coherence orders required during each part of the sequence; note that it is only the pulses which cause a change in the coherence order. In addition, the values of Δp are shown for each pulse. In this example, as is commonly the case, more than one order of coherence is present at a particular time. Each pulse is required to cause different changes to the coherence order – for example the second pulse is required to bring about no less than four values of Δp . Again, this is a common feature of pulse sequences.

It is important to realise that the CTP specified with the pulse sequence is just the *desired* pathway. We would need to establish separately (for example using a product operator calculation) that the pulse sequence is indeed capable of generating the coherences specified in the CTP. Also, the spin system which we apply the sequence to has to be capable of supporting the coherences. For example, if there are no couplings, then no double quantum will be generated and thus selection of the above pathway will result in a null spectrum.

The coherence transfer pathway must start with $p = 0$ as this is the order to which equilibrium magnetization (z-magnetization) belongs. In addition, the pathway has to end with $|p| = 1$ as it is only single quantum coherence that is observable. If we use quadrature detection (section 9.1.2) it turns out that only one of $p = \pm 1$ is observable; we will follow the usual convention of assuming that $p = -1$ is the detectable signal.

9.3 Lineshapes and frequency discrimination

9.3.1 Phase and amplitude modulation

The selection of a particular CTP has important consequences for lineshapes and frequency discrimination in two-dimensional NMR. These topics are illustrated using the NOESY experiment as an example; the pulse sequence and CTP is illustrated opposite.

If we imagine starting with I_z , then at the end of t_1 the operators present are

$$-\cos \Omega t_1 I_y + \sin \Omega t_1 I_x$$

The term in I_y is rotated onto the z-axis and we will assume that only this term survives. Finally, the z-magnetization is made observable by the last pulse (for convenience set to phase $-y$) giving the observable term present at $t_2 = 0$ as

$$\cos \Omega t_1 I_x$$

As was noted in section 9.2.1, I_x is in fact a mixture of coherence orders $p = \pm 1$, something which is made evident by writing the operator in terms of I_+ and I_-

$$\frac{1}{2} \cos \Omega t_1 (I_+ + I_-)$$

Of these operators, only I_- leads to an observable signal, as this corresponds to $p = -1$. Allowing I_- to evolve in t_2 gives

$$\frac{1}{2} \cos \Omega t_1 \exp(i\Omega t_2) I_-$$

The final detected signal can be written as

$$S_c(t_1, t_2) = \frac{1}{2} \cos \Omega t_1 \exp(i\Omega t_2)$$

This signal is said to be *amplitude modulated* in t_1 ; it is so called because the evolution during t_1 gives rise, via the cosine term, to a modulation of the amplitude of the observed signal.

The situation changes if we select a different pathway, as shown opposite. Here, only coherence order -1 is preserved during t_1 . At the start of t_1 the operator present is $-I_y$ which can be written

$$-\frac{1}{2i} (I_+ - I_-)$$

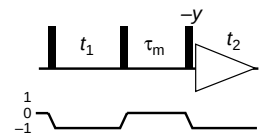
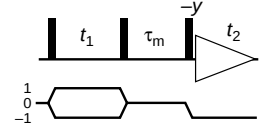
Now, in accordance with the CTP, we select only the I_- term. During t_1 this evolves to give

$$\frac{1}{2i} \exp(i\Omega t_1) I_-$$

Following through the rest of the pulse sequence as before gives the following observable signal

$$S_p(t_1, t_2) = \frac{1}{4} \exp(i\Omega t_1) \exp(i\Omega t_2)$$

This signal is said to be *phase modulated* in t_1 ; it is so called because the evolution during t_1 gives rise, via exponential term, to a modulation of the phase of the observed signal. If we had chosen to select $p = +1$ during t_1 the signal would have been



$$S_N(t_1, t_2) = \frac{1}{4} \exp(-i\Omega t_1) \exp(i\Omega t_2)$$

which is also phase modulated, except in the opposite sense. Note that in either case the phase modulated signal is one half of the size of the amplitude modulated signal, because only one of the two pathways has been selected.

Although these results have been derived for the NOESY sequence, they are in fact general for any two-dimensional experiment. Summarising, we find

- If a single coherence order is present during t_1 the result is phase modulation in t_1 . The phase modulation can be of the form $\exp(i\Omega t_1)$ or $\exp(-i\Omega t_1)$ depending on the sign of the coherence order present.
- If both coherence orders $\pm p$ are selected during t_1 , the result is amplitude modulation in t_1 ; selecting both orders in this way is called *preserving symmetrical pathways*.

9.3.2 Frequency discrimination

The amplitude modulated signal contains no information about the sign of Ω , simply because $\cos(\Omega t_1) = \cos(-\Omega t_1)$. As a consequence, Fourier transformation of the time domain signal will result in each peak appearing twice in the two-dimensional spectrum, once at $F_1 = +\Omega$ and once at $F_1 = -\Omega$. As was commented on above, we usually place the transmitter in the middle of the spectrum so that there are peaks with both positive and negative offsets. If, as a result of recording an amplitude modulated signal, all of these appear twice, the spectrum will hopelessly be confused. A spectrum arising from an amplitude modulated signal is said to *lack frequency discrimination* in F_1 .

On the other hand, the phase modulated signal is sensitive to the sign of the offset and so information about the sign of Ω in the F_1 dimension is contained in the signal. Fourier transformation of the signal $S_p(t_1, t_2)$ gives a peak at $F_1 = +\Omega$, $F_2 = \Omega$, whereas Fourier transformation of the signal $S_N(t_1, t_2)$ gives a peak at $F_1 = -\Omega$, $F_2 = \Omega$. Both spectra are said to be frequency discriminated as the sign of the modulation frequency in t_1 is determined; in contrast to amplitude modulated spectra, each peak will only appear once.

The spectrum from $S_p(t_1, t_2)$ is called the P-type (P for positive) or echo spectrum; a diagonal peak appears with the same sign of offset in each dimension. The spectrum from $S_N(t_1, t_2)$ is called the N-type (N for negative) or anti-echo spectrum; a diagonal peak appears with opposite signs in the two dimensions.

It might appear that in order to achieve frequency discrimination we should deliberately select a CTP which leads to a P- or an N-type spectrum. However, such spectra show a very unfavourable lineshape, as discussed in the next section.

9.3.3 Lineshapes

In section 9.1.4 we saw that Fourier transformation of the signal

$$S(t) = \exp(i\Omega t) \exp(-t/T_2)$$

gave a spectrum whose real part is an absorption lorentzian and whose imaginary part is a dispersion lorentzian:

$$S(\omega) = A(\omega) + iD(\omega)$$

We will use the shorthand that A_2 represents an absorption mode lineshape at $F_2 = \Omega$ and D_2 represents a dispersion mode lineshape at the same frequency. Likewise, A_{1+} represents an absorption mode lineshape at $F_1 = +\Omega$ and D_{1+} represents the corresponding dispersion lineshape. A_{1-} and D_{1-} represent the corresponding lines at $F_1 = -\Omega$.

The time domain signal for the P-type spectrum can be written as

$$S_p(t_1, t_2) = \frac{1}{4} \exp(i\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

where the damping factors have been included as before. Fourier transformation with respect to t_2 gives

$$S_p(t_1, F_2) = \frac{1}{4} \exp(i\Omega t_1) \exp(-t_1/T_2) [A_2 + iD_2]$$

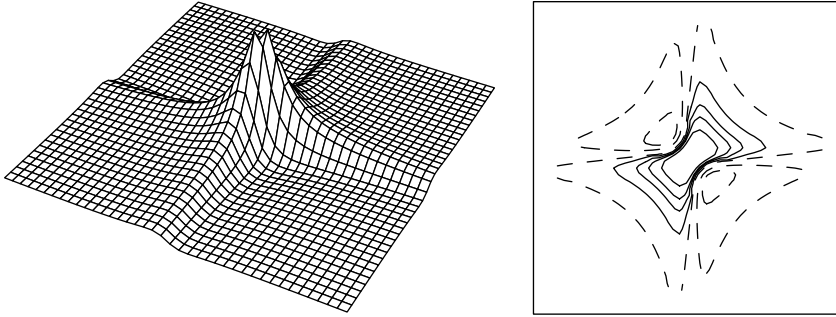
and then further transformation with respect to t_1 gives

$$S_p(F_1, F_2) = \frac{1}{4} [A_{1+} + iD_{1+}] [A_2 + iD_2]$$

The real part of this spectrum is

$$\text{Re}\{S_p(F_1, F_2)\} = \frac{1}{4} [A_{1+}A_2 - D_{1+}D_2]$$

The quantity in the square brackets on the right represents a phase-twist lineshape at $F_1 = +\Omega$, $F_2 = \Omega$



Perspective view and contour plot of the phase-twist lineshape. Negative contours are shown dashed.

This lineshape is an inextricable mixture of absorption and dispersion, and it is very undesirable for high-resolution NMR. So, although a phase modulated signal gives us frequency discrimination, which is desirable, it also results in a phase-twist lineshape, which is not.

The time domain signal for the amplitude modulated data set can be written as

$$S_c(t_1, t_2) = \frac{1}{2} \cos(\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

Fourier transformation with respect to t_2 gives

$$S_c(t_1, F_2) = \frac{1}{2} \cos(\Omega t_1) \exp(-t_1/T_2) [A_2 + iD_2]$$

which can be rewritten as

$$S_C(t_1, F_2) = \frac{1}{4} [\exp(i\Omega t_1) + \exp(-i\Omega t_1)] \exp(-t_1/T_2) [A_2 + iD_2]$$

Fourier transformation with respect to t_1 gives, in the real part of the spectrum

$$\text{Re}\{S_C(F_1, F_2)\} = \frac{1}{4} [A_{1+}A_2 - D_{1+}D_2] + \frac{1}{4} [A_{1-}A_2 - D_{1-}D_2]$$

This corresponds to two phase-twist lineshapes, one at $F_1 = +\Omega$, $F_2 = \Omega$ and the other at $F_1 = -\Omega$, $F_2 = \Omega$; the lack of frequency discrimination is evident. Further, the undesirable phase-twist lineshape is again present.

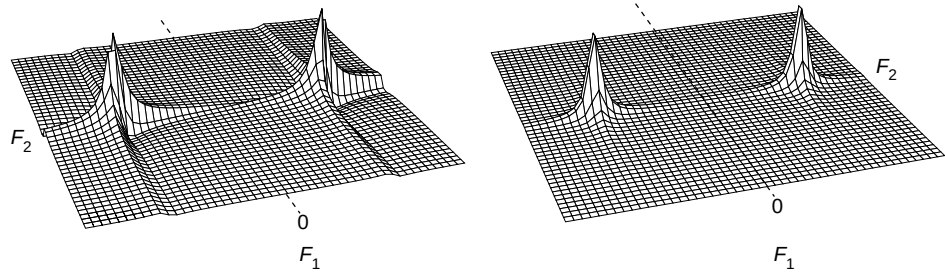
The lineshape can be restored to the absorption mode by discarding the imaginary part of the time domain signal after the transformation with respect to t_2 , i.e. by taking the real part

$$\text{Re}\{S_C(t_1, F_2)\} = \frac{1}{2} \cos(\Omega t_1) \exp(-t_1/T_2) A_2$$

Subsequent transformation with respect to t_1 gives, in the real part

$$\frac{1}{4} A_{1+}A_2 + \frac{1}{4} A_{1-}A_2$$

which is two double absorption mode lineshapes. Frequency discrimination is lacking, but the lineshape is now much more desirable. The spectra with the two phase-twist and two absorption mode lines are shown below on the left and right, respectively.



9.3.4 Frequency discrimination with retention of absorption mode lineshapes

For practical purposes it is essential to be able to achieve frequency discrimination and at the same time retain the absorption mode lineshape. There are a number of ways of doing this.

9.3.4.1 States-Haberkorn-Ruben (SHR) method

The key to this method is the ability to record a cosine modulated data set and a sine modulated data set. The latter can be achieved simply by changing the phase of appropriate pulses. For example, in the case of the NOESY experiment, all that is required to generate the sine data set is to shift the phase of the first 90° pulse by 90° (in fact in the NOESY sequence the pulse needs to shift from x to $-y$). The two data sets have to kept separate.

