

Physics 383 Laboratory Manual

Hobart & William Smith Colleges



HWS Physics Department

Fall 2007

Contents

General Instructions For Physics Lab	iii
Recording data	iii
Lab reports	iii
Graphing	iv
Reporting of Experimental Quantities	iv
Collaborative and Individual Work	iv
General Admonition	v
Single- and Double-Slit Interference	1
MiniLab	1
MesoLab	3
Two-Slit Interference, One Photon at a Time	1
Introduction	1
Introduction to the Apparatus	3
Operation of the Apparatus	8
Theoretical Modelling for the Experiments	20
Conclusions	24
Appendix A: How to Read a Micrometer Drive	26
Appendix B: Quantitative Detection by Lock-In Techniques	27
Injecting Test Pulses into the PMT Amplifier	28
Magnetic Torque: Experimenting with the magnetic dipole	31
The Instrument	31
EXPERIMENT 1: Magnetic vs. gravitational torque	38
EXPERIMENT 2: Harmonic oscillation of a spherical pendulum	41

EXPERIMENT 3: Precession of a spinning sphere	43
EXPERIMENT 4: Magnetic field gradient force	46
EXPERIMENT 5: Demonstrating Magnetic Resonance	51
Pulsed NMR	55
Introduction and Theory	55
Experimental Set-up	66
Prelab exercises, part 1	70
Multiple pulse sequences	76
Prelab II	78
Lab II	79
Muon Physics	81
Introduction	81
The Muon Source	81
Muon Decay Time Distribution	83
Detector Physics	83
Interaction of μ^- 's with Matter	84
μ^+/μ^- Charge Ratio at Ground Level	85
Backgrounds	87
Fermi Coupling Constant G_F	88
Time Dilation Effect	88
Electronics	91
Data Acquisition	94
Procedure	99
Gamma Ray Attenuation in Matter	103
Introduction	103
Experimental Set-up	107

General Instructions For Physics Lab

Come to lab prepared – Bring quadrille-ruled, bound lab notebook, lab manual, calculator, and pen.

RECORDING DATA

Each student must record all data and observations in ink in his or her own lab notebook. Mistaken data entries should be crossed out with a single line, accompanied by a brief explanation of the problem, not thrown away. Do not record data on anything but your lab notebooks.

Put the date, your name, and the name of your lab partner at the top of your first data page. Your data and observations should be neat and well organized. It will be submitted as part of your laboratory report and must be legible.

LAB REPORTS

Specific advice on this matter is given at the end of each lab description, but your report will generally consist of:

Your data and observations, a summary of your experimental and calculational results, and a discussion of your conclusions. The summary and discussion may be written out right in the lab book (if done very legibly and neatly), or you may use word processors and/or spreadsheets if you prefer.

One typed page summarizing what you learned by doing the experiment.

Occasionally your instructor may give some additional instructions in class. It is important that the report be well organized and legible. If several measurements are to be compared and discussed, they should be entered into a table located near the discussion. Your data and results should always include appropriate units.

Repetitive calculations should not be written out in your report; however, you should include representative examples of any calculations that are important for your analysis of the experiment.

GRAPHING

Graphs should be drawn with a fine-lead pencil. A complete graph will have clearly labeled axes, a title, and will include units on the axes.

When you are asked to plot T^2 versus M , it is implied that T^2 will be plotted on the vertical axis and M on the horizontal axis.

Do not choose awkward scales for your graph; e.g., 3 graph divisions to represent 10 or 100. Always choose 1, 2, or 5 divisions to represent a decimal value. This will make your graphs easier to read and plot. When finding the slope of a best-fit straight line drawn through experimental points:

1. Use a best-fit line to your data. You may do this either by hand or using a graphing program, such as the one with Microsoft Excel, though it is recommended that you do your graphs by hand.
2. Do not use differences in data point values for the rise and the run. You should use points that are on your best-fit line, not actual data points, which have larger errors than your best-fit line.
3. Do use as large a triangle as possible that has the straight line as its hypotenuse. Then use the lengths of the sides of the triangle as the rise and the run. This will yield the most significant figures for the value of your slope.

REPORTING OF EXPERIMENTAL QUANTITIES

Use the rules for significant figures to state your results to the correct number of significant figures. Do not quote your calculated results to eight or more digits! Not only is this bad form, it is actually dishonest as it implies that your errors are very small! Generally you measure quantities to about 3 significant figures, so you can usually carry out your calculations to 4 significant figures and then round off to 3 figures at the final result. After you calculate the uncertainties in your results, you should round off your calculated result values so that no digits are given beyond the digit(s) occupied by the uncertainty value. The uncertainty value should usually be only one significant figure, perhaps two in the event that its most significant figure is 1 or 2. For example: $T = 1.54 \pm .06$ s, or $T = 23.8 \pm 1.5$ s. If you report a quantity without an explicit error, it is assumed that the error is in the last decimal place. Thus, $T = 2.17$ s implies that your error is ± 0.005 s. Because the number of significant figures is tied implicitly to the error, it is very important to take care in the use of significant figures, unlike in pure mathematics.

COLLABORATIVE AND INDIVIDUAL WORK

You may collaborate with your partner(s) and other students in collecting and analyzing the data, and discussing your results, but for the report each student must write his or her own summary and discussion in his or her own words.

GENERAL ADMONITION

Don't take your data and then leave the lab early. No one should leave the lab before 4:30 PM without obtaining special permission from the instructor. If you have completed data collection early, stay in the lab and begin your analysis of the data – you may even be able to complete the analysis before leaving. This is good practice because the analysis is easier while the work is still fresh in your mind, and perhaps more importantly, you may discover that some of the data is faulty and you will be able to repeat the measurements immediately.

Single- and Double-Slit Interference

by S. Penn

MINILAB

This *MiniLab* will allow you to become familiar with the basic structure of a laboratory experiment. The MiniLab has all of the components of a regular lab, but it requires only a single session. You will perform your data analysis and compare it to theory. Then you will write a short lab report due next class. In this way you will experience all aspects of a full lab, and you will know what is expected of you in performing those labs.

Purpose

To measure the two slit interference pattern produced by a laser passing through a double slit and to determine from these measurements the width and separation of the double slit.

Experimental Method

A HeNe laser, mounted on a labstand, is positioned on a lab table such that its light projects through a narrow double-slit onto a white board a few meters away. Check that the laser beam is normal to the plane of the white board, and then measure the distance from the slits to the white board. Tape a piece of white paper to the board and mark the locations of maximum and minimum intensity. Label the paper with the number of the slit and tape it into your logbook.

Remove the slit from its labstand and examine it under the micrometer microscope. Using the vernier scale on the micrometer, record the location of the four edges of the double slit. With the vernier scale you should achieve 1 micron accuracy. The micrometer has backlash, which is a looseness in the screw mechanism when the direction of motion is changed. Therefore, you should first position the micrometer off to one side of the slits. Then, moving only in one direction, when the micrometer is fully engaged (no backlash) record the position of the four sides.

The position of the slit edge is determined when the microscope vertical cross-hair best approximates the location of the edge. If you want a second set of measurement, you will need to back up the micrometer to the initial point and repeat your measurements.

For all measurements be sure to note your units and the precision of your measuring device. However sometimes the uncertainty is not determined by the precision of the instrument. For example, the edge of these slits are jagged. A very conservative estimate for the uncertainty would be half of the range of the uneven edge of the slit. However, we actually know the position better than that estimate suggests. Make several measurements of the slit edge positions. The mean and standard deviation of your measurements provide a good estimate of the position and uncertainty of a given slit edge. Use these values for the four edges to calculate the slit widths, the slit separation, and their uncertainties.

Theory

The interference pattern for two wide slits, assuming Fraunhofer diffraction (far field approximation), is given by

$$I(\theta) = I_0 \left[\frac{\sin(\Phi/2)}{\Phi/2} \right]^2 \cos^2(\Delta\phi/2)$$

where $\Phi = kD \sin(\theta)$, and $\Delta\phi = kd \sin(\theta)$. D is the slit width, d is the slit separation, and $k = \lambda/2\pi$ is the wave number. (See F. Crawford, *Waves: Berkeley Physics*, Vol. 3, Ch.9, Section 6)

Analysis

Notice that the intensity pattern consists of a series of narrow fringes within an envelope of a set of broad fringes. Use the equation above (and your knowledge of interference) to determine which slit dimension, width or separation, is responsible for which set of fringes, broad or narrow. Then determine the values of theta for each set of fringes that will result in maxima and minima. Note that the extrema for the two fringe need only be aligned at 0° where there must be a maximum.

Use your knowledge of Φ and $\Delta\phi$ at the extrema, plot each versus $k \sin(\theta)$ for the measurements associated with a given set of fringes. The slope of the line (and its uncertainty) will give the respective slit dimension (and its uncertainty).

Compare these results with the measurements using the microscope micrometer. Are they consistent?

Report

Write a brief report that includes all major sections of a lab report: Title, Abstract, Introduction, Experimental Apparatus and Method, Analysis,

Results and Conclusion. (Reference if necessary). You may use as a guide the following web sites:

- http://www.unc.edu/depts/wcweb/handouts/lab_report_complete.html
- <http://writing.colostate.edu/guides/processes/science/pop2a.cfm>
- <http://www.wisc.edu/writing/Handbook/ScienceReport.html>
- <http://www.unc.edu/depts/wcweb/handouts/sciences.html>

MESOLAB

In this exercise we will perform the initial measurements required to determine the slit dimensions for the *Two-Slit, One Photon* experiment. These measurements are broken out from that lab because a) it is a necessary measurement that has sometimes proved tricky in the past, and b) because the alignment of the slits in the experimental apparatus is difficult, we would not want to disturb the slits once they are in place. Therefore, we will all measure the slits now using the interference pattern and the microscope. We can then perform the analysis to check our work. Then the slits will be installed in the apparatus and not removed for the remainder of the semester.

This data will be recorded and analyzed separately by each group. We can perform the analysis and check the results in lab.

You will not have to write a report for this exercise. Instead these measurements and results will be incorporated into your report on the *Two-Slit, One Photon* experiment.

Purpose

To measure the width of a few single slits and the width and separation of a few double slits. These slits are use in the *Two-Slit, One Photon* experiment.