The cosine data set is transformed with respect to t_1 and the imaginary part discarded to give

$$\text{Re}\{S_C(t_1, F_2)\} = \frac{1}{2} \cos(\Omega t_1) \exp(-t_1/T_2) A_2 \quad [9]$$

The same operation is performed on the sine modulated data set

$$S_s(t_1, t_2) = \frac{1}{2} \sin(\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

$$\text{Re}\{S_s(t_1, F_2)\} = \frac{1}{2} \sin(\Omega t_1) \exp(-t_1/T_2) A_2 \quad [10]$$

A new complex data set is now formed by using the signal from Eq. [9] as the real part and that from Eq. [10] as the imaginary part

$$S_{\text{SHR}}(t_1, F_2) = \text{Re}\{S_c(t_1, F_2)\} + i \text{Re}\{S_s(t_1, F_2)\}$$

$$= \frac{1}{2} \exp(i\Omega t_1) \exp(-t_1/T_2) A_2$$

Fourier transformation with respect to t_1 gives, in the real part of the spectrum

$$\text{Re}\{S_{\text{SHR}}(F_1, F_2)\} = \frac{1}{2} A_1 A_2$$

This is the desired frequency discriminated spectrum with a pure absorption lineshape.

As commented on above, in NOESY all that is required to change from cosine to sine modulation is to shift the phase of the first pulse by 90° . The general recipe is to shift the phase of *all* the pulses that precede t_1 by $90^\circ/|p_1|$, where p_1 is the coherence order present during t_1 . So, for a double quantum spectrum, the phase shift needs to be 45° . The origin of this rule is that, taken together, the pulses which precede t_1 give rise to a pathway with $\Delta p = p_1$.

In heteronuclear experiments it is not usually necessary to shift the phase of all the pulses which precede t_1 ; an analysis of the sequence usually shows that shifting the phase of the pulse which generates the transverse magnetization which evolves during t_1 is sufficient.

9.3.4.2 Echo anti-echo method

We will see in later sections that when we use gradient pulses for coherence selection the natural outcome is P- or N-type data sets. Individually, each of these gives a frequency discriminated spectrum, but with the phase-twist lineshape. We will show in this section how an absorption mode lineshape can be obtained provided *both* the P- and the N-type data sets are available.

As before, we write the two data sets as

$$S_p(t_1, t_2) = \frac{1}{4} \exp(i\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

$$S_n(t_1, t_2) = \frac{1}{4} \exp(-i\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

We then form the two combinations

$$S_c(t_1, t_2) = S_p(t_1, t_2) + S_n(t_1, t_2)$$

$$= \frac{1}{2} \cos(\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

$$S_s(t_1, t_2) = \frac{1}{i} [S_p(t_1, t_2) - S_n(t_1, t_2)]$$

$$= \frac{1}{2} \sin(\Omega t_1) \exp(i\Omega t_2) \exp(-t_1/T_2) \exp(-t_2/T_2)$$

These cosine and sine modulated data sets can be used as inputs to the SHR method described in the previous section.

An alternative is to Fourier transform the two data sets with respect to t_2 to give

$$S_p(t_1, F_2) = \frac{1}{4} \exp(i\Omega t_1) \exp(-t_1/T_2) [A_2 + iD_2]$$

$$S_N(t_1, F_2) = \frac{1}{4} \exp(-i\Omega t_1) \exp(-t_1/T_2) [A_2 + iD_2]$$

We then take the complex conjugate of $S_N(t_1, F_2)$ and add it to $S_p(t_1, F_2)$

$$S_N(t_1, F_2)^* = \frac{1}{4} \exp(i\Omega t_1) \exp(-t_1/T_2) [A_2 - iD_2]$$

$$\begin{aligned} S_+(t_1, F_2) &= S_N(t_1, F_2)^* + S_p(t_1, F_2) \\ &= \frac{1}{2} \exp(i\Omega t_1) \exp(-t_1/T_2) A_2 \end{aligned}$$

Transformation of this signal gives

$$S_+(F_1, F_2) = \frac{1}{2} [A_{1+} + iD_{1+}] A_2$$

which is frequency discriminated and has, in the real part, the required double absorption lineshape.

9.3.4.3 Marion-Wüthrich or TPPI method

The idea behind the TPPI (time proportional phase incrementation) or Marion–Wüthrich (MW) method is to arrange things so that all of the peaks have positive offsets. Then, frequency discrimination is not required as there is no ambiguity.

One simple way to make all offsets positive is to set the receiver carrier frequency deliberately at the edge of the spectrum. Simple though this is, it is not really a very practical method as the resulting spectrum would be very inefficient in its use of data space and in addition off-resonance effects associated with the pulses in the sequence will be accentuated.

In the TPPI method the carrier can still be set in the middle of the spectrum, but it is made to *appear* that all the frequencies are positive by phase shifting some of the pulses in the sequence *in concert* with the incrementation of t_1 .

It was noted above that shifting the phase of the first pulse in the NOESY sequence from x to $-y$ caused the modulation to change from $\cos(\Omega t_1)$ to $\sin(\Omega t_1)$. One way of expressing this is to say that shifting the pulse causes a phase shift ϕ in the signal modulation, which can be written $\cos(\Omega t_1 + \phi)$. Using the usual trigonometric expansions this can be written

$$\cos(\Omega t_1 + \phi) = \cos \Omega t_1 \cos \phi - \sin \Omega t_1 \sin \phi$$

If the phase shift, ϕ , is $-\pi/2$ radians the result is

$$\begin{aligned} \cos(\Omega t_1 + \pi/2) &= \cos \Omega t_1 \cos(-\pi/2) - \sin \Omega t_1 \sin(-\pi/2) \\ &= \sin \Omega t_1 \end{aligned}$$

This is exactly the result we found before.

In the TPPI procedure, the phase ϕ is made proportional to t_1 *i.e.* each time t_1 is incremented, so is the phase. We will suppose that

$$\phi(t_1) = \omega_{\text{add}} t_1$$

The constant of proportion between the time dependent phase, $\phi(t_1)$, and t_1 has been written ω_{add} ; ω_{add} has the dimensions of rad s^{-1} *i.e.* it is a frequency. Following the same approach as before, the time-domain function with the inclusion of this incrementing phase is thus

$$\begin{aligned}\cos(\Omega t_1 + \phi(t_1)) &= \cos(\Omega t_1 + \omega_{\text{add}} t_1) \\ &= \cos(\Omega + \omega_{\text{add}}) t_1\end{aligned}$$

In words, the effect of incrementing the phase in concert with t_1 is to add a frequency ω_{add} to *all* of the offsets in the spectrum. The TPPI method utilizes this in the following way.

In one-dimensional pulse-Fourier transform NMR the free induction signal is sampled at regular intervals Δ . After transformation the resulting spectrum displays correctly peaks with offsets in the range $-(SW/2)$ to $+(SW/2)$ where SW is the spectral width which is given by $1/\Delta$ (this comes about from the Nyquist theorem of data sampling). Frequencies outside this range are not represented correctly.

Suppose that the required frequency range in the F_1 dimension is from $-(SW_1/2)$ to $+(SW_1/2)$. To make it appear that all the peaks have a positive offset, it will be necessary to add $(SW_1/2)$ to all the frequencies. Then the peaks will be in the range 0 to (SW_1) .

As the maximum frequency is now (SW_1) rather than $(SW_1/2)$ the sampling interval, Δ_1 , will have to be halved *i.e.* $\Delta_1 = 1/(2SW_1)$ in order that the range of frequencies present are represented properly.

The phase increment is $\omega_{\text{add}} t_1$, but t_1 can be written as $n\Delta_1$ for the n th increment of t_1 . The required value for ω_{add} is $2\pi(SW_1/2)$, where the 2π is to convert from frequency (the units of SW_1) to rad s^{-1} , the units of ω_{add} . Putting all of this together $\omega_{\text{add}} t_1$ can be expressed, for the n th increment as

$$\begin{aligned}\omega_{\text{additional}} t_1 &= 2\pi \left(\frac{SW_1}{2} \right) (n\Delta_1) \\ &= 2\pi \left(\frac{SW_1}{2} \right) \left(n \frac{1}{2SW_1} \right) \\ &= n \frac{\pi}{2}\end{aligned}$$

In words this means that each time t_1 is incremented, the phase of the signal should also be incremented by 90° , for example by incrementing the phase of one of the pulses.

A data set from an experiment to which TPPI has been applied is simply amplitude modulated in t_1 and so can be processed according to the method described above for cosine modulated data so as to obtain absorption mode lineshapes. As the spectrum is symmetrical about $F_1 = 0$, it is usual to use a modified Fourier transform routine which saves effort and space by only calculating the positive frequency part of the spectrum.

9.3.4.4 States-TPPI

When the SHR method is used, axial peaks (arising from magnetization which has not evolved during t_1) appear at $F_1 = 0$; such peaks can be a nuisance as they may obscure other wanted peaks. We will see below (section 9.4.6) that axial peaks can be suppressed with the aid of phase cycling, all be it at the cost of doubling the length of the phase cycle.

The States-TPPI method does not suppress these axial peaks, but moves

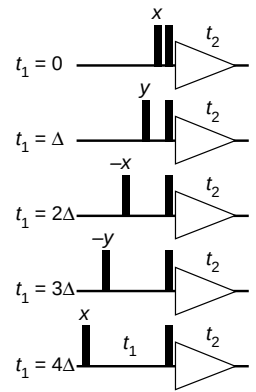


Illustration of the TPPI method. Each time that t_1 is incremented, so is the phase of the pulse preceding t_1 .

them to the edge of the spectrum so that they are less likely to obscure wanted peaks. All that is involved is that, each time t_1 is incremented, both the phase of the pulse which precedes t_1 and the receiver phase are advanced by 180° *i.e.* the pulse goes $x, -x$ and the receiver goes $x, -x$.

For non-axial peaks, the two phase shifts cancel one another out, and so have no effect. However, magnetization which gives rise to axial peaks does not experience the first phase shift, but does experience the receiver phase shift. The sign alternation in concert with t_1 incrementation adds a frequency of $SW_1/2$ to each peak, thus shifting it to the edge of the spectrum. Note that in States-TPPI the spectral range in the F_1 dimension is $-(SW_1/2)$ to $+(SW_1/2)$ and the sampling interval is $1/2SW_1$, just as in the SHR method.

The nice feature of States-TPPI is that it moves the axial peaks out of the way without lengthening the phase cycle. It is therefore convenient to use in complex three- and four-dimensional spectra where phase cycling is at a premium.

9.4 Phase cycling

In this section we will start out by considering in detail how to write a phase cycle to select a particular value of Δp and then use this discussion to lead on to the formulation of general principles for constructing phase cycles. These will then be used to construct appropriate cycles for a number of common experiments.

9.4.1 Selection of a single pathway

To focus on the issue at hand let us consider the case of transferring from coherence order $+2$ to order -1 . Such a transfer has $\Delta p = (-1 - (2)) = -3$. Let us imagine that the pulse causing this transformation is cycled around the four cardinal phases ($x, y, -x, -y$, *i.e.* $0^\circ, 90^\circ, 180^\circ, 270^\circ$) and draw up a table of the phase shift that will be experienced by the transferred coherence. This is simply computed as $-\Delta p \phi$, in this case $= -(-3)\phi = 3\phi$.

step	pulse phase	phase shift experienced by transfer with $\Delta p = -3$	equivalent phase
1	0	0	0
2	90	270	270
3	180	540	180
4	270	810	90

The fourth column, labelled "equivalent phase", is just the phase shift experienced by the coherence, column three, reduced to be in the range 0 to 360° by subtracting multiples of 360° (*e.g.* for step 3 we subtracted 360° and for step 4 we subtracted 720°).

If we wished to select $\Delta p = -3$ we would simply shift the phase of the receiver in order to match the phase that the coherence has acquired; these are the phases shown in the last column. If we did this, then each step of the cycle would give an observed signal of the same phase and so the four contributions would all add up. This is precisely the same thing as we did

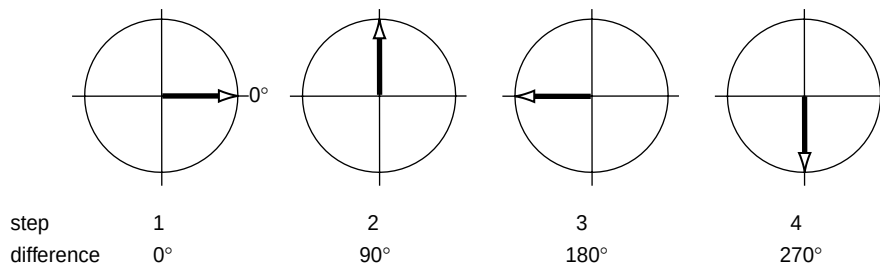
when considering the CYCLOPS sequence in section 9.1.6; in both cases the receiver phase follows the phase of the desired magnetization or coherence.

We now need to see if this four step phase cycle eliminates the signals from other pathways. As an example, let us consider a pathway with $\Delta p = 2$, which might arise from the transfer from coherence order -1 to $+1$. Again we draw up a table to show the phase experienced by a pathway with $\Delta p = 2$, that is computed as $-(2)\phi$

step	pulse phase	phase shift experienced by transfer with $\Delta p = 2$	equivalent phase	rx. phase to select $\Delta p = -3$	difference
1	0	0	0	0	0
2	90	-180	180	270	$270 - 180 = 90$
3	180	-360	0	180	$180 - 0 = 180$
4	270	-540	180	90	$90 - 180 = -90$

As before, the equivalent phase is simply the phase in column 3 reduced to the range 0 to 360° . The fifth column shows the receiver (abbreviated to rx.) phases determined above for selection of the transfer with $\Delta p = -3$. The question we have to ask is whether or not these phase shifts will lead to cancellation of the transfer with $\Delta p = 2$. To do this we compute the difference between the receiver phase, column 5, and the phase shift experienced by the transfer with $\Delta p = 2$, column 4. The results are shown in column 6, labelled "difference". Looking at this difference column we can see that step 1 will cancel with step 3 as the 180° phase shift between them means that the two signals have opposite sign. Likewise step 2 will cancel with step 4 as there is a 180° phase shift between them. We conclude, therefore, that this four step cycle cancels the signal arising from a pathway with $\Delta p = 2$.

An alternative way of viewing the cancellation is to represent the results of the "difference" column by vectors pointing at the indicated angles. This is shown below; it is clear that the opposed vectors cancel one another.



Next we consider the coherence transfer with $\Delta p = +1$. Again, we draw up the table and calculate the phase shifts experience by this transfer, which are given by $-(+1)\phi = -\phi$.

step	pulse phase	phase shift experienced by transfer with $\Delta p = +1$	equivalent phase	rx. phase to select $\Delta p = -3$	difference
1	0	0	0	0	0
2	90	-90	270	270	270 - 270 = 0
3	180	-180	180	180	180 - 180 = 0
4	270	-270	90	90	90 - 90 = 0

Here we see quite different behaviour. The equivalent phases, that is the phase shifts experienced by the transfer with $\Delta p = 1$, match exactly the receiver phases determined for $\Delta p = -3$, thus the phases in the "difference" column are all zero. We conclude that the four step cycle selects transfers both with $\Delta p = -3$ and $+1$.

Some more work with tables such as these will reveal that this four step cycle suppresses contributions from changes in coherence order of -2 , -1 and 0 . It selects $\Delta p = -3$ and 1 . It also selects changes in coherence order of 5 , 9 , 13 and so on. This latter sequence is easy to understand. A pathway with $\Delta p = 1$ experiences a phase shift of -90° when the pulse is shifted in phase by 90° ; the equivalent phase is thus 270° . A pathway with $\Delta p = 5$ would experience a phase shift of $-5 \times 90^\circ = -450^\circ$ which corresponds to an equivalent phase of 270° . Thus the phase shifts experienced for $\Delta p = 1$ and 5 are identical and it is clear that a cycle which selects one will select the other. The same goes for the series $\Delta p = 9, 13 \dots$

The extension to negative values of Δp is also easy to see. A pathway with $\Delta p = -3$ experiences a phase shift of 270° when the pulse is shifted in phase by 90° . A transfer with $\Delta p = +1$ experiences a phase of -90° which corresponds to an equivalent phase of 270° . Thus both pathways experience the same phase shifts and a cycle which selects one will select the other. The pattern is clear, this four step cycle will select a pathway with $\Delta p = -3$, as it was designed to, and also it will select any pathway with $\Delta p = -3 + 4n$ where $n = \pm 1, \pm 2, \pm 3 \dots$

9.4.2 General Rules

The discussion in the previous section can be generalised in the following way. Consider a phase cycle in which the phase of a pulse takes N evenly spaced steps covering the range 0 to 2π radians. The phases, ϕ_k , are

$$\phi_k = 2\pi k/N \text{ where } k = 0, 1, 2 \dots (N - 1).$$

To select a change in coherence order, Δp , the receiver phase is set to $-\Delta p \times \phi_k$ for each step and the resulting signals are summed. This cycle will, in addition to selecting the specified change in coherence order, also select pathways with changes in coherence order $(\Delta p \pm nN)$ where $n = \pm 1, \pm 2 \dots$

The way in which phase cycling selects a series of values of Δp which are related by a harmonic condition is closely related to the phenomenon of aliasing in Fourier transformation. Indeed, the whole process of phase cycling can be seen as the computation of a discrete Fourier transformation with respect to the pulse phase; in this case the Fourier co-domains are phase and coherence order.

The fact that a phase cycle inevitably selects more than one change in coherence order is not necessarily a problem. We may actually wish to select more than one pathway, and examples of this will be given below in relation to specific two-dimensional experiments. Even if we only require one value of Δp we may be able to discount the values selected at the same time as being improbable or insignificant. In a system of m coupled spins one-half, the maximum order of coherence that can be generated is m , thus in a two spin system we need not worry about whether or not a phase cycle will discriminate between double quantum and six quantum coherences as the latter simply cannot be present. Even in more extended spin systems the likelihood of generating high-order coherences is rather small and so we may be able to discount them for all practical purposes. If a high level of discrimination between orders is needed, then the solution is simply to use a phase cycle which has more steps *i.e.* in which the phases move in smaller increments. For example a six step cycle will discriminate between $\Delta p = +2$ and $+6$, whereas a four step cycle will not.

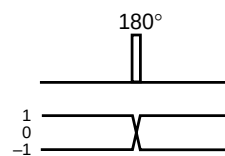
9.4.3 Refocusing Pulses

A 180° pulse simply changes the sign of the coherence order. This is easily demonstrated by considering the effect of such a pulse on the operators I_+ and I_- . For example:

$$I_+ \equiv (I_x + iI_y) \xrightarrow{\pi I_x} (I_x - iI_y) \equiv I_-$$

corresponds to $p = +1 \rightarrow p = -1$. In a more complex product of operators, each raising and lowering operator is affected in this way so overall the coherence order changes sign.

We can now derive the EXORCYCLE phase cycle using this property. Consider a 180° pulse acting on single quantum coherence, for which the CTP is shown opposite. For the pathway starting with $p = 1$ the effect of the 180° pulse is to cause a change with $\Delta p = -2$. The table shows a four-step cycle to select this change



A 180° pulse simply changes the sign of the coherence order. The EXORCYCLE phase cycling selects both of the pathways shown.

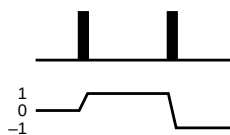
Step	phase of 180° pulse	phase shift experienced by transfer with $\Delta p = -2$	Equivalent phase = rx. phase
1	0	0	0
2	90	180	180
3	180	360	0
4	270	540	180

The phase cycle is thus $0, 90^\circ, 180^\circ, 270^\circ$ for the 180° pulse and $0^\circ, 180^\circ, 0^\circ, 180^\circ$ for the receiver; this is precisely the set of phases deduced before for EXORCYCLE in section 9.1.7.

As the cycle has four steps, a pathway with $\Delta p = +2$ is also selected; this is the pathway which starts with $p = -1$ and is transferred to $p = +1$. Therefore, the four steps of EXORCYCLE select both of the pathways shown in the diagram above.

A two step cycle, consisting of $0^\circ, 180^\circ$ for the 180° pulse and $0^\circ, 0^\circ$ for

the receiver, can easily be shown to select all *even* values of Δp . This reduced form of EXORCYCLE is sometimes used when it is necessary to minimise the number of steps in a phase cycle. An eight step cycle, in which the 180° pulse is advanced in steps of 45° , can be used to select the refocusing of double-quantum coherence in which the transfer is from $p = +2$ to -2 (i.e. $\Delta p = -4$) or *vice versa*.



9.4.4 Combining phase cycles

Suppose that we wish to select the pathway shown opposite; for the first pulse Δp is 1 and for the second it is -2 . We can construct a four-step cycle for each pulse, but to select the overall pathway shown these two cycles have to be completed *independently* of one another. This means that there will be a total of sixteen steps. The table shows how the appropriate receiver cycling can be determined

Step	phase of 1st pulse	phase for $\Delta p = 1$	phase of 2nd pulse	phase for $\Delta p = -2$	total phase	equivalent phase = rx. phase
1	0	0	0	0	0	0
2	90	-90	0	0	-90	270
3	180	-180	0	0	-180	180
4	270	-270	0	0	-270	90
5	0	0	90	180	180	180
6	90	-90	90	180	90	90
7	180	-180	90	180	0	0
8	270	-270	90	180	-90	270
9	0	0	180	360	360	0
10	90	-90	180	360	270	270
11	180	-180	180	360	180	180
12	270	-270	180	360	90	90
13	0	0	270	540	540	180
14	90	-90	270	540	450	90
15	180	-180	270	540	360	0
16	270	-270	270	540	270	270

In the first four steps the phase of the second pulse is held constant and the phase of the first pulse simply goes through the four steps 0° 90° 180° 270° . As we are selecting $\Delta p = 1$ for this pulse, the receiver phases are simply 0° , 270° , 180° , 90° .

Steps 5 to 8 are a repeat of steps 1–4 except that the phase of the second pulse has been moved by 90° . As Δp for the second pulse is -2 , the required pathway experiences a phase shift of 180° and so the receiver phase must be advanced by this much. So, the receiver phases for steps 5–8 are just 180° ahead of those for steps 1–4.

In the same way for steps 9–12 the first pulse again goes through the same four steps, and the phase of the second pulse is advanced to 180° . Therefore, compared to steps 1–4 the receiver phases in steps 9–12 need to be advanced by $-(-2) \times 180^\circ = 360^\circ = 0^\circ$. Likewise, the receiver phases for steps 13–16 are advanced by $-(-2) \times 270^\circ = 540^\circ = 180^\circ$.

Another way of looking at this is to consider each step individually. For example, compared to step 1, in step 14 the first pulse has been advanced by 90° so the phase from the first pulse is $-(1) \times 90^\circ = -90^\circ$. The second pulse has been advanced by 270° so the phase from this is $-(-2) \times 270^\circ = 540^\circ$. The total phase shift of the required pathway is thus $-90 + 540 = 450^\circ$ which is an equivalent phase of 90° . This is the receiver phase shown in the final column.

The key to devising these sequences is to simply work out the two four-step cycles independently and then merge them together rather than trying to work on the whole cycle. One writes down the first four steps, and then duplicates this four times as the second pulse is shifted. We would find the same steps, in a different sequence, if the phase of the *second* pulse is shifted in the *first* four steps.

We can see that the total size of a phase cycle grows at an alarming rate. With four phases for each pulse the number of steps grows as 4^l where l is the number of pulses in the sequence. A three-pulse sequence such as NOESY or DQF COSY would therefore involve a 64 step cycle. Such long cycles put a lower limit on the total time of an experiment and we may end up having to run an experiment for a long time not to achieve the desired signal-to-noise ratio but simply to complete the phase cycle.

Fortunately, there are several "tricks" which we can use in order to shorten the length of a phase cycle. To appreciate whether or not one of these tricks can be used in a particular sequence we need to understand in some detail what the sequence is actually doing and what the likely problems are going to be.

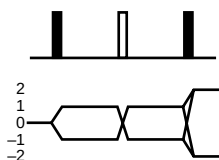
9.4.5 Tricks

9.4.5.1 The first pulse

All pulse sequences start with equilibrium magnetization, which has coherence order 0. It can easily be shown that when a pulse is applied to equilibrium magnetization the only coherence orders that can be generated are ± 1 . If retaining both of these orders is acceptable (which it often is), it is therefore not necessary to phase cycle the first pulse in a sequence.