Method

A HeNe laser is set-up to shine through the double slit onto the white board. Measure the distance between the slit and the white board. Mark the positions of the fringes on the white board. In this case you should mark the dark fringe and the center of the bright fringe. Only when you have completely finished with this measurement should you remove the fringe and place it under the microscope. Use the micrometer on the microscope to measure the fringe width and separation. For all measurements be sure to note your units and the precision of your measuring device.

Theory

The interference pattern for two wide slits, assuming Fraunhofer diffraction (far field approximation), is given by

$$I(\theta) = I_0 \left[\frac{\sin(\Phi/2)}{\Phi/2} \right]^2 \cos^2(\Delta\phi/2)$$

where $\Phi = kD \sin(\theta)$, and $\Delta\phi = kd \sin(\theta)$. D is the slit width, d is the slit separation, and $k = \lambda/2\pi$ is the wave number.

The interference pattern for a single wide slit is obtained when the slit separation, d goes to 0.

$$I(\theta) = I_0 \left[\frac{\sin(\Phi/2)}{\Phi/2} \right]^2$$

(See F. Crawford, *Waves: Berkeley Physics*, Vol. 3, Ch.9, Section 6)

Analysis

Fit the results of the dark and bright fringe locations to the predicted values. Obtain values for D and d . Determine the uncertainties. Compare with the results from the microscope measurement.

Questions

When looking at a double slit, try covering one of the slits and noticing what happens to the interference pattern. What happens to the brightness of the central fringe? What happens to the intensity at the minimum adjacent to the central fringe? Can this be explained by the corpuscular (particle) theory of light?

Two-Slit Interference, One Photon at a Time

originally by David Van Baak, edited by S. Penn

INTRODUCTION

Wave-particle duality

Pick up any book about quantum mechanics and you're sure to read about *wave-particle duality*. What is this mysterious *duality*, and why should we believe that it's a feature of the real world? This manual describes the Teachspin apparatus which makes the concept of duality as concrete as possible, by letting you encounter it with photons, the quanta of light.

This apparatus makes it possible for you to perform the famous two-slit interference experiment with light, even in the limit of light intensities so low that you can record the arrival of individual photons at the detector. And that brings up the apparent paradox which has motivated the concept of duality – in the very interference experiment that makes possible the measurement of the wavelength of light, you will be seeing the arrival of the energy of light in particle-like quanta; in individual photon events. How can light act like waves and yet arrive as particles? This paradox has been used, by no less an authority than Richard Feynman, as the introduction to the fundamental issue of quantum mechanics:

In this chapter we shall tackle immediately the basic element of the mysterious behavior in its most strange form. We choose to examine a phenomenon which is impossible, *absolutely* impossible, to explain in any classical way, and which has in it the heart of quantum mechanics. In reality, it contains the only mystery. We cannot make the mystery go away by explaining how it works. We will just *tell* you how it works. In telling you how it works we will have told you about the basic peculiarities of all quantum mechanics." [R. P. Feynman, R. B. Leighton, and M. Sands, *The Feynman Lectures on Physics*, vol. I, ch. 37, or vol. III, ch. 1 (Addison-Wesley, 1965)]

You should find and read either of these famous chapters in the *Feynman Lectures* which introduce the central features of quantum mechanics

using the two-slit experiment as an example. Feynman notes that he discusses the experiment as if it were being done with particles with rest mass, such as electrons; he wrote in an era in which he was discussing a *thought experiment*. But since that time the two-slit experiment really has been done with neutrons [see *Reviews of Modern Physics* 60, 1067-1073 (1988)], and in this experiment, you too will translate a thought experiment into a real one, in this case with photons.

There are several technical advantages to the use of photons. They are easily produced and detected using the ordinary tools of optics. They propagate freely in air, and require no vacuum system. They are electrically neutral, and thus interact neither with each other nor with ambient electric and magnetic fields. They are (at high enough levels) directly visible to the eye. Finally, their wavelengths are (for typical visible-light photons) of a much more convenient size than the wavelength of electrons or neutrons. All these technical advantages make it possible to perform even the single-photon version of the two-slit experiment in a tabletop-sized and affordable instrument.

Historical context

There is a rich historical background behind the experiment you are about to perform. You may recall that Isaac Newton first separated white light into its colors, and in the 1680's hypothesized that light was composed of *corpuscles*, supposed to possess some properties of particles. This view reigned until the 1810's, when Thomas Young first performed the two-slit experiment now known by his name. In this experiment he discovered a property of destructive interference which seemed impossible to explain in terms of corpuscles, but is very naturally explained in terms of waves. His experiment not only suggested that such *light waves* existed, it also provided a result that could be used to determine the wavelength of light, measured in familiar units. Light waves became even more acceptable with dynamical theories of light, such as Fresnel's and Maxwell's, in the 19th century, until it seemed that the wave theory of light was incontrovertible.

And yet the discovery of the photoelectric effect, and its explanation in terms of light quanta by Einstein, threw the matter into dispute again. The explanations of blackbody radiation, of the photoelectric effect, and of the Compton effect seemed to point to the existence of *photons*, quanta of light which possessed definite and indivisible amounts of energy and momentum. These are very satisfactory explanations so far as they go, but they throw into question the destructive-interference explanation of Young's experiment. Does light have a dual nature, of waves and of particles? And if experiments force us to suppose that it does, how does the light know when to behave according to each of its natures? These are the sort of questions which lend a somewhat mystical air to the concept of duality.

Of course the deeper worry is that the properties of light might be not merely mysterious, but in some sense self-contradictory. You will be confronted with just the sort of evidence which has led some persons to worry that either light, or our theories for light, are not only surprising

but also inconsistent or incoherent. As you explore the phenomena, keep telling yourself that the light is doing what light does naturally, and keep asking yourself if the difficulties lie with light, or with our theories of light, or with our verbal, pictorial, or even mechanical interpretations of these theories.

Goals for this apparatus

It is the purpose of this experimental apparatus to make the phenomenon of light interference as concrete as possible, and to give you the hands-on familiarity which will allow you to confront duality in a precise and definite way. When you have finished, you might not fully understand the mechanism of duality – Feynman asserts that nobody really does – but you will certainly have direct experience of the actual phenomena that motivate all this discussion. Here, then, are the goals of the experiments that this apparatus makes possible:

1. You will be seeing two-slit interference visually, by opening up an apparatus and seeing the exact arrangements of light sources and apertures which operate to produce an 'interference pattern'. You'll be able to examine every part of the apparatus, and make all the measurements you'll need for theoretical modeling.
2. You will be able to perform the two-slit experiment quantitatively, recreating not only Young's measurement of the wavelength of light, but also getting detailed information about intensities in a two-slit interference pattern which can be compared to predictions of wave theories of light.
3. You will be able to perform the two-slit experiment one photon at a time, continuing the same kind of experiments, but now at a light level so low that you can assure yourself that there is at most one relevant photon in the apparatus at any time. Not only will this familiarize you with single-photon detection technology, it will also show you that however two-slit interference is to be explained, it must be explained in terms that can apply to single photons. [And how can a single photon involve itself with two slits?]
4. You will be exploring a pair of theoretical models which attempt, at differing levels of sophistication, to describe your experimental findings. You will thus encounter the distinctions between Fraunhofer and Fresnel diffraction theories in a concrete case, and you will also learn the difference between a mathematical, and a physical, description of what is going on.

INTRODUCTION TO THE APPARATUS

Familiarization with the Apparatus

The apparatus consists of a long, narrow, rectangular metal box, containing the light sources, attached to a squat metal box, which contains

the detectors and the signal electronics. Orient the long assembly on its wooden feet so that the detector box is at the righthand end.

One caution before we begin. The detector box contains a photomultiplier tube (PMT), which is an extremely sensitive, single-photon detector. It is quite easy to damage this detector if it is exposed to intense light or voltage spikes. Therefore, before turning on the apparatus, it is important to check that the PMT voltage is off and that the PMT's shutter is closed. On the electronics box, find the high voltage (HV) controls and ensure that the HV is off and that the voltage dial (10-turn) is set to 0.00, fully counter-clockwise. The PMT shutter is controlled by a rod that protrudes vertically from the junction between the light box and the detector box. With the light box closed, pull gently up on this rod just enough to convince yourself that the shutter moves, then close the shutter once more by pushing it all the way *down*.

Now, remove the cover of the light box by opening the four latches that secure the top using a slight lift and turn. Then grip the left end of the tube and lift the left end of the lid slightly. After you feel it disengage, slide the cover to the left. When the left end disengages from its slot, the lid can be removed.

[Take this opportunity to learn how to re-install the cover, making sure that you can engage its right end, lower its left end, and re-engage the four latches which hold it in place.]

With the cover open, you are ready to look over all the parts of the apparatus:

1. At the left end are two distinct light sources, one a red laser and the other a green-filtered light bulb; also at the left end are the controls for the light sources.
2. Found along the length of the long box are the various slits and apertures that form a Young's two-slit experiment. On the front of the box are two 'micrometer drives', which allow you to make mechanical adjustments to the two-slit apparatus.
3. Finally, there are two light detectors: the *photodiode*, which is used with the laser, and the *photomultiplier tube* (PMT), which is used with the low-intensity incandescent bulb.
 - (a) The photodiode produces a photoelectric current when exposed to visible light. This current is converted to a voltage and amplified in the electronics box. The photodiode is attached to the light shutter such that it is in position when the shutter is in its down position.
 - (b) The PMT is safe to use **only when light box lid is completely closed and the green-filtered incandescent light is in use**. It is exposed to light only when the shutter is in its up position. When a photon strikes the light sensitive emulsion of the PMT, it knocks out an electron. The PMT has a series of dynodes with a cascading electric field between them that grows in strength with each successive dynode. The ejected electron is accelerated in the cascading field such that when it strikes the dynode it knocks out

additional electrons. This multiplication process continued with each dynode. In this way a single photon is able to trigger a large voltage pulse on the PMT's anode.

Finally, there are electrical connections. Electric power comes from a power module, plugged into your main power outlet and connected by a cable to the left end of the apparatus. Two more cables run from this end to the detector box, and supply the photodiode and PMT detectors with the power they need for operation.