There are two additional points to make here. If the spins have not relaxed completely by the start of the sequence the initial magnetization will not be at equilibrium. Then, the above simplification does not apply. Secondly, the first pulse of a sequence is often cycled in order to suppress axial peaks in two-dimensional spectra. This is considered in more detail in section 9.4.6.

9.4.5.2 Grouping pulses together



Pulse sequence for generating double-quantum coherence. Note that the 180° pulse simply causes a change in the sign of the coherence order.

The sequence shown opposite can be used to generate multiple quantum coherence from equilibrium magnetization; during the spin echo anti-phase magnetization develops and the final pulse transfers this into multiple quantum coherence. Let us suppose that we wish to generate double quantum, with $p = \pm 2$, as show by the CTP opposite.

As has already been noted, the first pulse can only generate $p = \pm 1$ and the 180° pulse only causes a change in the sign of the coherence order. The only pulse we need to be concerned with is the final one which we want to generate only double quantum. We could try to devise a phase cycle for the last pulse alone or we could simply group all three pulses together and imagine that, as a group, they achieve the transformation $p = 0$ to $p = \pm 2$ i.e. $\Delta p = \pm 2$. The phase cycle would simply be for the three pulses together to go $0^\circ, 90^\circ, 180^\circ, 270^\circ$, with the receiver going $0^\circ, 180^\circ, 0^\circ, 180^\circ$.

It has to be recognised that by cycling a group of pulses together we are only selecting an *overall* transformation; the coherence orders present within the group of pulses are not being selected. It is up to the designer of the experiment to decide whether or not this degree of selection is sufficient.

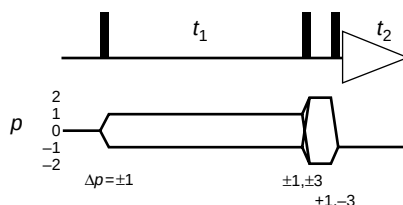
The four step cycle mentioned above also selects $\Delta p = \pm 6$; again, we would have to decide whether or not such high orders of coherence were likely to be present in the spin system. Finally, we note that the Δp values for the final pulse are $\pm 1, \pm 3$; it would not be possible to devise a four step cycle which selects *all* of these pathways.

9.4.5.3 The last pulse

We noted above that only coherence order -1 is observable. So, although the final pulse of a sequence may cause transfer to many different orders of coherence, only transfers to $p = -1$ will result in observable signals. Thus, if we have already selected, in an unambiguous way, a particular set of coherence orders present just before the last pulse, no further cycling of this pulse is needed.

9.4.5.4 Example – DQF COSY

A good example of the applications of these ideas is in devising a phase cycle for DQF COSY, whose pulse sequence and CTP is shown below.



Note that we have retained symmetrical pathways in t_1 so that absorption mode lineshapes can be obtained. Also, both in generating the double quantum coherence, and in reconverting it to observable magnetization, all possible pathways have been retained. If we do not do this, signal intensity is lost.

One way of viewing this sequence is to group the first two pulses

together and view them as achieving the transformation $0 \rightarrow \pm 2$ *i.e.* $\Delta p = \pm 2$. This is exactly the problem considered in section 9.4.5.2, where we saw that a suitable four step cycle is for the first two pulses to go $0^\circ, 90^\circ, 180^\circ, 270^\circ$ and the receiver to go $0^\circ, 180^\circ, 0^\circ, 180^\circ$. This unambiguously selects $p = \pm 2$ just before the last pulse, so phase cycling of the last pulse is not required (see section 9.4.5.3).

An alternative view is to say that as only $p = -1$ is observable, selecting the transformation $\Delta p = +1$ and -3 on the last pulse will be equivalent to selecting $p = \pm 2$ during the period just before the last pulse. Since the first pulse can only generate $p = \pm 1$ (present during t_1), the selection of $\Delta p = +1$ and -3 on the last pulse is sufficient to define the CTP completely.

A four step cycle to select $\Delta p = +1$ involves the pulse going $0^\circ, 90^\circ, 180^\circ, 270^\circ$ and the receiver going $0^\circ, 270^\circ, 180^\circ, 90^\circ$. As this cycle has four steps is automatically also selects $\Delta p = -3$, just as required.

The first of these cycles also selects $\Delta p = \pm 6$ for the first two pulses *i.e.* filtration through six-quantum coherence; normally, we can safely ignore the possibility of such high-order coherences. The second of the cycles also selects $\Delta p = +5$ and $\Delta p = -7$ on the last pulse; again, these transfers involve such high orders of multiple quantum that they can be ignored.

9.4.6 Axial peak suppression

Peaks are sometimes seen in two-dimensional spectra at co-ordinates $F_1 = 0$ and $F_2 =$ frequencies corresponding to the usual peaks in the spectrum. The interpretation of the appearance of these peaks is that they arise from magnetization which has not evolved during t_1 and so has not acquired a frequency label.

A common source of axial peaks is magnetization which recovers due to longitudinal relaxation during t_1 . Subsequent pulses make this magnetization observable, but it has no frequency label and so appears at $F_1 = 0$. Another source of axial peaks is when, due to pulse imperfections, not all of the original equilibrium magnetization is moved into the transverse plane by the first pulse. The residual longitudinal magnetization can be made observable by subsequent pulses and hence give rise to axial peaks.

A simple way of suppressing axial peaks is to select the pathway $\Delta p = \pm 1$ on the first pulse; this ensures that all signals arise from the first pulse. A two-step cycle in which the first pulse goes $0^\circ, 180^\circ$ and the receiver goes $0^\circ, 180^\circ$ selects $\Delta p = \pm 1$. It may be that the other phase cycling used in the sequence will also reject axial peaks so that it is not necessary to add an explicit axial peak suppression steps. Adding a two-step cycle for axial peak suppression doubles the length of the phase cycle.

9.4.7 Shifting the whole sequence – CYCLOPS

If we group *all* of the pulses in the sequence together and regard them as a unit they simply achieve the transformation from equilibrium magnetization, $p = 0$, to observable magnetization, $p = -1$. They could be cycled as a group to select this pathway with $\Delta p = -1$, that is the pulses

going 0° , 90° , 180° , 270° and the receiver going 0° , 90° , 180° , 270° . This is simple the CYCLOPS phase cycle described in section 9.1.6.

If time permits we sometimes add CYCLOPS-style cycling to all of the pulses in the sequence so as to suppress some artefacts associated with imperfections in the receiver. Adding such cycling does, of course, extend the phase cycle by a factor of four.

This view of the whole sequence as causing the transformation $\Delta p = -1$ also enables us to interchange receiver and pulse phase shifts. For example, suppose that a particular step in a phase cycle requires a receiver phase shift θ . The same effect can be achieved by shifting all of the pulses by $-\theta$ and leaving the receiver phase unaltered. The reason this works is that all of the pulses taken together achieve the transformation $\Delta p = -1$, so shifting their phases by $-\theta$ shift the signal by $-(-\theta) = \theta$, which is exactly the effect of shifting the receiver by θ . This kind of approach is sometimes helpful if hardware limitations mean that small angle phase-shifts are only available for the pulses.

9.4.8 Equivalent cycles

For even a relatively simple sequence such as DQF COSY there are a number of different ways of writing the phase cycle. Superficially these can look very different, but it may be possible to show that they really are the same.

For example, consider the DQF COSY phase cycle proposed in section 9.4.5.4 where we cycle just the last pulse

step	1st pulse	2nd pulse	3rd pulse	receiver
1	0	0	0	0
2	0	0	90	270
3	0	0	180	180
4	0	0	270	90

Suppose we decide that we do not want to shift the receiver phase, but want to keep it fixed at phase zero. As described above, this means that we need to subtract the receiver phase from all of the pulses. So, for example, in step 2 we subtract 270° from the pulse phases to give -270° , -270° and -180° for the phases of the first three pulses, respectively; reducing these to the usual range gives phases 90° , 90° and 180° . Doing the same for the other steps gives a rather strange looking phase cycle, but one which works in just the same way.

step	1st pulse	2nd pulse	3rd pulse	receiver
1	0	0	0	0
2	90	90	180	0
3	180	180	0	0
4	270	270	180	0

We can play one more trick with this phase cycle. As the third pulse is required to achieve the transformation $\Delta p = -3$ or $+1$ we can alter its phase by 180° and compensate for this by shifting the receiver by 180° also.

Doing this for steps 2 and 4 only gives

step	1st pulse	2nd pulse	3rd pulse	receiver
1	0	0	0	0
2	90	90	0	180
3	180	180	0	0
4	270	270	0	180

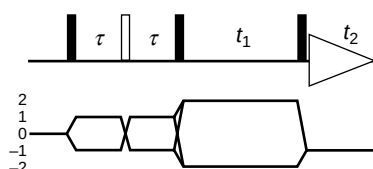
This is exactly the cycle proposed in section 9.4.5.4.

9.4.9 Further examples

In this section we will use a shorthand to indicate the phases of the pulses and the receiver. Rather than specifying the phase in degrees, the phases are expressed as multiples of 90° . So, EXORCYCLE becomes $\emptyset \ 1 \ 2 \ 3$ for the 180° pulse and $\emptyset \ 2 \ \emptyset \ 2$ for the receiver.

9.4.9.1 Double quantum spectroscopy

A simple sequence for double quantum spectroscopy is shown below

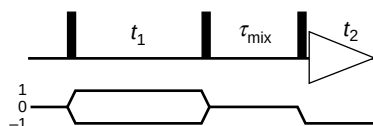


Note that both pathways with $p = \pm 1$ during the spin echo and with $p = \pm 2$ during t_1 are retained. There are a number of possible phase cycles for this experiment and, not surprisingly, they are essentially the same as those for DQF COSY. If we regard the first three pulses as a unit, then they are required to achieve the overall transformation $\Delta p = \pm 2$, which is the same as that for the first two pulses in the DQF COSY sequence. Thus the same cycle can be used with these three pulses going $\emptyset \ 1 \ 2 \ 3$ and the receiver going $\emptyset \ 2 \ \emptyset \ 2$. Alternatively the final pulse can be cycled $\emptyset \ 1 \ 2 \ 3$ with the receiver going $\emptyset \ 3 \ 2 \ 1$, as in section 9.4.5.4.

Both of these phase cycles can be extended by EXORCYCLE phase cycling of the 180° pulse, resulting in a total of 16 steps.

9.4.9.2 NOESY

The pulse sequence for NOESY (with retention of absorption mode lineshapes) is shown below



If we group the first two pulses together they are required to achieve the transformation $\Delta p = 0$ and this leads to a four step cycle in which the pulses go $\emptyset \ 1 \ 2 \ 3$ and the receiver remains fixed as $\emptyset \ \emptyset \ \emptyset \ \emptyset$. In this experiment axial peaks arise due to z -magnetization recovering during the mixing time, and this cycle will *not* suppress these contributions. Thus we need to add axial peak suppression, which is conveniently done by adding the simple

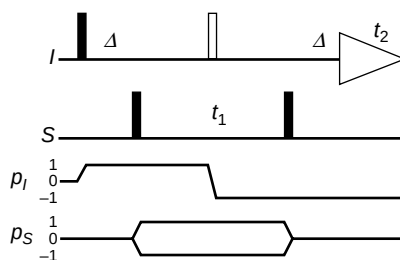
cycle 0 2 on the first pulse and the receiver. The final 8 step cycle is 1st pulse: 0 1 2 3 2 3 0 1, 2nd pulse: 0 1 2 3 0 1 2 3, 3rd pulse fixed, receiver: 0 0 0 0 2 2 2 2.

An alternative is to cycle the last pulse to select the pathway $\Delta p = -1$, giving the cycle 0 1 2 3 for the pulse and 0 1 2 3 for the receiver. Once again, this does not discriminate against z-magnetization which recovers during the mixing time, so a two step phase cycle to select axial peaks needs to be added.

9.4.9.3 Heteronuclear Experiments

The phase cycling for most heteronuclear experiments tends to be rather trivial in that the usual requirement is simply to select that component which has been transferred from one nucleus to another. We have already seen in section 9.1.8 that this is achieved by a 0 2 phase cycle on one of the pulses causing the transfer accompanied by the same on the receiver *i.e.* a difference experiment. The choice of which pulse to cycle depends more on practical considerations than with any fundamental theoretical considerations.

The pulse sequence for HMQC, along with the CTP, is shown below



Note that separate coherence orders are assigned to the *I* and *S* spins. Observable signals on the *I* spin must have $p_I = -1$ and $p_S = 0$ (any other value of p_S would correspond to a heteronuclear multiple quantum coherence). Given this constraint, and the fact that the *I* spin 180° pulse simply inverts the sign of p_I , the only possible pathway on the *I* spins is that shown.

The *S* spin coherence order only changes when pulses are applied to those spins. The first 90° *S* spin pulse generates $p_S = \pm 1$, just as before. As by this point $p_I = +1$, the resulting coherences have $p_S = +1$, $p_I = -1$ (heteronuclear zero-quantum) and $p_S = +1$, $p_I = +1$ (heteronuclear double-quantum). The *I* spin 180° pulse interconverts these midway during t_1 , and finally the last *S* spin pulse returns both pathways to $p_S = 0$. A detailed analysis of the sequence shows that retention of both of these pathways results in amplitude modulation in t_1 (provided that homonuclear couplings between *I* spins are not resolved in the F_1 dimension).

Usually, the *I* spins are protons and the *S* spins some low-abundance heteronucleus, such as ^{13}C . The key thing that we need to achieve is to suppress the signals arising from vast majority of *I* spins which are not coupled to *S* spins. This is achieved by cycling a pulse which affects the phase of the required coherence but which does not affect that of the

unwanted coherence. The obvious targets are the two S spin 90° pulses, each of which is required to give the transformation $\Delta p_s = \pm 1$. A two step cycle with either of these pulses going $\theta \rightarrow 2$ and the receiver doing the same will select this pathway and, by difference, suppress any I spin magnetization which has not been passed into multiple quantum coherence.

It is also common to add EXORCYCLE phase cycling to the I spin 180° pulse, giving a cycle with eight steps overall.

9.4.10 General points about phase cycling

Phase cycling as a method suffers from two major practical problems. The first is that the need to complete the cycle imposes a minimum time on the experiment. In two- and higher-dimensional experiments this minimum time can become excessively long, far longer than would be needed to achieve the desired signal-to-noise ratio. In such cases the only way of reducing the experiment time is to record fewer increments which has the undesirable consequence of reducing the limiting resolution in the indirect dimensions.

The second problem is that phase cycling always relies on recording all possible contributions and then cancelling out the unwanted ones by combining subsequent signals. If the spectrum has high dynamic range, or if spectrometer stability is a problem, this cancellation is less than perfect. The result is unwanted peaks and t_1 -noise appearing in the spectrum. These problems become acute when dealing with proton detected heteronuclear experiments on natural abundance samples, or in trying to record spectra with intense solvent resonances.

Both of these problems are alleviated to a large extent by moving to an alternative method of selection, the use of field gradient pulses, which is the subject of the next section. However, as we shall see, this alternative method is not without its own difficulties.

9.5 Selection with field gradient pulses

9.5.1 Introduction

Like phase cycling, field gradient pulses can be used to select particular coherence transfer pathways. During a pulsed field gradient the applied magnetic field is made spatially inhomogeneous for a short time. As a result, transverse magnetization and other coherences dephase across the sample and are apparently lost. However, this loss can be reversed by the application of a subsequent gradient which undoes the dephasing process and thus restores the magnetization or coherence. The crucial property of the dephasing process is that it proceeds at a different rate for different coherences. For example, double-quantum coherence dephases twice as fast as single-quantum coherence. Thus, by applying gradient pulses of different strengths or durations it is possible to refocus coherences which have, for example, been changed from single- to double-quantum by a radiofrequency pulse.

Gradient pulses are introduced into the pulse sequence in such a way that

only the wanted signals are observed in each experiment. Thus, in contrast to phase cycling, there is no reliance on subtraction of unwanted signals, and it can thus be expected that the level of t_1 -noise will be much reduced. Again in contrast to phase cycling, no repetitions of the experiment are needed, enabling the overall duration of the experiment to be set strictly in accord with the required resolution and signal-to-noise ratio.

The properties of gradient pulses and the way in which they can be used to select coherence transfer pathways have been known since the earliest days of multiple-pulse NMR. However, in the past their wide application has been limited by technical problems which made it difficult to use such pulses in high-resolution NMR. The problem is that switching on the gradient pulse induces currents in any nearby conductors, such as the probe housing and magnet bore tube. These induced currents, called *eddy currents*, themselves generate magnetic fields which perturb the NMR spectrum. Typically, the eddy currents are large enough to disrupt severely the spectrum and can last many hundreds of milliseconds. It is thus impossible to observe a high-resolution spectrum immediately after the application of a gradient pulse. Similar problems have beset NMR imaging experiments and have led to the development of *shielded gradient coils* which do not produce significant magnetic fields outside the sample volume and thus minimise the generation of eddy currents. The use of this technology in high-resolution NMR probes has made it possible to observe spectra within tens of microseconds of applying a gradient pulse. With such apparatus, the use of field gradient pulses in high resolution NMR is quite straightforward, a fact first realised and demonstrated by Hurd whose work has pioneered this whole area.

9.5.2 Dephasing caused by gradients

A field gradient pulse is a period during which the B_0 field is made spatially inhomogeneous; for example an extra coil can be introduced into the sample probe and a current passed through the coil in order to produce a field which varies linearly in the z -direction. We can imagine the sample being divided into thin discs which, as a consequence of the gradient, all experience different magnetic fields and thus have different Larmor frequencies. At the beginning of the gradient pulse the vectors representing transverse magnetization in all these discs are aligned, but after some time each vector has precessed through a different angle because of the variation in Larmor frequency. After sufficient time the vectors are disposed in such a way that the net magnetization of the sample (obtained by adding together all the vectors) is zero. The gradient pulse is said to have dephased the magnetization.

It is most convenient to view this dephasing process as being due to the generation by the gradient pulse of a *spatially dependent phase*. Suppose that the magnetic field produced by the gradient pulse, B_g , varies linearly along the z -axis according to

$$B_g = Gz$$

where G is the gradient strength expressed in, for example, T m^{-1} or G cm^{-1} ;

the origin of the z-axis is taken to be in the centre of the sample. At any particular position in the sample the Larmor frequency, $\omega_L(z)$, depends on the applied magnetic field, B_0 , and B_g

$$\omega_L = \gamma(B_0 + B_g) = \gamma(B_0 + Gz) ,$$

where γ is the gyromagnetic ratio. After the gradient has been applied for time t , the phase at any position in the sample, $\Phi(z)$, is given by $\Phi(z) = \gamma(B_0 + Gz)t$. The first part of this phase is just that due to the usual Larmor precession in the absence of a field gradient. Since this is constant across the sample it will be ignored from now on (which is formally the same result as viewing the magnetization in a frame of reference rotating at γB_0). The remaining term γGzt is the *spatially dependent phase* induced by the gradient pulse.

If a gradient pulse is applied to pure x-magnetization, the following evolution takes place at a particular position in the sample

$$I_x \xrightarrow{\gamma Gzt I_z} \cos(\gamma Gzt) I_x + \sin(\gamma Gzt) I_y .$$

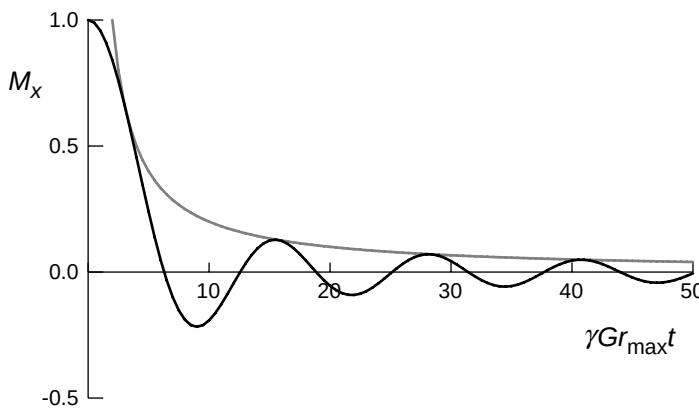
The total x-magnetization in the sample, M_x , is found by adding up the magnetization from each of the thin discs, which is equivalent to the integral

$$M_x(t) = \frac{1}{r_{\max}} \int_{-\frac{1}{2}r_{\max}}^{\frac{1}{2}r_{\max}} \cos(\gamma Gzt) dz$$

where it has been assumed that the sample extends over a region $\pm \frac{1}{2}r_{\max}$. Evaluating the integral gives an expression for the decay of x-magnetization during a gradient pulse

$$M_x(t) = \frac{\sin(\frac{1}{2} \gamma G r_{\max} t)}{\frac{1}{2} \gamma G r_{\max} t} \quad [11]$$

The plot below shows $M_x(t)$ as a function of time



The black line shows the decay of magnetization due to the action of a gradient pulse. The grey line is an approximation, valid at long times, for the envelope of the decay.

Note that the oscillations in the decaying magnetization are imposed on an overall decay which, for long times, is given by $2/(\gamma G r_{\max})$. Equation [11] embodies the obvious points that the stronger the gradient (the larger G) the faster the magnetization decays and that magnetization from nuclei with higher gyromagnetic ratios decays faster. It also allows a quantitative

assessment of the gradient strengths required: the magnetization will have decayed to a fraction α of its initial value after a time of the order of $2/(\gamma G \alpha r_{\max})$ (the relation is strictly valid for $\alpha \ll 1$). For example, if it is assumed that $r_{\max}$ is 1 cm, then a 2 ms gradient pulse of strength 0.37 T m^{-1} (37 G cm^{-1}) will reduce proton magnetization by a factor of 1000. Gradients of such strength are readily obtainable using modern shielded gradient coils that can be built into high resolution NMR probes

This discussion now needs to be generalised for the case of a field gradient pulse whose amplitude is not constant in time, and for the case of dephasing a general coherence of order p . The former modification is of importance as for instrumental reasons the amplitude envelope of the gradient is often shaped to a smooth function. In general after applying a gradient pulse of duration τ the spatially dependent phase, $\Phi(r, \tau)$ is given by

$$\Phi(r, \tau) = sp\gamma B_g(r)\tau \quad [12]$$

The proportionality to the coherence order comes about due to the fact that the phase acquired as a result of a z-rotation of a coherence of order p through an angle ϕ is $p\phi$, (see Eqn. [2] in section 9.2.1). In Eqn. [12] s is a shape factor: if the envelope of the gradient pulse is defined by the function $A(t)$, where $|A(t)| \leq 1$, s is defined as the area under $A(t)$

$$s = \frac{1}{\tau} \int_0^{\tau} A(t) dt$$

The shape factor takes a particular value for a certain shape of gradient, regardless of its duration. A gradient applied in the opposite sense, that is with the magnetic field decreasing as the z-coordinate increases rather than *vice versa*, is described by reversing the sign of s . The overall amplitude of the gradient is encoded within B_g .

In the case that the coherence involves more than one nuclear species, Eqn. [12] is modified to take account of the different gyromagnetic ratio for each spin, γ_i , and the (possibly) different order of coherence with respect to each nuclear species, p_i :

$$\Phi(r, \tau) = sB_g(r)\tau \sum_i p_i \gamma_i$$

From now on we take the dependence of Φ on r and τ , and of B_g on r as being implicit.

9.5.3 Selection by refocusing

The method by which a particular coherence transfer pathway is selected using field gradients is illustrated opposite. The first gradient pulse encodes a spatially dependent phase, Φ_1 and the second a phase Φ_2 where

$$\Phi_1 = s_1 p_1 \gamma B_{g,1} \tau_1 \quad \text{and} \quad \Phi_2 = s_2 p_2 \gamma B_{g,2} \tau_2 .$$

After the second gradient the net phase is $(\Phi_1 + \Phi_2)$. To select the pathway involving transfer from coherence order p_1 to coherence order p_2 , this net phase should be zero; in other words the dephasing induced by the first gradient pulse is undone by the second. The condition $(\Phi_1 + \Phi_2) = 0$ can be rearranged to

$$\frac{s_1 B_{g,1} \tau_1}{s_2 B_{g,2} \tau_2} = \frac{-p_2}{p_1} . \quad [13]$$

For example, if $p_1 = +2$ and $p_2 = -1$, refocusing can be achieved by making the second gradient either twice as long ($\tau_2 = 2\tau_1$), or twice as strong ($B_{g,2} = 2B_{g,1}$) as the first; this assumes that the two gradients have identical shape factors. Other pathways remain dephased; for example, assuming that we have chosen to make the second gradient twice as strong and the same duration as the first, a pathway with $p_1 = +3$ to $p_2 = -1$ experiences a net phase

$$\Phi_1 + \Phi_2 = 3sB_{g,1}\tau_1 - sB_{g,2}\tau_1 = sB_{g,1}\tau_1 .$$

Provided that this spatially dependent phase is sufficiently large, according to the criteria set out in the previous section, the coherence arising from this pathway remains dephased and is not observed. To refocus a pathway in which there is no sign change in the coherence orders, for example, $p_1 = -2$ to $p_2 = -1$, the second gradient needs to be applied in the opposite sense to the first; in terms of Eqn. [13] this is expressed by having $s_2 = -s_1$.