Alignment of the Apparatus

Before you take on the task of aligning the apparatus, it is worth examining all the slits. Each slit is stamped in a metal foil, whose edge is attached to a magnetic fixture which allows it to be placed conveniently in a slit-holder in the apparatus. Practice taking the source slit right off of its holder, and look it over with a magnifying glass. You're looking at a single slit, and you'll want to know its width; this can be measured by single-slit diffraction using a separate tabletop set-up and a HeNe laser, or much more directly by viewing the slit in a low-power microscope whose sample holder is equipped with a travelling stage.

By either technique, you can also characterize the detector slit, and the wide slit whose two edges are used to execute the slit-blocking functions. Finally, examine one or more of the double-slit assemblies that are available; for these, you will want not only to measure the width of the two individual slits, but also to characterize their separation. However you measure this separation, it is conventional to characterize a double slit by the center-to-center separation of the two slits.

As you replace the slits, you'll note that they can be translated both vertically and horizontally to a limited extent; it is also possible to rotate them to some extent in their own plane. For best results, you want all the slits to have their long edges set to be accurately vertical; for the moment, use an 'eyeball accuracy' criterion for this.

Now it's time to align the whole apparatus; the goal will be to get the apparatus aligned well enough so that you can switch from laser to light-bulb optical source without needing to disturb anything at all in the whole arrangement of slits. Here's the procedure:

1. Ensure that the PMT shutter is closed, and the PMT bias (voltage) is turned off (by toggle switch) and down (to 0.00 on the 10-turn dial). Now connect the line cord to mains power outlet, and learn how to turn on the laser source, and the light-bulb source – use a paper card to see each source's output beam.
2. Learn how to raise the black cylindrical light-bulb source up into the beam path; you may raise its left-hand (bulb-holding) end until it reaches the level of the red laser beam. This brings the light bulb up to the height at which it will be used.
3. Remove the green filter holder from the right-hand (output) end of the bulb source – the black plastic cylinder will slide off the black

tube. Level the black tube, turn on the bulb source, and set it to about 6 on its adjustable intensity scale. Use a paper card, and a darkened room, to follow the white light it sends to and through the source slit. Follow the widening ribbon of white light downstream to the vicinity of the double-slit, and (by lateral adjustment of the position of the source slit on its holder) ensure that the cm-wide ribbon of light is approximately centered in the box. This will ensure that light from the bulb can reach the double slits.

4. Turn the bulb source down in intensity and then off, and push the whole bulb-holder assembly down into its stowed position at the bottom of the box. Now turn on the laser and confirm, by following it with a paper card, that its output beam passes over the top of the bulb holder and reaches the vicinity of the source slit.
5. Acquaint yourself with the mechanical adjustments for the laser. There are two brass thumbscrews emerging from the front wall of the apparatus. The right screw is sealed at a set position. Please do not attempt to adjust it. The left screw will adjust the laser angle in the horizontal plane. Beneath the light box is a nylon thumbscrew that adjusts the laser angle in the vertical plane. Use the thumbscrews until the black block is approximately centered in the apparatus and aligned with it.
6. Now you're ready to use the laser beam itself as an alignment tool. By making similar adjustments to both brass screws, translate the black block laterally in the box until some of the laser beam passes through the source slit. (Do not worry yet if that light reaches the double-slit area.) Instead, use that light as a diagnostic for the vertical alignment of the laser's black block: use the nylon thumbscrew to raise the back of the block, thereby lowering the direction of the laser beam; or conversely, back off the nylon thumbscrew while pushing down the back of the block, thereby raising the direction of the laser beam. You may have to fuss a bit to keep some diagnostic light passing through the source slit; your goal is to get the bright part of the laser beam to reach the double-slit area at the correct height, to a 1-cm tolerance or better.
7. Finally, you want to make the lateral adjustments of the laser block which will not only send the maximum power through the source slit, but will also do so in such a direction that the beam will continue and reach the double-slit area. Here are two procedures which may be used—
 - (a) Think of the source slit as a fixed fulcrum of a lever, and think of the laser beam as a ray that is pivoting about the fulcrum. By seeing where light passing through the source slit reaches the double-slit area, decide how you need to rotate this ray. This will require translating the laser block (by making identical adjustments to both brass screws) and also rotating the laser block (by making differential adjustments to the two brass screws) to keep its output beam passing through the source slit.
 - (b) When you are close (ie. when light passing through the source slit comes within a few mm of being laterally centered on the double slits), you can use an alternative technique. Remove the source slit altogether, and then use the brass screws differentially to rotate the laser block a trifle, until you have the laser beam reaching

the double slits, centered to a mm or so. Now replace the source slit, and hand-translate it laterally in its holder, by tiny steps, until it is centered on the laser beam – your diagnostic is to maximize the throughput of laser power through the source slit. You should now see the source slit's single-slit diffraction pattern projected on, and centered on, the double-slit holder; and ideally, you should have changed the lateral position of the source slit only a trifle to achieve this. Still, you'll need to go back to alignment step 3) to check that with the new position of the source slit, white light can pass from the bulb, through it, and reach the double slit.

8. Now replace the green filter holder on the bulb source tube.
9. There is next an alignment task for the slit-blocker; your goal is to get its two long edges aligned to be vertical. You may start by removing the slit-blocker's (wide) slit from its magnetic mount; remember the slit-blocker is mounted on the micrometer-adjusted, flexure-supported moveable block just downstream of the double slits. The best diagnostic for the proper positioning of the slit-blocker's wide slit is to use the two ribbons of light emerging beyond the double slit when they are illuminated by the laser source. Find these ribbons on a paper card, and now re-install the wide slit on its mount. Use the micrometer to translate the wide slit laterally until both ribbons emerge, and now dial the micrometer until you see the knife-edge function of the slit-blocker in cutting off one of the ribbons of light. If the wide slit is in place with its edges not correctly vertical, this knife-edging of the ribbon will proceed not all at once, but bottom-to-top or top-to-bottom. Hand-adjust the wide slit by small in-plane rotations on its magnetic holder until you achieve the desired all-at-once chopping-off of the ribbon of light.
10. The equivalent alignment task for the detector slit is to rotate it in its own plane, until its length is accurately parallel to the double-slit fringes which it is to scan. The best diagnostic for this is the contrast of the interference pattern that you will record in section 1II.B. below; so for now, just be sure that the detector slit is vertical to 'eyeball accuracy'.

Ancillary equipment needed for operation

There are a few more electrical functions of the apparatus that you should now encounter, and a few electrical tools that you will need to assemble for operation of the apparatus.

You have already operated the light-bulb source; it has an on-off switch and an intensity adjustment. You've also operated the laser source, thus far merely on or off. For some modes of operation, you might want to turn the laser on and off repeatedly, and not just by toggling its switch. For that purpose, there is a laser-modulation input connector at the source end of the apparatus; it's used in the optional lock-in mode of detection described in Appendix B.

At the detector end of the apparatus, you'll need a digital multimeter for reading the output voltage which emerges from the photodiode-amplifier mode of detection. You can understand this signal by noting that the photodiode translates input optical power to output electrical current with a sensitivity of about 0.4 A/W, and then the photodiode amplifier translates photodiode current to output voltage with a conversion constant of 22×10^6 V/A.

Exercise: Show that the conversion of photons of red light to electron-hole pairs in a semiconductor with 100% efficiency would yield about 0.5 A of output current for an input of 1 Watt of optical power.

So if we see an output of 1.0 Volt, we can infer an input current of 4.5×10^{-8} A = $0.045 \mu\text{A}$, and further infer an incident optical power of 11×10^{-8} W, or $0.11 \mu\text{W}$, on the photodiode. The photodiode amplifier has a time constant of 0.4 ms, which sets the timescale for the response time of this detector system.

Also at the detector end of the system are the electronics for running the photomultiplier tube. A PMT requires a high-voltage power supply for its operation, and this is integrated into the socket of the PMT. There is a toggle switch to activate this supply, and a 10-turn dial to set the high voltage somewhere in the range 0 – 1000 V dc. There is also a pair of monitor outputs, at which a potential difference of 10^{-3} of the PMT bias voltage is available for inspection. You may use the approximate conversion constants of 1.00 turn of 10-turn dial = 0.10 V at monitor output = 100 V of PMT bias. The *gain* of the phototube, or the number of multiplied electrons out per single photoelectron event, rises exponentially with this bias voltage, and reaches about 10^6 at full bias. You will not be seeing this charge pulse directly, since it is sent directly into a pulse amplifier & discriminator module. What you will need is a moderately fast digital oscilloscope to monitor the amplified pulses, and a TTL-compatible counter to count those charge pulses which activate the discriminator and lead to the generation of a single-shot, fixed-width output *event* pulse.

If you wish to investigate the operation of the discriminator in detail, there is a provision to inject *test pulses* into the pulse amplifier. To do this, follow the directions of Appendix C.

OPERATION OF THE APPARATUS

This section of the manual will introduce you to the functions of the apparatus, leading you through three stages of understanding until you have seen two-slit interference, one photon at a time.

Visual mode of operation

For this mode of operation, you will be working with the cover of the apparatus open. You will need to use neither light detector, since you will be using your eyes as the only detector needed. You'll be using the laser as light source, and you'll find a supply of business cards or other small white paper screens to be convenient tools.

Use the controls at the left end of the apparatus to turn on the solid-state diode-laser light source – it's contained within a black metal block. Use a paper screen to find its bright red output beam; the diode laser manufacturer asserts that its output wavelength is 670 ± 5 nm [or 0.670 ± 0.005 μm], and its output power is about 5 mW. [So long as you don't allow the full beam to fall directly into your eye, it presents no safety hazard.]

The laser beam propagates to the right, and it will run into the light-bulb source if that's in its up position. So locate the black cylindrical tube which contains the light bulb, and push down the left end of that tube until the whole tube lies flat against the bottom of the box. You'll need to unscrew the adjustable support post underneath the bulb source. Now follow the laser beam until it reaches the entrance slit, or *source slit* of the two-slit apparatus; this is a single slit, of height about 1 cm, but width of only 0.085 mm. If your apparatus is aligned, the slit will be neatly straddling the laser beam, so that a good fraction of the laser light will be passing through the slit – use your paper card to see if this is so.

Light passing through this narrow source slit will undergo single-slit diffraction, and you can follow this process by moving your viewing card downstream. You will see the red beam spread out horizontally, reaching a width of about a cm by the time it reaches the middle of the apparatus.