The procedure can easily be extended to select a more complex coherence transfer pathway by applying further gradient pulses as the coherence is transferred by further pulses, as illustrated opposite. The condition for refocusing is again that the net phase acquired by the required pathway be zero, which can be written formally as

$$\sum_i s_i p_i \gamma_i B_{g,i} \tau_i = 0 .$$

With more than two gradients in the sequence, there are many ways in which a given pathway can be selected. For example, the second gradient may be used to refocus the first part of the required pathway, leaving the third and fourth to refocus another part. Alternatively, the pathway may be consistently dephased and the magnetization only refocused by the final gradient, just before acquisition.

At this point it is useful to contrast the selection achieved using gradient pulses with that achieved using phase cycling. From Eqn. [13] it is clear that a particular pair of gradient pulses selects a particular *ratio* of coherence orders; in the above example any two coherence orders in the ratio $-2 : 1$ or $2 : -1$ will be refocused. This selection according to ratio of

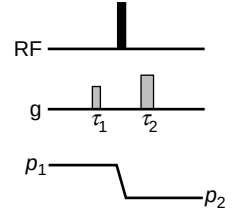
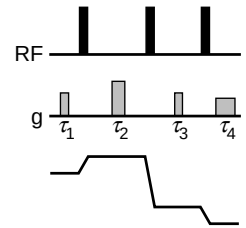


Illustration of the use of a pair of gradients to select a single pathway. The radiofrequency pulses are on the line marked "RF" and the field gradient pulses are denoted by shaded rectangles on the line marked "g".

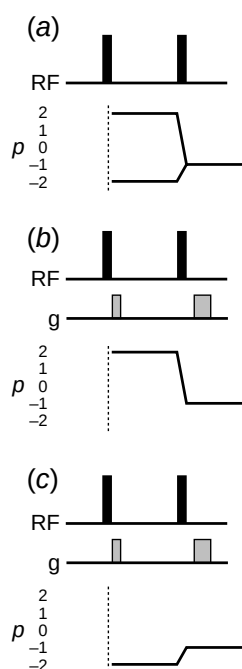


coherence orders is in contrast to the case of phase cycling in which a phase cycle consisting of N steps of $2\pi/N$ radians selects a particular *change* in coherence order $\Delta p = p_2 - p_1$, and further pathways which have $\Delta p = (p_2 - p_1) \pm mN$, where $m = 0, 1, 2 \dots$

It is straightforward to devise a series of gradient pulses which will *select* a single coherence transfer pathway. It cannot be assumed, however, that such a sequence of gradient pulses will *reject* all other pathways *i.e.* leave coherence from all other pathways dephased at the end of the sequence. Such assurance can only be given by analysing the fate of *all* other possible coherence transfer pathways under the particular gradient sequence proposed. In complex pulse sequences there may also be several different ways in which gradient pulses can be included in order to achieve selection of the desired pathway. Assessing which of these alternatives is the best, in the light of the requirement of suppression of unwanted pathways and the effects of pulse imperfections may be a complex task.

9.5.3.1 Selection of multiple pathways

As we have seen earlier, it is not unusual to want to select two or more pathways *simultaneously*, for example either to maximise the signal intensity or to retain absorption-mode lineshapes. A good example of this is the double-quantum filter pulse sequence element, shown opposite.



The ideal pathway, shown in (a), preserves coherence orders $p = \pm 2$ during the inter-pulse delay. Gradients can be used to select the pathway -2 to -1 or $+2$ to -1 , shown in (b) and (c) respectively. However, no combination of gradients can be found which will select simultaneously both of these pathways. In contrast, it is easy to devise a phase cycle which selects both of these pathways (section 9.4.5.4). Thus, selection with gradients will in this case result in a loss of *half* of the available signal when compared to an experiment of equal length which uses selection by phase cycling. Such a loss in signal is, unfortunately, a very common feature when gradients are used for pathway selection.

9.5.3.2 Selection versus suppression

Coherence order zero, comprising z -magnetization, zz -terms and homonuclear zero-quantum coherence, does not accrue any phase during a gradient pulse. Thus, it can be separated from all other orders simply by applying a single gradient. In a sense, however, this is not a gradient selection process; rather it is a *suppression* of all other coherences. A gradient used in this way is often called a *purge gradient*. In contrast to experiments where selection is achieved, there is no inherent sensitivity loss when gradients are used for suppression. We will see examples below of cases where this suppression approach is very useful.

9.5.3.3 Gradients on other axes

The simplest experimental arrangement generates a gradient in which the magnetic field varies in the z direction, however it is also possible to generate gradients in which the field varies along x or y . Clearly, the spatially dependent phase generated by a gradient applied in one direction

cannot be refocused by a gradient applied in a different direction. In sequences where more than one pair of gradients are used, it may be convenient to apply further gradients in different directions to the first pair, so as to avoid the possibility of accidentally refocusing unwanted coherence transfer pathways. Likewise, a gradient which is used to destroy all coherences can be applied in a different direction to gradients used for pathway selection.

9.5.4 Refocusing and inversion pulses

Refocusing and inversion pulses play an important role in multiple-pulse NMR experiments and so the interaction between such pulses and field gradient pulses will be explored in some detail. As has been noted above in section 9.4.3, a perfect refocusing pulse simply changes the sign of the order of any coherences present, $p \rightarrow -p$. If the pulse is imperfect, there will be transfer to coherence orders other than $-p$.

A perfect inversion pulse simply inverts z-magnetization or, more generally, all z-operators: $I_z \rightarrow -I_z$. If the pulse is imperfect, it will generate transverse magnetization or other coherences. Inversion pulses are used extensively in heteronuclear experiments to control the evolution of heteronuclear couplings.

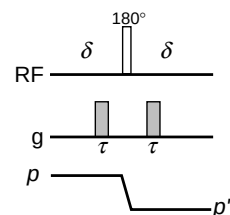
We start out the discussion by considering the refocusing of coherences, illustrated opposite. The net phase, Φ , at the end of such a sequence is

$$\Phi = \Omega^{(p)}\delta + sp\gamma B_g\tau + \Omega^{(p')}\delta + sp'\gamma B_g\tau$$

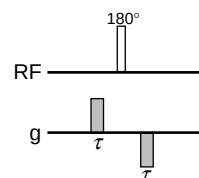
where $\Omega^{(p)}$ is the frequency with which coherence of order p evolves in the absence of a gradient; note that $\Omega^{(-p)} = -\Omega^{(p)}$. The net phase is zero if, and only if, $p' = -p$. With sufficiently strong gradients all other pathways remain dephased and the gradient sequence thus selects the refocused component. As is expected for a spin echo, the underlying evolution of the coherence (as would occur in the absence of a gradient) is also refocused by the selection of the pathway shown. Any transverse magnetization which an imperfect refocusing pulse might create is also dephased.

Placing equal gradients either side of a refocusing pulse therefore selects the coherence transfer pathway associated with a perfect refocusing pulse. This selection works for all coherence orders so, in contrast to the discussion in section 9.5.3.1, there is no loss of signal. Such a pair of gradients are often described as being used to "clean up" a refocusing pulse, referring to their role in eliminating unwanted pathways.

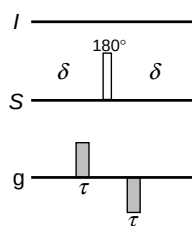
We cannot use gradients to select the pathway associated with an inversion pulse as $p = 0$ both before and after the pulse. However, we can apply a gradient *after* the pulse to dephase any magnetization which might be created by an imperfect pulse. Taking the process a step further, we can apply a gradient both *before* and *after* the pulse, with the two gradients in *opposite* directions. The argument here is that this results in the maximum dephasing of unwanted coherences – both those present before the pulse and those that might be generated by the pulse. Again, this sequence is often described as being used to "clean up" an inversion pulse.



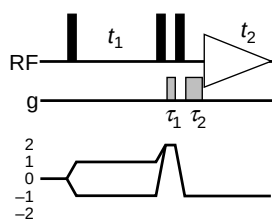
Gradient sequence used to "clean up" a refocusing pulse. Note that the two gradients are of equal area. The refocused pathway has $p' = -p$.



Gradient sequence used to "clean up" an inversion pulse.



In a heteronuclear experiment (separate pulses shown on the *I* and *S* spins) a 180° pulse on the *S* spin refocuses the *I*/*S* coupling over the period 2δ. The gradient pulses shown are used to "clean up" the inversion pulse.



Simple pulse sequence for DQF COSY with gradient selection and retention of symmetrical pathways in t_1 . The second gradient is twice the area of the first.

In heteronuclear experiments an inversion pulse applied to one nucleus is used to refocus the evolution of a coupling to another nucleus. For example, in the sequence shown opposite the centrally placed *S* spin 180° pulse refocuses the *I*/*S* coupling over the period 2δ. The pair of gradients shown have no net effect on *I* spin coherences as the dephasing due to the first gradient is exactly reversed by the second. The gradient sequence can be thought of as "cleaning up" the *S* spin inversion pulse.

9.5.4.1 Phase errors due to gradient pulses

For the desired pathway, the spatially dependent phase created by a gradient pulse is refocused by a subsequent gradient pulse. However, the underlying evolution of offsets (chemical shifts) and couplings is not refocused, and phase errors will accumulate due to the evolution of these terms. Since gradient pulses are typically of a few milliseconds duration, these phase errors are substantial.

In multi-dimensional NMR the uncompensated evolution of offsets during gradient pulses has disastrous effects on the spectra. This is illustrated here for the DQF COSY experiment, for which a pulse sequence using the gradient pulses is shown opposite. The problem with this sequence is that the double quantum coherence will evolve during delay τ_1 . Normally, the delay between the last two pulses would be a few microseconds at most, during which the evolution is negligible. However, once we need to apply a gradient the delay will need to be of the order of milliseconds and significant evolution takes place. The same considerations apply to the second gradient.

We can investigate the effect of this evolution by analysing the result for a two-spin system. In the calculation, it will be assumed that only the indicated pathway survives and that the spatially dependent part of the evolution due to the gradients can be ignored as ultimately it is refocused. The coherence with order of + 2 present during τ_1 evolves according to

$$I_{1+}I_{2+} \xrightarrow{\Omega_1\tau_1 I_{1z} + \Omega_2\tau_1 I_{2z}} I_{1+}I_{2+} \exp(-i(\Omega_1 + \Omega_2)\tau_1) ,$$

where Ω_1 and Ω_2 are the offsets of spins 1 and 2, respectively. After the final 90° pulse and the second gradient the observable terms on spin 1 are

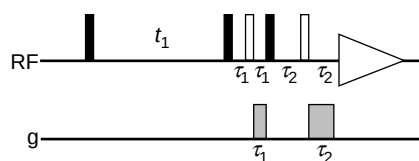
$$\frac{1}{2} \exp(-i(\Omega_1 + \Omega_2)\tau_1) \left[\cos \Omega_1 \tau_2 \ 2I_{1x}I_{2z} + \sin \Omega_1 \tau_2 \ 2I_{1y}I_{2z} \right] \quad [14]$$

where it has been assumed that τ_2 is sufficiently short that evolution of the coupling can be ignored. It is clearly seen from Eqn. [14] that, due to the evolution during τ_2 , the multiplet observed in the F_2 dimension will be a mixture of dispersion and absorption anti-phase contributions. In addition, there is an overall phase shift due to the evolution during τ_1 . The phase correction needed to restore this multiplet to absorption depends on both the frequency in F_2 and the double-quantum frequency during the first gradient. Thus, no single linear frequency dependent phase correction could phase correct a spectrum containing many multiplets. The need to control these phase errors is plain.