In the middle of the apparatus is the holder for the double slit, which is a structure with two rectangular apertures, again each about a cm high, and each with width of only 0.085 mm, and with center-to-center separation of 0.353 mm. [All of the slit structures are mounted on magnetic holders so that they can be removed, examined, and re-positioned; if your apparatus is already aligned, it would be a pity to spoil this alignment by disturbing the slits now.] Instead, put your viewing card just downstream of the two-slit structure, and look (close up!) to see if you can observe the two ribbons of light, just 0.35 mm apart, which emerge from the two slits.

This is your chance to understand the function of the slit-blocker, just downstream of the double-slit system, which will allow you selectively to block the light coming from either of the two slits. This slit-blocker is controlled by the micrometer screw on the front center of the apparatus; rotating the micrometer's knob will move the slit-blocker laterally across the two ribbons of light emerging from the two-slit system. For the present, find a position for the micrometer adjustment which permits the two ribbons of light to emerge and continue rightwards in the apparatus.

Each of the narrow ribbons of light emerging from a slit will continue to diffract; follow these broadening ribbons downstream until they visibly overlap. By the time your viewing card reaches the right-hand end of the apparatus, the overlap will be nearly complete; and you'll see that the two overlapping ribbons of light combine to form a pattern of illumination displaying the celebrated *fringes* named after Thomas Young. How would you characterize these fringes? Can you describe them qualitatively? Can you distinguish between your description of the phenomenon you see and your hypothesis for its cause?

Now position a viewing card at the downstream end so you can refer

to it for a view of the fringes, and take another card back upstream to the vicinity of the slit-blocker. Learn to dial the slit-blocker's micrometer adjuster [the one at the *middle* of the apparatus] until you see how to use it to block the ribbon of light coming from either the farther, or the nearer, of the two slits. Start by rotating the multi-turn micrometer screw fully clockwise, and watch this adjustment push the slit-blocker away from you. Take this opportunity to learn how to read a micrometer dial – see Appendix A if you haven't done this before – and now, as you dial the micrometer counter-clockwise, find and record five settings for the slit-blocker's micrometer:

- one position for which both slits are blocked;
- another for which light emerges only from the farther of the two slits;
- a third (anywhere in a wide range) that allows both ribbons of light to emerge;
- a fourth for which light emerges only from the nearer of the two slits; and finally,
- a fifth setting (and highest reading) which again blocks the light from both slits.

It is essential that you are confident enough in your ability to read, and to set, these five positions that you can do so even when the box cover is closed (when you can't, as now, confirm your results by checking with a white viewing card).

Once you've gotten these five settings, let the light reach your viewing card at the far-right end of the apparatus, and find out what happens there, qualitatively, for each of the five settings. In two of them, no light at all will arrive; in the third of them, light will arrive from both slits to form Young's two-slit interference fringes. What happens in the other two cases, when light from only one slit arrives?

In particular, observe what occurs at some *particular locations* on your screen:

- at a bright fringe, or *interference maximum*, what happens to the light intensity when you use the slit-blocker to cover one slit?
- at a dark fringe, or *interference minimum*, what happens to the light intensity *at that location* when you use the slit-blocker to cover one slit?

You should be able to use the language of interfering light waves to describe what you see, and to explain why it happens. You'll be investigating this behavior quantitatively in later parts of the experiment, but for now you need to be familiar with it qualitatively, and in terms of causes.

There's one last piece of the apparatus which you can now learn to use. At the far-right end of the apparatus is a final optical element; it's an exit slit, or *detector slit*, of the same size and character as the source slit, except that it's mounted on a moveable structure like the slit-blocker, so it too can have its position adjusted by a micrometer screw drive. The

purpose of this detector slit is to allow light from a narrow slice of the interference pattern to pass along to the end of the long apparatus and into the detector box. By translating the detector slit laterally along the interference pattern in space, you can select which part of the pattern will have its light sent on to the detector. Thus by scanning the micrometer screw of the detector slit, you can scan over the interference pattern, eventually mapping out its intensity distribution quantitatively. For now, ensure that the detector slit is located somewhere near the middle of the two-slit interference pattern.

You are now acquainted with every part of this version of Young's experiment, and with the two *independent variables* you can control: one is the position of the slit-blocker, and you have recorded settings corresponding to each of five slit conditions; the other sets the location of the detector slit. Now you're ready to go on to quantitative measurements.

Quantitative mode of operation

This mode of operation of the apparatus continues to use the laser light source, but it begins to use the photodiode detector to survey quantitatively the intensity distribution of the interference pattern, by varying the position of the detector slit. You might conduct these measurements with the box cover open, but room light will contribute excessive and variable contributions to your signals, so now is the time either to dim the room lights or to close up the cover of the apparatus. For convenience, have the slit-blocker set to that previously determined setting which allows light from both slits to emerge and interfere.

The shutter of the detector box will still be in its closed, or down, position. This blocks any light from reaching the PMT, but the shutter in its down position correctly centers a 1-cm² solid-state photodiode, which acts just like a solar cell in actively generating electric current when it's illuminated. The device is equally sensitive everywhere over its area, so it would record *all* the light in the whole interference pattern if it were not for the detector slit. But with the detector slit at a fixed position, the only light reaching the detector is that from a selected part of the interference pattern; by this means, a single-element, spatially-fixed, large-area detector can measure (serially in time) the intensities at various places in the interference pattern. The method of course relies on the fact that the rest of the apparatus is stable in time; happily, the diode-laser source has an output power varying by $< 0.1\%$ in time, and the mechanical stability of the rest of the apparatus is also adequate.

The electric current from the photodiode is conducted by a thin coaxial cable to the INPUT BNC connector of the photodiode-amplifier section of the detector box. At the OUTPUT BNC connector adjacent to it, there appears a voltage signal derived from that photocurrent. Connect to this output a digital multimeter set to 2-Volt sensitivity; you should see a stable positive reading.

To determine if this reading means anything, go back to the left end of the apparatus and use the 3-position toggle switch to turn the laser source off. This should reduce the voltage signal you've been seeing, but perhaps not to zero; record the value you see, and take it to be the *zero*

offset of the photodiode-detector system. You might turn off the room lights to confirm that the signal you see is actually an electronic offset, and not the leakage of light into your apparatus. The zero-offset reading will eventually need to be subtracted from all the other reading you make of this output voltage.

Turn your laser source back on, and now watch the photodiode's voltage-output signal as you vary the setting of the detector-slit micrometer. If all is well, you will see a systematic variation of the signal as you dial the micrometer; you are scanning over the interference pattern. You'll find a variety of maxima, and you should try to find the highest of the maxima – this is the *central fringe* or 0^{th} -order fringe which theory predicts. Between the various maxima, you should see minima; and if your alignment is good, the signal at these minima should drop nearly to the zero-offset signal you previously recorded. These deep minima are of course the manifestation of destructive interference.

The size of the signals you observe will be somewhere in the vicinity of 5 Volts at the central maximum, but the number will depend a great deal on the details of your alignment procedure. If the (offset-corrected) signal is less than 1 Volt, you will probably want to improve the alignment. Here's a very approximate calculation that suggests why signals of order 1-5 Volts are to be expected. Start with 5 mW of laser power, and suppose that 1 mW of that makes it through the source slit. This power diffracts laterally to about 10 mm width, and thus one of the double slits, of width about 0.1 mm, can pass only about 1% of this, or $10\ \mu\text{W}$. This power in turn diffracts out to a stripe about 10 mm wide at the plane of the detector slit, which in turn can pass only about 1%, or $0.1\ \mu\text{W}$, to the photodiode detector. That yields about $0.04\ \mu\text{A}$ of photodiode current, and hence about 0.9 Volts at the photodiode signal output. If we take into account not only one of the double slits, but interference of light from both of them, we can understand why signals of order 1-5 Volts are to be expected. Nevertheless, these are very large signals in terms of rate of photon arrival.

Exercise: compute the rate of arrival of red photons which will deliver energy at a rate of $0.1\ \mu\text{W}$.

Before you go on to record data systematically, park the detector slit at the location of the central maximum, and then go over to the slit-blocker micrometer screw, and set it to a (previously determined) reading that you know will permit light to pass through only *one* of the slits. Check what photodiode signal you now are getting – it should be less than before. Again, set the slit-blocker to permit light only from the *other* of the two slits, and record another diminished signal. Can you explain why the signals have diminished? Corrected by subtraction of zero-offset, by what factor should they have diminished? [Answer: *Not* to 50%, but to a *smaller* fraction, of the original intensity. Why?]

To see another and even more dramatic manifestation of the wave nature of light, set the slit-

blocker again to permit light from both slits to pass along the apparatus, and now place the detector slit at either of the *minima* immediately adjacent to the central maximum; take some care to find the very bottom

of this minimum. Now you're seeing the effect of destructive interference – what will happen when you use the slit-blocker to block the light from one, or the other, of the two slits? [Answer; Blocking fully half the light coming through the two-slit assembly is going to *raise* the signal you're looking at – to what level? and why?]

Once you have performed these spot-checks, and have understood the motivation for them, and the explanation of their results, you are ready to conduct systematic measurements of one dependent variable (the photodiode voltage-output signal) as a function of two independent variables. What you want are three graphs, each giving the voltage-output signal as a function of the detector-slit position; the graphs will be for one slit open, the other slit open, and both slits open. You will want occasionally to block both slits to get a measure of the zero-offset signal which needs to be subtracted from all the readings you take. You will learn, by trying, what spacing of detector-slit positions to use. You will learn this fastest if you, or your partner, plot the data as you take it – nothing beats an emerging graph for teaching you what is going on.

The data you have obtained are along a calibrated horizontal scale – if you read Appendix A, you will know what the readings of the detector-slit micrometer imply, quantitatively, for position along the interference pattern. The data are also obtained on a quantitative vertical scale, though this one is in *arbitrary units* – we have not translated voltage-output units into light-intensity units, but you may be assured that the (offset-corrected) signal you obtain is linear in light intensity. Thus your data can be directly compared with theoretical models of single-slit diffraction and two-slit interference; section IV of this manual discusses two models which can be used to explain your data. In particular, the spacing of your interference minima and maxima gives a quite direct measure of the wavelength of the red laser light you are using.