The general way to minimise these problems is to associate each gradient with a refocusing pulse, as shown opposite. In sequence (a) the gradient is placed in one of the delays of a spin echo; the evolution of the offset during the first delay τ is refocused during the second delay τ . So, overall there are no phase errors due to the evolution of the offset.

An alternative is the sequence (b). Here, as in (a), the offset is refocused over the whole sequence. The first gradient results in the usually spatially dependent phase and then the 180° pulse changes the sign of the coherence order. As the second gradient is opposite to the first, it causes *further* dephasing; effectively, it is as if a gradient of length 2τ is applied. Sequence (b) will give the same dephasing effect as (a) if each gradient in (b) is of duration $\tau/2$; the overall sequence will then be of duration τ . If losses due to relaxation are a problem, then clearly sequence (b) is to be preferred as it takes less time than (a).

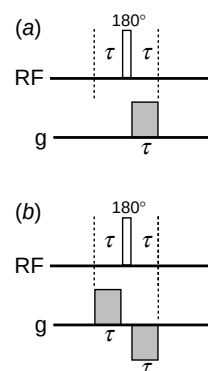
The sequence below shows the gradient-selected DQF COSY pulse sequence modified by the inclusion of extra 180° pulses to remove phase errors. Note that although the extra 180° pulses are effective at refocusing offsets, they do not refocus the evolution of homonuclear couplings. It is essential, therefore, to keep the gradient pulses as short as is feasible.



In many pulse sequences there are already periods during which the evolution of offsets is refocused. The evolution of offsets during a gradient pulse placed within such a period will therefore also be refocused, making it unnecessary to include extra refocusing pulses. Likewise, a gradient may be placed during a "constant time" evolution period of a multi-dimensional pulse sequence without introducing phase errors in the corresponding dimension; the gradient simply becomes part of the constant time period. This approach is especially useful in three- and four-dimensional experiments used to record spectra of ^{15}N , ^{13}C labelled proteins.

9.5.5 Sensitivity

The use of gradients for coherence selection has consequences for the signal-to-noise ratio of the spectrum when it is compared to a similar spectrum recorded using phase cycling. If a gradient is used to *suppress* all coherences other than $p = 0$, *i.e.* it is used simply to remove all coherences, leaving just z -magnetization or zz terms, there is no inherent loss of sensitivity when compared to a corresponding phase cycled experiment. If, however, the gradient is used to *select* a particular order of coherence the signal which is subsequently refocused will almost always be half the intensity of that which can be observed in a phase cycled experiment. This factor comes about simply because it is likely that the phase cycled experiment will be able to retain two symmetrical pathways, whereas the gradient selection method will only be able to refocus one of these.



The foregoing discussion applies to the case of a selection gradient placed in a *fixed* delay of a pulse sequence. The matter is different if the gradient is placed within the incrementable time of a multi-dimensional experiment, *e.g.* in t_1 of a two-dimensional experiment. To understand the effect that such a gradient has on the sensitivity of the experiment it is necessary to be rather careful in making the comparison between the gradient selected and phase cycled experiments. In the case of the latter experiments we need to include the SHR or TPPI method in order to achieve frequency discrimination with absorption mode lineshapes. If a gradient is used in t_1 we will need to record separate P- and N-type spectra so that they can be recombined to give an absorption mode spectrum. We must also ensure that the two spectra we are comparing have the same limiting resolution in the t_1 dimension, that is they achieve the same maximum value of t_1 and, of course, the total experiment time must be the same.

The detailed argument which is needed to analyse this problem is beyond the scope of this lecture; it is given in detail in *J. Magn. Reson. Ser. A*, **111**, 70-76 (1994)[†]. The conclusion is that the signal-to-noise ratio of an absorption mode spectrum generated by recombining P- and N-type gradient selected spectra is lower, by a factor $1/\sqrt{2}$, than the corresponding phase cycled spectrum with SHR or TPPI data processing.

The potential reduction in sensitivity which results from selection with gradients may be more than compensated for by an improvement in the *quality* of the spectra obtained in this way. Often, the factor which limits whether or not a cross peak can be seen is not the *thermal* noise level by the presence of other kinds of "noise" associated with imperfect cancellation *etc.*

9.5.6 Diffusion

The process of refocusing a coherence which has been dephased by a gradient pulse is inhibited if the spins move either during or between the defocusing and refocusing gradients. Such movement alters the magnetic field experienced by the spins so that the phase acquired during the refocusing gradient is not exactly opposite to that acquired during the defocusing gradient.

In liquids there is a translational diffusion of both solute and solvent which causes such movement at a rate which is fast enough to cause significant effects on NMR experiments using gradient pulses. As diffusion is a random process we expect to see a smooth attenuation of the intensity of the refocused signal as the diffusion contribution increases. These effects have been known and exploited to measure diffusion constants since the very earliest days of NMR.

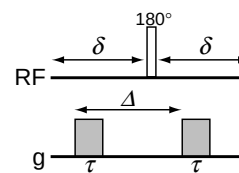
[†] There is an error in this paper: in Fig. 1(b) the penultimate S spin 90° pulse should be phase y and the final S spin 90° pulse is not required.

The effect of diffusion on the signal intensity from the simple echo sequence shown opposite is relatively simple to analyse and captures all of the essential points. Note that the two gradient pulses can be placed anywhere in the intervals δ either side of the 180° pulse. For a single uncoupled resonance, the intensity of the observed signal, S , expressed as a fraction of the signal intensity in the absence of a gradient, S_0 is given by

$$\frac{S}{S_0} = \exp\left(-\gamma^2 G^2 \tau^2 \left(\Delta - \frac{\tau}{3}\right) D\right) \quad [15]$$

where D is the diffusion constant, Δ is the time between the start of the two gradient pulses and τ is the duration of the gradient pulses; relaxation has been ignored. For a given pair of gradient pulses it is diffusion during the interval between the two pulses, Δ , which determines the attenuation of the echo. The stronger the gradient the more rapidly the phase varies across the sample and thus the more rapidly the echo will be attenuated. This is the physical interpretation of the term $\gamma^2 G^2 \tau^2$ in Eqn. [15].

Diffusion constants generally decrease as the molecular mass increases. A small molecule, such as water, will diffuse up to twenty times faster than a protein with molecular weight 20,000. The table shows the loss in intensity due to diffusion for typical gradient pulse pair of 2 ms duration and of strength 10 G cm^{-1} for a small, medium and large sized molecule; data is given for $\Delta = 2 \text{ ms}$ and $\Delta = 100 \text{ ms}$. It is seen that even for the most rapidly diffusing molecules the loss of intensity is rather small for $\Delta = 2 \text{ ms}$, but becomes significant for longer delays. For large molecules, the effect is small in all cases.



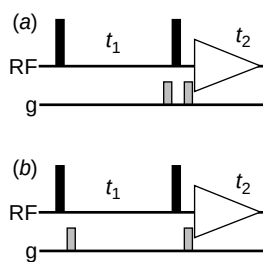
Fraction of transverse magnetization refocused after a spin echo with gradient refocusing^a

Δ/ms	small molecule ^b	medium sized molecule ^c	macro molecule ^d
2	0.99	1.00	1.00
100	0.55	0.88	0.97

^a Calculated for the pulse sequence shown above for two gradients of strength 10 G cm^{-1} and duration, τ , 2 ms; relaxation is ignored. ^b Diffusion constant, D , taken as that for water, which is $2.1 \times 10^{-9} \text{ m}^2 \text{ s}^{-1}$ at ambient temperatures. ^c Diffusion constant taken as $0.46 \times 10^{-9} \text{ m}^2 \text{ s}^{-1}$. ^d Diffusion constant taken as $0.12 \times 10^{-9} \text{ m}^2 \text{ s}^{-1}$.

9.5.6.1 Minimisation of Diffusion Losses

The foregoing discussion makes it clear that in order to minimise intensity losses due to diffusion the product of the strength and durations of the gradient pulses, $G^2 \tau^2$, should be kept as small as is consistent with achieving the required level of suppression. In addition, a gradient pulse pair should be separated by the shortest time, Δ , within the limits imposed by the pulse sequence. This condition applies to gradient pairs the first of which is responsible for dephasing, and the second for rephasing. Once the coherence is rephased the time that elapses before further gradient pairs is irrelevant from the point of view of diffusion losses.



In two-dimensional NMR, diffusion can lead to line broadening in the F_1 dimension if t_1 intervenes between a gradient pair. Consider the two alternative pulse sequences for recording a simple COSY spectrum shown opposite. In (a) the gradient pair are separated by the very short time of the final pulse, thus keeping the diffusion induced losses to an absolute minimum. In (b) the two gradients are separated by the incrementable time t_1 ; as this increases the losses due to diffusion will also increase, resulting in an extra decay of the signal in t_1 . The extra line broadening due to this decay can be estimated from Eqn. [15], with $\Delta = t_1$, as $\gamma^2 G^2 \tau^2 D / \pi$ Hz. For a pair of 2 ms gradients of strength 10 G cm⁻¹ this amounts ≈ 2 Hz in the case of a small molecule.

This effect by which diffusion causes an extra line broadening in the F_1 dimension is usually described as *diffusion weighting*. Generally it is possible to avoid it by careful placing of the gradients. For example, the sequences (a) and (b) are in every other respect equivalent, thus there is no reason *not* to chose (a). It should be emphasised that diffusion weighting occurs only when t_1 intervenes between the dephasing and refocusing gradients.

9.5.7 Some examples of gradient selection

9.5.7.1 Introduction

Reference has already been made to the two general advantages of using gradient pulses for coherence selection, namely the possibility of a general improvement in the quality of spectra and the removal of the requirement of completing a phase cycle for each increment of a multi-dimensional experiment. This latter point is particularly significant when dealing with three- and four-dimensional experiments.

The use of gradients results in very significant improvement in the quality of proton-detected heteronuclear experiments, especially when unlabelled samples are used. In such experiments, gradient selection results in much lower dynamic range in the free induction decay as compared to phase cycled experiments.

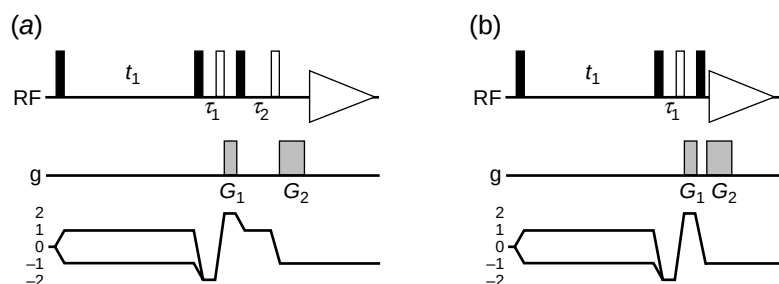
As has been discussed above, special care needs to be taken in experiments which use gradient selection in order to retain absorption mode lineshapes.

In the following sections the use of gradient selection in several different experiments will be described. The gradient pulses used in these sequences will be denoted G_1 , G_2 etc. where G_i implies a gradient of duration τ_i , strength $B_{g,i}$ and shape factor s_i . There is always the choice of altering the duration, strength or, conceivably, shape factor in order to establish refocusing. Thus, for brevity we shall from now on write the spatially dependent phase produced by gradient G_i acting on coherence of order p as $\gamma p G_i$ in the homonuclear case or

$$\sum_j \gamma_j p_j G_i$$

in the heteronuclear case; the sum is over all types of nucleus.

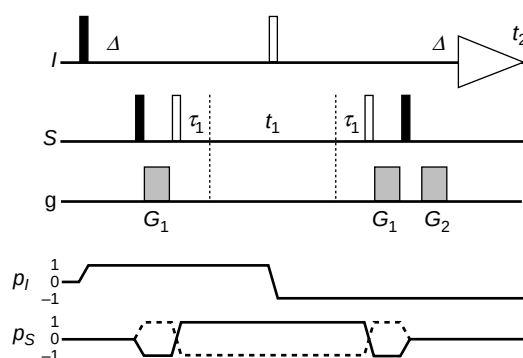
9.5.7.2 Double-quantum Filtered COSY



This experiment has already been discussed in detail in previous sections; sequence (a) is essentially that described already and is suitable for recording absorption mode spectra. The refocusing condition is $G_2 = 2G_1$; frequency discrimination in the F_1 dimension is achieved by the SHR or TPPI procedures. Multiple quantum filters through higher orders can be implemented in the same manner.