The data you obtain are also serve as another diagnostic of alignment procedures. You have data plotted for both cases of one-slit-blocked, and these should display characteristic single-slit- diffraction features. In particular, the heights of these two graphs should be equal; if they are not, it could be that the source slit is unequally illuminating the two slits of the double-slit assembly. You also have data for the Young's-experiment two-slit-interference case, and the depth of the minima of this curve are another searching test of alignment. With moderate care, you can achieve a contrast ratio,

$$\frac{\text{offset corrected central maximum signal}}{\text{offset corrected first minimum signal}} \approx 30$$

If your contrast ratio is markedly worse, you might wonder if the double slits are being unequally illuminated, if stray light is reaching your detector (which would wash out the contrast), or if the scanning detector slit might fail to have its long dimension carefully parallel to the long straight interference fringes. You want the detector slit to be oriented to permit an alignment error of only a fraction of its width, say 0.05 mm, over the 10 mm of fringe height that can illuminate the detector. This means the fringe contrast is sensitive to a rotation of the detector slit in its own plane by less than 1°.

Since the slit-blocker provides a way to give the signal when light from neither slit is reaching the detector directly, you have a fine operational method for finding the background level that should be subtracted from the data. Then you may be guided by theory to expect, with the detector slit parked at the central maximum, background-corrected signals in the proportions 1:4:1 for the use of one:both:other of the two slits. You may also be guided by theory to expect, with the detector slit parked at the locations of the innermost minima, background-corrected signals in the proportions $1:\varepsilon:1$ for the use of one:both:other of the two slits; here ε is related to the inverse of the contrast ratio previously discussed. There's no need to obsess yet about departures from these theoretical ideals, since they apply only in certain limiting cases; rather, you should compare your results to the expectations from a corpuscular view of light, in which the addition of 1 unit of light from one slit, and 1 unit of light from the other slit, ought always to yield 2 units of light. If instead you have displayed a result much nearer to 4 than to 2, you have falsified some corpuscular views. More dramatically still, you have shown that the addition of 1 unit to 1 unit not only can yield a result smaller than 2, it can yield a result markedly smaller than 1! This is naturally explained by a wave theory, but you should put into words in just what sense it is evidence against a corpuscular view of light.

Single-photon mode of operation

If you have gotten the visual and quantitative modes of operation of this apparatus to work, you are ready for this section; but if you have not found interference fringes, trying single-photon detection is not going to cure your problems. So this section assumes

- that you know how to use the slit-blocker, and have left it in position to pass light from both slits; and
- that you know how to find the central maximum of the interference pattern, and have left the detector slit parked to pass light from that central fringe to the detector box.

Now you're ready to go on to the use of the photomultiplier tube, or PMT, which makes available to you electrical pulses that correspond to the detection of light, one photon at a time.

But first a WARNING: a photomultiplier tube is so sensitive a device that it should not be exposed even to moderate levels of light when turned off, and must not be exposed to anything but the dimmest of lights when turned on. In this context, ordinary room light is intolerably bright even to a PMT turned off, and light as dim as moonlight is much too bright for a PMT turned on. That is why the PMT box is equipped with a shutter, and the apparatus as a whole is equipped with a buzzer alarm, to help protect the PMT from misadventure. The goal of these safety measures is to assure that the PMT is used

- *only* when its box is coupled to the two-slit apparatus, and
- *only* when the cover of the apparatus is closed, and

- *only* when the light-bulb (and *not* the laser) source is being used inside the box.

Exercise: To understand the implications of light leaks, and to use some numbers typical of the PMT in this apparatus, suppose that *full sunlight* delivers 10^3 W/m^2 to the earth, and that *full moonlight* is a million times dimmer. Suppose that this moonlight is delivered by photons of yellow light, that they fall on a photocathode of area $8 \times 25 \text{ mm}^2$, and that they are converted to photoelectrons with efficiency 4%. Suppose finally that each photoelectron emitted at the photocathode is amplified to a pulse of 10^6 electrons arriving at the anode. Calculate:

1. the arrival rate of photons at the photocathode,
2. the emission rate of photoelectrons,
3. the arrival rate of electrons at the anode,
4. the average anode current this represents, and
5. compare that current with the manufacturer's suggested limit of $1 \mu\text{A}$ for average anode current.

Before you do anything with the PMT, locate on the detector box the HIGH-VOLTAGE toggle switch which activates its power supply, and ensure that it's turned off, also turn down to 0.00 the 10-turn dial which will later set the voltage that enables the process of electron-multiplication inside it. Also, ensure that the shutter of the box is still in its down position.

Now it's both safe and necessary to open the cover of the apparatus, because you're about to change from laser to light-bulb illumination of the two-slit apparatus. When you open the cover, you might find the laser still running; use the 3-position toggle switch to turn it off. Find the light-bulb source (the black cylindrical enclosure) and lift its right-hand end until it's tilted up by about 20° into the air. The thickest, output, end of the source tube is a structure which holds a green optical filter. You should now hold the source tube with your left hand, and use your right hand to rotate and slide that filter-structure off the source tube. (Keep all fingerprints off the glass narrow-band interference filter.) Now turn the bulb power adjustment dial down to 0 on its scale, and set the 3-position toggle switch to the BULB position; dial the bulb adjustment up from 0 until you see the bulb light up.

The humble #47 flashlight bulb you're using will live longest if you minimize the time you spend with it dialed above 6 on its scale, and if you toggle its power switch only when the dial is set to low values.

Now it's time to get the bulb into position to send light along the apparatus in place of the laser source. The goal is to achieve that without changing the position of any of the slits (source, double, or detector slits). To do this, you can use the cylindrical bulb housing as a sort of handle to pull upwards the left-hand end of the bulb housing, until it reaches the height previously attained by the laser beam. (You might even temporarily turn on the laser, and watch its beam strike the rear end of the

bulb structure, to achieve this.) Now preserve the height of the left end of the bulb structure, while lowering the right end, until the black tube is approximately level; no great precision is required. Finally, you can re-position the adjustable support structure which holds the light source in position.

To confirm that the bulb is in the correct position, you will need to be able to darken the room completely, and you'll need a dim flashlight and a white paper viewing card. Set the bulb's brightness to about half scale, and follow its splatter of white-light emission to the source slit. Now find the narrow ribbon of light which emerges downstream of the source slit, and trace it along the apparatus in the direction of the double-slit structure. As you travel downstream, the white band will widen horizontally, not only because of diffraction, but also because the filament of the light bulb is extended a few mm horizontally, and the source slit is acting like a one-dimensional pinhole camera. So the ribbon of light will be about a cm wide when it reaches the double-slit holder; what you need to ensure is that this cm-wide stripe of light is on-center in the apparatus and is thus falling on the double-slit structure. If the cm-wide ribbon is off-center and entirely missing the slits, you will need to go through the alignment procedure discussed in section II.B. of this manual. But if the ribbon illuminates both slits, you're all set; it is not necessary that it be perfectly centered in the apparatus.

Now carefully put the green filter-holding structure onto the right-hand, downstream end of the light-bulb source, being sure not to disturb the position of the left-hand, bulb-holding end of the structure. The green filter blocks nearly all the light emerging from the bulb, allowing to pass only wavelengths in the range 541 to 551 nm. This is a small fraction of the total light, and while you will now be able to see a ribbon of green light emerging from the source slit, you will probably not be able to follow it very far downstream. No matter; plenty of green-light photons will still be reaching the double-slit structure – in fact, you should now dim the bulb even more, by setting its intensity control down to about 3 on its dial.

Exercise: To see why you need to turn down the light bulb, and to gain familiarity with intensities expressed in units of photons/second, try to estimate – to the nearest order of magnitude – some photon transport rates:

1. Assume a #47 bulb runs at 6.0 V, 0.3 A, and converts electrical energy into light with 5% efficiency. Further suppose that it puts out (on average) photons of red light. What is the total production rate of photons?
2. Assume that the output spectrum is in fact spread over the range 500 to 1500 nm, and that the green filter passes only a 10-nm bandwidth; what is the production rate of photons which would pass through the green filter?
3. Suppose these photons are emitted in all directions, but that all are absorbed except for those which reach a slit of size $0.1 \times 10 \text{ mm}^2$, located about 100 mm from the bulb; what is the rate at which green photons will pass through this entrance slit?

4. Suppose that photons passing through the entrance slit spread to cover a 1 cm^2 area by the time they reach the double slits; what is the rate at which green photons pass through the double slits?
5. Suppose that photons passing through the double slits diffract to cover a 1 cm^2 area by the time they reach the detector slit; what is the rate at which green photons pass through the detector slit?
6. Finally, given the PMT efficiency of about 4%, what rate of photon events would result from all these assumptions?

You should find that further dimming of the bulb is necessary; that's why the apparatus makes it possible to adjust the bulb's intensity. If we turn it down to half voltage, it gets about half the current, hence $1/4$ the power; more importantly, the bulb cools off and the peak of its blackbody spectrum moves toward the infrared. The fraction which passes the green filter, on the short-wavelength side of the blackbody peak, plummets exponentially for lower bulb temperatures, so it's very easy to reduce the green emission not just by 10-fold but by a factor of 10^3 or more. The result is that it's feasible to achieve the enormous dilution of photon count rate down to the level desired.]

You may now return to ordinary room illumination, and you must close up the cover of the box. Be sure that first the right end, and then the left end, of the cover are fully engaged, and then that all four latches are in place. Now, and only now, are you ready to start activating the photomultiplier tube (PMT) to detect the photons that are flying through the box.

You'll need a reasonably fast ($> 20 \text{ MHz}$ bandwidth), preferably digital, oscilloscope for first examination, and a digital counter sensitive to TTL-level(+4 V positive-going) pulses for eventual counting, of photon events. Set the oscilloscope to about 50 mV/division vertically, and 200 ns/division horizontally, and set it to trigger on positive-going pulses of perhaps $> 20 \text{ mV}$ height. Now find the PHOTOMULTIPLIER OUTPUT of the detector box, and connect it via a BNC cable to the vertical input of the oscilloscope; use a 50Ω terminator at the oscilloscope. You are about to look at pulses, each of which starts with the light-induced ejection of a single electron from the light-sensitive photocathode of the PMT. Those photoelectrons will be amplified by about 10^5 inside the PMT, and arrive as a pulse of about 10^5 electrons, or about $-2 \times 10^{-14} \text{ Coulombs}$, at the input of a charge-sensitive amplifier. That device will convert that pulse of negative charge to a positive-going voltage pulse, which you will see on the oscilloscope.

When you are ready with these electronics, keep the shutter closed, set the HIGH-VOLTAGE 10- turn dial to 0.00, and turn on the HIGH-VOLTAGE toggle switch. You are now ready to start dialing up the high-voltage supply which provides the potentials inside the PMT needed for the electron-multiplication process; watch the oscilloscope display as you start to dial up the voltage, about a turn at a time. [Each full turn raises the PMT bias by about 100 Volts, and the gain the PMT grows exponentially with this voltage setting. A pair of terminals on the panel allow you to monitor the potential difference (PMT bias/1000) using an ordinary digital multimeter.] As you raise the bias voltage, you will see occasional

transient electronic noise, but the oscilloscope display should be quiet at each steady voltage. Somewhere around a setting of 4 or 5 turns of the dial, you should get occasional positive-going pulses on the oscilloscope, occurring at a modest rate of 1–10 per second.

- If you see some sinusoidal modulation of a few mV amplitude, and of about 200 kHz frequency, in the baseline of the PMT signal, this is normal.
- If you see a continuing high rate (> 10 kHz) of pulses from the PMT, this is not normal, and you should turn down, or off, the bias level and start fresh – you may have a malfunction, or a light leak.)

If you see this low rate of pulses, you have discovered the *dark rate* of the PMT, its output pulse rate even in the total absence of light. You also now have the PMT ready to look at photons from your two-slit apparatus, so finally you may open the shutter (by grasping that vertically-emerging rod and lifting it a few cm). The oscilloscope should now show a much greater rate of pulses, perhaps of order 10^3 per second, and that rate should vary systematically with the setting of the bulb intensity.

If you can see those analog pulses, you are ready to look at the digital pulses they trigger. Inside the PMT pulse amplifier is a pulse generator which emits, at the OUTPUT TTL connector, a single pulse, of fixed height and duration, each time the analog pulse you're viewing exceeds an adjustable threshold. Connect the OUTPUT TTL to the second vertical channel of your oscilloscope using a $50\ \Omega$ terminator. Set the vertical scale to 2 V/div. Now start with the discriminator control on your PMT box set to 0, and watch the oscilloscope display the analog, and TTL-level, pulse outputs in dual-trace operation. You should be able to find a discriminator setting, low on the dial, for which the oscilloscope shows one TTL pulse for each of, and for only, those analog pulses which reach about a +50 mV level.

- If your analog pulses are mostly not this high, you can raise the PMT bias by half a turn (50 Volts) to gain more electron multiplication.
- If your TTL pulses come much more frequently than the analog pulses, set the discriminator dial lower on its scale.

Now send the TTL pulses to a counter, arranged to display successive readings of the number of pulses that occur in successive 1-second time intervals. These represent photon-induced events (which are converted in the PMT to photoelectrons, then multiplied to electron pulses, then amplified to voltage pulses meeting a threshold criterion). *Because of the $< 100\%$ efficiency of the PMT, not every photon is detected; but you are seeing a representative subset of events due to the arrivals of individual photons.*

To confirm that this is true, record a series of *dark counts* obtained with the light bulb dialed all the way down to 0 on its scale. Now choose a setting that gives an adequate photon count rate (about 10^3 /second) and use the slit-blocker, according to your previously obtained settings, to block the light from both slits. This should reduce the count rate to a background rate, probably somewhat higher than the dark rate. Next,

open up both slits, and try moving the detector slit to see if you can see interference fringes in the photon count rate. You will need to pick a detector-slit location, then wait for a second or more, then read the photon count in one or more 1-second intervals, before trying a new detector-slit location. If you can see maxima and minima, you are ready to take data.

A first sort of data will enable you to be assured that you have set the PMT bias voltage to a suitable range. For this data, the independent variable is the PMT voltage setting, or one of its surrogates, the 10-turn dial setting (1.00 turn = 100 V of bias) or the monitor voltage (0.10 V = 100 V of bias). The dependent variables are both count rates as read by the TTL counter: one is the count rate at the central maximum of the interference pattern, with both slits open; the other is the dark rate, the count rate obtained under identical conditions but with the PMT shutter closed. Try recording both rates over the range 300 – 650 V of PMT bias, and plot the two count rates on a semi-logarithmic graph. You should see the light rate reach a plateau, with the interpretation that you have reached a PMT bias which suffices to allow each genuine photoelectron to trigger the whole chain of electronics all the way to the TTL counter; you should also see the (much lower) dark rate also rising with PMT bias. See if your graph motivates a setting at which you are counting substantially all of the honest-to-goodness photon events, but minimizing the number of dark events.

Hereafter you may leave the PMT bias fixed, and you may set the bulb intensity to yield some convenient count rate ($10^3 - 10^4$ events/second) at the central maximum. You will now note that you seem to be getting an *unstable* count rate, with apparent fluctuations – this is in contrast to the rock-steady signal that you previously got in the quantitative mode of operation observed previously. In fact, the fluctuations you are seeing are a necessary consequence of the independent arrival of photons, and the fluctuations are expected to obey a statistical law. So with everything else stable, take a dozen or more successive readings of the number of counts in 1-second interval; call these samples N_i . Consider a set of n measurements $\{N_i\}$. The expectation value for the *true physical value* will be the mean value, N_{mean} . The standard deviation, σ_{n-1} , which measures the typical deviation from the mean, should obey Poisson statistics, where $\sigma_{n-1} = N_{\text{mean}}^{1/2}$. Compute these numbers for your list, and test this prediction of the *photon statistics*.

By changing the light bulb's intensity, you can easily arrange to make N_{mean} vary from order 10^2 to 10^4 for 1-second count intervals; if you take a set of measurements $\{N_i\}$ for each of several bulb settings, you'll be able to test whether σ_{n-1} does behave as $N_{\text{mean}}^{1/2}$. You might try a log-log plot of σ_{n-1} as a function of N_{mean} and look for a power-law fit of the form $y = x^{1/2}$; note there are no free parameters in this prediction from photon statistics.

With this sort of evidence that you are detecting the arrivals of independent individual photons at your detector, you are now ready to take just the same sort of data as in part III.B. of this manual, except that now the dependent variable is photon count rate. You'll record this as a function of detector-slit position, making records for three cases: one slit,

the other slit, and both slits open. You'll also need to make occasional readings with both slits blocked, to establish a *background rate* of photons reaching the detector but not passing through either of the two slits.

Your graphs will again have a calibrated horizontal axis of known scale; this time their vertical axis will be in absolute units, of photon events detected per unit time. It is perfectly true that the PMT is not 100% efficient; in fact, for green light, it produces countable electronic events for only about 4% of the incident photons. Now take your central-maximum reading for the event rate, and use this efficiency estimate to infer the rate of arrival of all photons at the photocathode of the PMT. From this arrival rate, compute the average time interval between the arrivals of successive photons. Next, compute the *time of flight* of a photon in the apparatus, the time it takes to traverse the distance from the bulb source to the detector. Comparing the time of flight of a given photon to the typical time interval between the arrival of successive photons, you can compute the fraction of the time there is even one photon in flight through the apparatus. What is this fraction for your data? Since you've shown the photon arrivals follow the statistics of independent events, the probability that there are ever *two* relevant photons in the box is truly tiny, of the order of the *square* of this probability. This is the sense in which this apparatus allows you to work *one photon at a time*; and yet even at this low rate of photons, you can see phenomena that you attribute to constructive and destructive interference. The agonizing question is – interference of what, with what?

When you have graphed your data, you'll be able to see a number of new phenomena. First of all, from the spacing of the interference maxima, you'll be able to see immediately that the light in question has a different wavelength than the red laser light you used previously. Second, you'll be able to find detector locations at which you can raise the photon arrival rate by closing one of the slits. Here you'll be in the closest possible contact with the central question of quantum mechanics: how can light, which so clearly propagates as a wave that we can (with this very apparatus) measure its wavelength, also be detected as individual photon events? Or alternatively, how can individual photons in flight through this apparatus nevertheless *know* whether one, or both, slits are open, in the sense of giving photon arrival rates which decrease when a second slit is opened? A good oral dialog between advocates of *wave-nature* and of *particle-nature* for light, each using experimental evidence to poke holes in the other's assertions, will do a great deal to teach participants why concepts as slippery as duality have had to be invented.

THEORETICAL MODELLING FOR THE EXPERIMENTS

Fraunhofer Models for the Interference and Diffraction

The simplest models for the data you have now taken are based on the assumptions of Fraunhofer diffraction, and are described in generic textbooks on electromagnetism. Find a reference you like, and now compare the assumptions of the theory with the facts of your experiment:

- Theory assumes light reaches the double-slit assembly as a plane

wave, ie. from a source infinitely far away; your light reaches the double slits from the source slit, about 38 cm away.

- Theory assumes that the double slits have the same width, D , a center-to-center separation, d , and that the slit lengths are much greater than either the width or separation. Thus the problem is essentially two-dimensional. Your apparatus has slits with approximate values $D = 0.085$ mm, $d = 0.353$ mm, and slit lengths of about 10 mm.
- Theory assumes the light is of a definite wavelength λ ; in your experiment, the laser light is monochromatic, with λ somewhere in the range 0.670 ± 0.005 μm , while the filtered bulb light is not monochromatic, but is distributed over a range from about 0.541 – 0.551 μm .
- Theory assumes the light spreads from the double-slit assembly toward a point detector also infinitely far away, so the radiation pattern can be expressed in terms of an angular variable θ , measured in radians away from the central axis of the apparatus. Your detector is not point-like, but is effectively 0.085 mm wide; nor is it infinitely far away, but a distance of about 50 cm downstream. Nevertheless, it is approximately measuring light intensity as a function of an angular variable; make a model that relates the theoretical parameter θ to the detector-slit position you can set with its micrometer.

Now find a reference that gives a predicted intensity distribution analogous to

$$I(\theta) = I_0 \left(\frac{\sin(\Phi/2)}{\Phi/2} \right)^2 \cos^2(\Delta\varphi/2)$$

where $\Phi = kD \sin \theta$ and $\Delta\varphi = kd \sin \theta$. Get some practice graphing this expression, until you understand the role of the two factors in the intensity expression; also, get familiar enough with its derivation that you can understand how the theoretical prediction changes when one, or the other, of the two slits is blocked. The theory doesn't pretend to give the intensity I_0 at central maximum, whether that's in units of photodiode voltage or photon count rate; so you can supply this value to the theory by hand. The theory also can't know what setting of your detector-slit micrometer corresponds to the location of the 2-slit pattern's central maximum, so you can provide this value by hand as well. But with these vertical-scale and horizontal-zero adjustments, you should be able to intercompare theoretical predictions and experimental data for signals as a function of detector-slit position. You could even indulge in a bit of optimization:

- Does your model predict the fringes of the right height? Change the central-intensity parameter.
- Does your model predict fringes centered about the same center of symmetry as your data shows? Change your central-location parameter.
- Does your model display fringes of the right spacing? Change your slit-spacing parameter, or your wavelength parameter.