In sequence (b) the final spin echo is not required as data acquisition is started *immediately* after the final radiofrequency pulse; phase errors which would accumulate during the second gradient pulse are thus avoided. Of course, the signal only rephases towards the end of the final gradient, so there is little signal to be observed. However, the crucial point is that, as the magnetization is all in antiphase at the start of t_2 , the signal grows from zero at a rate determined by the size of the couplings on the spectrum. Provided that the gradient pulse is much shorter than $1/J$, where J is a typical proton-proton coupling constant, the part of the signal missed during the gradient pulse is not significant and the spectrum is not perturbed greatly. An alternative procedure is to start to acquire the data *after* the final gradient, and then to right shift the free induction decay, bringing in zeroes from the left, by a time equal to the duration of the gradient.

9.5.7.3 HMQC



There are several ways of implementing gradient selection into the HMQC experiment, one of which, which leads to absorption mode spectra, is shown above. The centrally placed I spin 180° pulse results in no net dephasing of the I spin part of the heteronuclear multiple quantum coherence by the two gradients G_1 i.e. the dephasing of the I spin coherence caused by the first is undone by the second. However, the S spin coherence experiences a net dephasing due to these two gradients and this coherence is subsequently

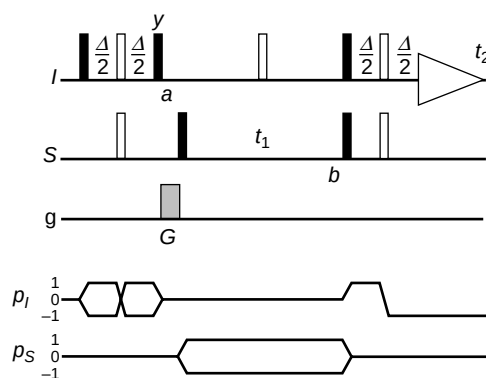
refocused by G_2 . Two 180° S spin pulses together with the delays τ_1 refocus shift evolution during the two gradients G_1 . The centrally placed 180° I spin pulse refocuses chemical shift evolution of the I spins during the delays Δ and all of the gradient pulses (the last gradient is contained within the final delay, Δ). The refocusing condition is

$$\mp 2\gamma_s G_1 - \gamma_I G_2 = 0$$

where the $+$ and $-$ signs refer to the P- and N-type spectra respectively. The switch between recording these two types of spectra is made simply by reversing the sense of G_2 . The P- and N-type spectra are recorded separately and then combined in the manner described in section 9.3.4.2 to give a frequency discriminated absorption mode spectrum.

In the case that I and S are proton and carbon-13 respectively, the gradients G_1 and G_2 are in the ratio $2:\pm 1$. Proton magnetization not involved in heteronuclear multiple quantum coherence, *i.e.* magnetization from protons not coupled to carbon-13, is refocused after the second gradient G_1 but is then dephased by the final gradient G_2 . Provided that the gradient is strong enough these unwanted signals, and the t_1 -noise associated with them, will be suppressed.

9.5.7.4 HSQC

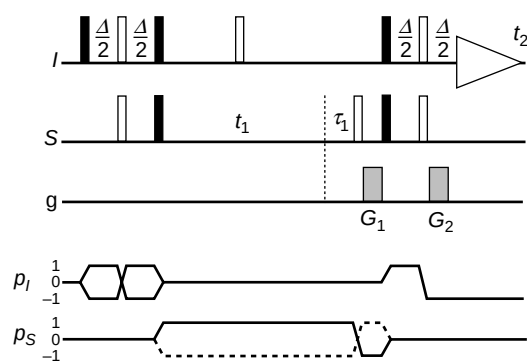


The sequence above shows the simplest way of implementing gradients into the HSQC experiment. An analysis using product operators shows that at point a the required signal is present as the operator $2I_z S_z$ whereas the undesired signal (from I spins not coupled to S spins) is present as I_y . Thus, a field gradient applied at point a will dephase the unwanted magnetization and leave the wanted term unaffected. This is an example of using gradients not for selection, but for suppression of unwanted magnetization (see section 9.5.3.2).

The main practical difficulty with this approach is that the unwanted magnetization is only along y at point a provided all of the pulses are perfect; if the pulses are imperfect there will be some z -magnetization present which will not be eliminated by the gradient. In the case of observing proton- ^{13}C or proton- ^{15}N HSQC spectra from natural abundance samples, the magnetization from uncoupled protons is very much larger than the wanted magnetization, so even very small imperfections in the

pulses can give rise to unacceptably large residual signals. However, for globally labelled samples the degree of suppression is often sufficient and such an approach is used successfully in many three- and four-dimensional experiments applied to globally ^{13}C and ^{15}N labelled proteins.

The key to obtaining the best suppression of the uncoupled magnetization is to apply a gradient when transverse magnetization is present on the S spin. An example of the HSQC experiment utilising such a principle is shown below



HSQC pulse sequence with gradient selection. The CTP for an N-type spectrum is shown by the full line and for the P-type spectrum by the dashed line.

Here, G_1 dephases the S spin magnetization present at the end of t_1 and, after transfer to the I spins, G_2 refocuses the signal. An extra 180° pulse to S in conjunction with the extra delay τ_1 ensures that phase errors which accumulate during G_1 are refocused; G_2 is contained within the final spin echo which is part of the usual HSQC sequence. The refocusing condition is

$$\mp \gamma_S G_1 - \gamma_I G_2 = 0$$

where the $-$ and $+$ signs refer to the N- and P-type spectra respectively. As before, an absorption mode spectrum is obtained by combining the N- and P-type spectra, which can be selected simply by reversing the sense of G_2 .

9.6 Zero-quantum dephasing and purge pulses

Both z-magnetization and homonuclear zero-quantum coherence have coherence order 0, and thus neither are dephased by the application of a gradient pulse. Selection of coherence order zero is achieved simply by applying a gradient pulse which is long enough to dephase all other coherences; no refocusing is used. In the vast majority of experiments it is the z-magnetization which is required and the zero-quantum coherence that is selected at the same time is something of a nuisance.

A number of methods have been developed to suppress contributions to the spectrum from zero-quantum coherence. Most of these utilise the property that zero-quantum coherence evolves in time, whereas z-magnetization does not. Thus, if several experiments in which the zero-quantum has been allowed to evolve for different times are co-added, cancellation of zero-quantum contributions to the spectrum will occur. Like phase cycling, such a method is time consuming and relies on a difference

procedure. However, it has been shown that if a field gradient is combined with a period of spin-locking the coherences which give rise to these zero-quantum coherences can be dephased. Such a process is conveniently considered as a modified purging pulse.

9.6.1 Purging pulses

A purging pulse consists of a relatively long period of spin-locking, taken here to be applied along the x -axis. Magnetization not aligned along x will precess about the spin-locking field and, because this field is inevitably inhomogeneous, such magnetization will dephase. The effect is thus to purge all magnetization except that aligned along x . However, in a coupled spin system certain anti-phase states aligned *perpendicular* to the spin-lock axis are also preserved. For a two spin system (with spins k and l), the operators preserved under spin-locking are I_{kx} , I_{lx} and the anti-phase state $2I_{ky}I_{lz} - 2I_{kz}I_{ly}$. Thus, in a coupled spin system, the purging effect of the spin-locking pulse is less than perfect.

The reason why these anti-phase terms are preserved can best be seen by transforming to a tilted co-ordinate system whose z -axis is aligned with the *effective* field seen by each spin. For the case of a strong B_1 field placed close to resonance the effective field seen by each spin is along x , and so the operators are transformed to the tilted frame simply by rotating them by -90° about y

$$\begin{aligned} I_{kx} &\xrightarrow{-\pi/2 I_{ky}} I_{kz}^T & I_{lx} &\xrightarrow{-\pi/2 I_{ly}} I_{lz}^T \\ 2I_{ky}I_{lz} - 2I_{kz}I_{ly} &\xrightarrow{-\pi/2(I_{ky}+I_{ly})} 2I_{ky}^T I_{lx}^T - 2I_{kx}^T I_{ly}^T \end{aligned}$$

Operators in the tilted frame are denoted with a superscript T. In this frame the x -magnetization has become z , and as this is parallel with the effective field, it clearly does not dephase. The anti-phase magnetization along y has become

$$2I_{ky}^T I_{lx}^T - 2I_{kx}^T I_{ly}^T$$

which is recognised as *zero-quantum coherence in the tilted frame*. Like zero-quantum coherence in the normal frame, this coherence does not dephase in a strong spin-locking field. There is thus a connection between the inability of a field gradient to dephase zero-quantum coherence and the preservation of certain anti-phase terms during a purging pulse.

Zero-quantum coherence in the tilted frame evolves with time at a frequency, Ω_{zQ}^T , given by

$$\Omega_{zQ}^T = \left| \sqrt{(\Omega_k^2 + \omega_1^2)} - \sqrt{(\Omega_l^2 + \omega_1^2)} \right|$$

where Ω_i is the offset from the transmitter of spin i and ω_1 is the B_1 field strength. If a field gradient is applied during the spin-locking period the zero quantum frequency is modified to

$$\Omega_{zQ}^T(r) = \left| \sqrt{\left\{ \Omega_k + \gamma B_g(r) \right\}^2 + \omega_1^2} - \sqrt{\left\{ \Omega_l + \gamma B_g(r) \right\}^2 + \omega_1^2} \right|$$

This frequency can, under certain circumstances, become spatially

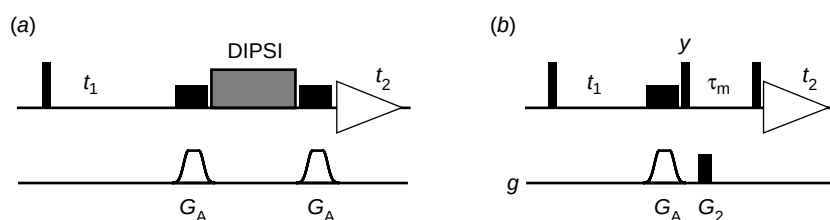
dependent and thus the zero-quantum coherence in the tilted frame will dephase. This is in contrast to the case of zero-quantum coherence in the laboratory frame which is not dephased by a gradient pulse.

The principles of this dephasing procedure are discussed in detail elsewhere (*J. Magn. Reson. Ser. A* **105**, 167-183 (1993)). Here, we note the following features. (a) The optimum dephasing is obtained when the extra offset induced by the gradient at the edges of the sample, $\gamma B_g(r_{\max})$, is of the order of ω_1 . (b) The rate of dephasing is proportional to the zero-quantum frequency in the absence of a gradient, $(\Omega_k - \Omega_l)$. (c) The gradient must be switched on and off adiabatically. (d) The zero-quantum coherences may also be dephased using the inherent inhomogeneity of the radio-frequency field produced by typical NMR probes, but in such a case the optimum dephasing rate is obtained by spin locking off-resonance so that

$$\tan^{-1}(\omega_1/\Omega_{k,l}) \approx 54^\circ.$$

(e) Dephasing in an inhomogeneous B_1 field can be accelerated by the use of special composite pulse sequences.

The combination of spin-locking with a gradient pulse allows the implementation of essentially perfect purging pulses. Such a pulse could be used in a two-dimensional TOCSY experiment whose pulse sequence is shown below as (a).



Pulse sequences using purging pulses which comprise a period of spin locking with a magnetic field gradient. The field gradient must be switched on and off in an adiabatic manner.

In this experiment, the period of isotropic mixing transfers in-phase magnetization (say along x) between coupled spins, giving rise to cross-peaks which are absorptive and in-phase in both dimensions. However, the mixing sequence also both transfers and generates anti-phase magnetization along y , which gives rise to undesirable dispersive anti-phase contributions in the spectrum. In sequence (a) these anti-phase contributions are eliminated by the use of a purging pulse as described here. Of course, at the same time all magnetization other than x is also eliminated, giving a near perfect TOCSY spectrum without the need for phase cycling or other difference measures.

These purging pulses can be used to generate pure z -magnetization without contamination from zero-quantum coherence by following them with a $90^\circ(y)$ pulse, as is shown in the NOESY sequence (b). Zero-quantum coherences present during the mixing time of a NOESY experiment give rise to troublesome dispersive contributions in the spectra, which can be eliminated by the use of this sequence.