- Does your model display a fringe envelope of the right character? (Look at the intensities of the maxima other than the central maximum.) Try changing your slit-width parameter.
- Does your model include the effects of a distribution of wavelength values? Test the effect of including a range of wavelengths appropriate to your experiment.
- Does your model include the effect of the widths of the source and detector slits? Try to model at least the nonzero width of the detector slit.

By these tests, you will not only learn how the model's predictions depend on its input parameters, you will also learn how well your data constrain these parameters' values.

You also have taken data for signals obtained with one, or the other, of the two slits closed. What does your Fraunhofer theory predict for these signals? Note that there are now *no free parameters* at all, since your two-slit data have determined all of them; so overlay the predictions on your data, and have a look at one of the deficiencies of a Fraunhofer model.

Fresnel and Other Models for Interference and Diffraction

There are various deficiencies of the simple Fraunhofer models discussed above, some of which show up directly in comparisons with experimental data. This section suggests a more general method for modelling interference and diffraction phenomena, one which is free of some of these limitations; in particular, it does not require the assumption that source and detector are located at *infinity*. Even so, it is still not a fundamental electromagnetic calculation, since it fails to consider the polarization of actual electromagnetic fields.

Advanced electromagnetism textbooks suggest a variety of methods for treating diffraction, some based on Huygen's principle for wave propagation, and some based on Kirchhoff's integral for scalar wave fields. Both of these theoretical approaches reveal the crucial role of phase, and particularly of the phase variation of a wave field with respect to position. So crucial is this role of phase that nearly all the features of light's behavior can be captured in a truly remarkable model described in Feynman's book *QED: the strange theory of light and matter*. In this book, Feynman models the behavior of light by considering *only* its phase variation, taken to depend on positional displacement Δs according to the complex factor

$$e^{2\pi i \Delta s / \lambda}$$

for a monochromatic wave field whose free-space wavelength is λ . More surprisingly, Feynman uses the *sum over paths* approach, and given light source at P, assigns to another location Q a wave disturbance which is just the (complex) sum of the results computed along *any and all* imaginable paths which lead from P to Q. Along each path, the overall phase is taken to be the product of all the phase factors for each infinitesimal part of the

path. Finally, he supposes that a physical observable, such as the probability of detecting a photon at location Q, is given by the absolute square of the complex amplitude computed by the sum-over-paths approach.

Independent of the justification of this model (which emerges naturally from the Feynman picture of quantum field theory), it provides a readily calculated model for interference phenomena, particularly if we make the daring assumptions that we may use an essentially two-dimensional treatment of the apparatus, and that we may consider only those paths which proceed in straight lines from one slit to another. With these approximations, we can locate slits along a longitudinal central axis, with separation D_1 from source-slit plane to double-slit plane, and further separation D_2 from double-slit plane to detector-slit plane. We can also introduce locations P, Q, and R in the apertures of the source, double, and detector slits, respectively, located off the central axis by horizontal displacements x , y , and z . (In the spirit of the 2D approximation, we ignore altogether the vertical coordinates of P, Q, and R.) Then we compute the phase factors for the (assumed straight-line) paths from P to Q, and from Q to R, according to displacements

$$s_1 = [D_1^2 + (x - y)^2]^{1/2} \text{ and } s_2 = [D_2^2 + (y - z)^2]^{1/2}$$

For a path from source-point P to detector-point R, via intermediate point Q in the double-slit plane, we get an overall phase factor

$$e^{2\pi i s_1(x,y)/\lambda} \cdot e^{2\pi i s_2(y,z)/\lambda}$$

and in the spirit of the sum-over-paths approach, we take the overall field disturbance at R to be the sum (here, the integral) over all possible locations for Q. This yields the integral

$$I(x, z) = \int_{y \in S} dy e^{2\pi i s_1(x,y)/\lambda} \cdot e^{2\pi i s_2(y,z)/\lambda}$$

where the domain of integration S is the combination of the two slits: S_1 ranging from $(-d/2 - D/2)$ to $(-d/2 + D/2)$, and S_2 ranging from $(d/2 - D/2)$ to $(d/2 + D/2)$. Finally, the absolute square of this integral is taken to be an observable signal, corresponding to the photon detection rate for point detector at s , given point source at x .

The experimental apparatus in fact has some locations spread from $x = -D/2$ to $x = +D/2$, and detector locations spread from $z = t - D/2$ to $z = t + D/2$ (where t gives the location, relative to the central axis, of the center of the detector slit), so final integrations over x and z will account for the nonzero widths of the source and detector slits.

The computational burden of this approach can be reduced by series expansions of the path lengths s_1 and s_2 , in powers of the (rather small) off-axis distances x , y , and z ; for example,

$$s_1(x, y) = D_1 \left[1 + \left(\frac{x - y}{D_1} \right)^2 \right]^{1/2} = D_1 \left[1 + \frac{1}{2} \left(\frac{x - y}{D_1} \right)^2 \right]$$

Taking only the lowest-order (quadratic) terms in the series reproduces the Fresnel approximation in optics, and can be justified for the parameters appropriate to this experiment.

Exercise: Compute a typical next-order term, and show that even when exponentiated, it yields only a negligible correction.]

This approximation also allows two factors to be removed from the integral above, giving

$$I(x, z) = e^{2\pi i D_1/\lambda} \cdot e^{2\pi i D_2/\lambda} \int_{y \in S} dy e^{2\pi i \frac{(x-y)^2}{2D_1 l}} \cdot e^{2\pi i \frac{(y-z)^2}{2D_2 l}}$$

in which form it is clear that the two leading factors will disappear upon taking the absolute square. The integration remaining is still daunting, but can actually be performed analytically in terms of the complex error function $\text{Erf}(\times)$, allowing a reasonably efficient evaluation of the model's predictions, especially in the context of a symbolic-computation program such as Mathematica.

One notable success of this model is that it can produce predictions for the data obtained with one slit blocked, merely by performing the integration only over the region S_1 (or only over region S_2). Since the model makes no assumption about small angles, or detector at infinity, it gives predictions not limited by the Fraunhofer assumptions, which more successfully match the single-slit diffraction data readily taken with this apparatus. Nevertheless, this more sophisticated model does not change any qualitative conclusions you have established, and it makes only the slightest quantitative differences to your predictions; the need for deep thought about duality for light remains just as urgent after the models for its behavior have been improved.

It is worth noting that this Feynman model is silent about what the fundamental exponential varying in space really corresponds to. There is not even an appeal to the electromagnetic fields supposed to represent light; instead, the foundations of this model lie in quantum field theory, and the spatial variations represent mathematically the variation of these quantum fields. What is actually propagating through the air in your apparatus is a question not immediately addressed by the mathematics of the theory; it might be described as the propagation of a quantum amplitude, whose interpretation is that its absolute square gives the probability of a quantum event. However vague this picture might be physically, it at least has the virtue of showing that something can propagate as a wave disturbance, yet nevertheless predict the probability of observing a photon event. This is perhaps as close to an 'explanation' of wave-particle duality as can be achieved by a translation from the mathematical language of quantum field theory into a more mechanical or visualizable picture.

CONCLUSIONS

You have performed the two-slit experiment in three ways, and are now fully familiar with the apparatus required to perform it, the data which emerge, and some models that explain the data. You might even have

some numerical conclusions, or deduced values for experimental parameters. But apart from this, you have some added insight into light's propagation and detection.

What is the best evidence your experiment provides that light has a wave-like nature?

What is the best evidence your experiment provides that light has a particle-like nature?

What is the concept of duality, and why was it invented?

Does duality satisfy you, or is the behavior of light paradoxical? Or, is it the character of your *description* of light that is paradoxical?

APPENDIX A: HOW TO READ A MICROMETER DRIVE

Two micrometer screws allow precise mechanical adjustments to the positions of the slit-blocker and the detector slit in this apparatus, and you should learn how they work, and how to read their scales. The two micrometers, and the mechanical flexture mounts they drive, are identical.

Each micrometer consists of a very carefully made metric screw thread of pitch exactly 0.50 mm, so that the rotating shaft of the micrometer advances 0.50 mm for each full (clockwise) turn of the screw. The markings on the micrometer are in place so you can

- keep track of the number of turns you've made, and thereby read the position of the shaft's working end to the nearest 0.50 mm; and
- interpolate within a full turn, so that you can finally quote the position of the shaft's working end to the nearest 0.01 mm.

Here's how to read the number of turns. Call the fixed part of the micrometer the *barrel*, and the rotating external part the *drum*. Note that on the barrel there is a printed longitudinal *stem*, along which there are *branches* emerging alternately on either side. On the one side, every fifth mark is labelled with an integer, 0, 5, 10, and so on: these are at 5-mm spacing, and between them are the 1-mm marks. On the other side of the stem are also branches at 1-mm spacing, but these lie halfway between the mm-marks, and form the *half-mm* marks. Now turn the drum until the 0-mark on its circumference lies right along the line of the stem (this marks the 0° point on one of its 360° rotations). Find the last branch exposed to view on the barrel, and use it to read the micrometer to the nearest 0.50 mm.

For example, if the 5-branch and the next branch on the other side of the stem, are the last exposed to view, the micrometer is set to 5.50 mm.

Now from that position, further counter-clockwise rotation of the drum will withdraw the screw, by another 0.50 mm for a full turn. If you rotate by a fraction $N/50$ of a full turn, you'll withdraw the micrometer by $(N/50) \times 0.50$ mm, or $(0.01 \times N)$ mm. The drum's periphery is conveniently printed with 50 marks around the circumference, and every fifth one of these is labelled, so that you can read the integer N directly. This provides the rest of the information you need to read the micrometer to 0.01 mm.

APPENDIX B: QUANTITATIVE DETECTION BY LOCK-IN TECHNIQUES

This appendix describes an alternative method for processing the photodiode signal which you used quantitatively in section III.B. of this manual; at the cost of a bit of electronic complexity, it offers higher sensitivity to the laser-derived light signal, and immunity from certain kinds of light-induced and electronic background signals. The idea is to isolate that part of the signal due to the laser alone, and to discriminate against any other signals, by turning the laser on and off repeatedly, and electronically isolating the difference that this makes. This technique requires an oscillator with TTL-compatible output to modulate the laser, and a lock-in amplifier to process the photodiode signal.

To use this technique, note that the LASER MOD. input on the light-source end of your apparatus accepts a TTL-level input (0 to +4 Volt square-edged waveform), which turns the laser on for TTL level high (or connector left open) and turns the laser off for TTL level low (or connector grounded). You can use any signal generator with TTL-compatible square-wave output to drive this modulation input; you will need modulation rates only in the 40 - 400 Hz range.

At the detector end of the apparatus, you will note there is the photodiode-amplifier output connector, which produces a voltage signal proportional to the optical power incident on the photodiode. (In addition there may be a small d.c. offset in this signal.) The photodiode's intrinsic response time is well under 1 ms, and the amplifier which turns it sub- μ A output current into the voltage you've been using has a time constant of 0.4 ms.

If you do modulate the laser source, then this photodiode-voltage output will also be modulated in an approximate square wave. You can use lock-in detection of this output voltage, synchronously triggering the lock-in amplifier with the same TTL waveform which is modulating the laser, and thereby achieve total immunity from the d.c. offset of the photodiode amplifier, and nearly total immunity from all noise signals in the photodiode output. This will permit operation of the apparatus in the quantitative mode of section II.B. with the cover open, and/or higher-sensitivity detection of weaker signals on the photodiode.

Because of the 0.4-ms time constant in the photodiode voltage signal, you should use a laser modulation frequency below 400 Hz in this lock-in mode. Some lock-in amplifiers have a current-mode input, in which case you can send the raw photodiode current directly to the lock-in; this permits somewhat faster modulation. In either case, you'll want to set the lock-in's time constant to at least 100 ms, so as to average over many on-off modulation cycles of the laser.

Note that the lock-in technique will do nothing to discriminate against scattered laser light that might be reaching the photodiode; nor will it compensate for poor alignment of the apparatus. But given the averaging features and sensitivity of a typical lock-in amplifier, it can markedly improve the signal-to-noise ratio and the dynamic range of your measurements.

INJECTING TEST PULSES INTO THE PMT AMPLIFIER

This appendix describes a way to use artificial stimulation of the pulse amplifier and discriminator in the detection box at the end of the apparatus; it requires an electronic generator and oscilloscope, and offers the chance to understand and calibrate the pulse-processing electronics in a context separate from actual photon events.

In normal operation, the photomultiplier tube turns single-photoelectron events at the photocathode into much-amplified pulses of electrons arriving at the anode; depending on the high-voltage setting of the PMT supply, the gain of the PMT's dynode structure might be 10^5 or more. A charge pulse of this size is sufficient to excite the charge-pulse amplifier, whose output is in turn sufficient to trigger the TTL-output pulse generator in the detection electronics. The goal here is to test these electronics with a source of artificial events, which are more frequent and regular than those due to individually and randomly arriving photons.

The artificial events will be stimulated by a signal generator; you need a square-wave generator capable of delivering sharp-edged transitions (< 20 ns transition time) of amplitude smoothly adjustable down to say 25 mV. You'll also need a coaxial attenuator of $50\ \Omega$ impedance and 20 dB (or more) attenuation, and an adaptor to deliver its output to the SMA-type input connection marked INPUT TEST PULSE on the detector box.

To conduct these tests, ensure that the PMT high-voltage supply is set to 0 V, so that no light-induced pulses will be reaching the amplifier input. Now connect the square-wave generator's output to a BNC cable to bring it to a tee-junction at an oscilloscope's input; use a high-impedance input on the oscilloscope to monitor the voltage swings ΔV of the square wave as it goes by. Send the square waveform onward from the tee by another BNC cable to the coaxial attenuator, and couple that attenuator's output to the SMA test-input connector.

Now suppose you adjust the square-wave generator amplitude to give voltage swings visible at the oscilloscope of ± 40 mV. Since a 20-dB attenuator reduces power by a factor of 100, in a system of $50\ \Omega$ impedance it attenuates both voltage and current by a factor of 10. Thus you can be confident that clean fast transitions of $\Delta V = \pm 4.0$ mV are present at the test-input connector.

These transitions are internally terminated by a $50\ \Omega$ resistor, and then coupled by a 2-pF capacitor to the input of the pulse amplifier. The negative-going voltage edge will thus couple into the amplifier a charge pulse of size $q = C\Delta V = (2 \times 10^{-12}\text{ F})(-4 \times 10^{-3}\text{ V}) = -8 \times 10^{-15}\text{ C}$, which is the charge carried by 5×10^4 electrons. This charge pulse will be your surrogate for the pulse of electrons arriving at the anode of the PMT, and it will excite the pulse amplifier and discriminator in exactly the same way as genuine photon events do. The advantage of using the square-wave generator to stimulate these events is that the test pulses can be frequent (say at rate 10 kHz), regularly spaced (say every 100 μs), and accompanied by a synchronizing pulse (to trigger an oscilloscope or other equipment).

With such excitation of the input of the pulse amplifier, monitor with an oscilloscope the analog pulse at the PHOTOMULTIPLIER OUT BNC connector. You'll be able to see just what response it gives to a fast charge pulse of a quantifiable size; alternatively, you'll be able to say just how much charge must be delivered in a fast pulse to its input in order that its output pulse reaches an amplitude of (say) +50 mV. That, in turn, will tell you how high the gain of the dynode structure of the PMT needs to be, for single-photoelectron events at the photocathode to be able to give 50-mV analog pulses at this monitor output.

Furthermore, with this steady train of (say) 50-mV pulses at the analog output, you will be able to find just what discriminator setting is required for them to trigger reliably the discriminator, and thus produce a train of TTL-level pulses at its OUTPUT TTL BNC connector.

Finally, with all these electronic tests completed, understood, and documented, you can remove the artificial test-pulse excitation, and be assured that the whole train of electronics will respond in the same way to actual electron pulses from the PMT anode as it has to your test pulses.

Magnetic Torque: Experimenting with the magnetic dipole

originally by Douglas LaFountain, edited by S. Penn

Most of us have childhood memories of playing with magnets and being fascinated by their behavior. You were probably taught that this behavior could be explained by north and south "poles" that either attract or repel each other. However, since your introduction to classical electricity and magnetism, you have learned that a more fundamental model has been developed to explain magnetic interactions. This model involves the magnetic dipole, which itself can be modeled as a loop of current. By viewing permanent magnets as being composed of tiny magnetic dipoles, the magnetic field of permanent magnets can be predicted, as well as those magnets' behavior in the presence of external magnetic fields.

You have been taught these concepts from a theoretical standpoint. Now you need to test these theories in the laboratory. The Magnetic Torque instrument has been designed just for this purpose, to cement your knowledge of the magnetic dipole through experiments. With Magnetic Torque, you'll be examining the behavior of a magnetic dipole in both uniform and non-uniform magnetic fields. You will also be taking data in order to calculate the dipole moment of your magnetic dipole. There are five different experiments that can be performed; each of these experiments involves different combinations of mechanics and E&M principles. Your instructor may pick and choose which experiments he or she wishes you to perform, but you might want to survey all five of the experiments anyway. They'll help your understanding of physics, and may even come in handy on a test!

THE INSTRUMENT

In order to use the Magnetic Torque instrument and learn the physics principles involved with it, it is first necessary to understand the instrument itself. This section of the manual provides a description of the various component parts of the Magnetic Torque apparatus. *Study it carefully before beginning your laboratory session.*

The Magnet

What is referred to as *the magnet* is the component of the instrument that houses the two co-axial coils, the air bearing, and the strobe light.

Coils

The coils are composed of copper wire (#18 gauge) that is wound on bobbins. Each coil has 195 turns. The two coils are always connected in series so that the same current flows through each turn. This current is displayed on the analog ammeter. It is important to note that the coils have some resistance, and that resistance is temperature-dependent. If current is allowed to flow through the coils for a long time, or if the current is high (3–4 amps), the coils' temperature begins to rise. You can feel the increase in temperature. As the coils heat up, their resistance rises, and the current subsequently begins to decrease since the power supply is not current-regulated. It is therefore a good idea to turn the current down to zero when the magnet is not being used. Try to avoid using high currents for any appreciable length of time. The instrument is designed to sustain the full output power of the supply without any danger. However, since the maximum current will be decreased as the temperature of the coils increases, you may not be able to obtain the highest fields if you allow the coils to get too hot.

In order to calculate the magnetic field at the center of the apparatus, one needs to perform an integral, since each of the turns has a different radius and a different distance from the center of the pair. Such a calculation may be a bit tedious. Therefore, given below are an equivalent radius and equivalent distance between the two coils, such that the 390 turns can be thought of as two separate current loops.

equivalent radius $\langle R \rangle = 0.109 \text{ m}$

equivalent separation between the coils $\langle d \rangle = 0.138 \text{ m}$

Using the equivalent radius and separation along with the Biot-Savart Law, you should be able to calculate the magnetic field at the center, which is where the magnetic dipole will reside. *It is only the magnetic field in a small region around the midpoint of the coils axis that is important in all of these experiments.*

The value for the magnetic field *gradient* at the center will be needed for one of the experiments. The field gradient can be calculated by differentiating the expression of the magnetic field with respect to z , where z is the axial distance from the center of one of the coils.

To do the calculation, start with the Biot-Savart Law,

$$d\vec{B} = \frac{I d\vec{\ell} \times \hat{r}}{r^2}$$

where I is the total current in the current ring (i.e. the current in a given wire times the number of windings), $\vec{\ell}$ is the vector along the wire, and $\vec{r}$ is the position vector for the point where the field is being calculated. If

we consider only the field along the axis of the wire coil, then the integral simplifies to:

$$B_z = \frac{2\pi R^2 I}{r^3} = \frac{2\pi R^2 I}{(R^2 + z^2)^{3/2}}$$

where z is the distance along the axis from the center of the ring. In your case, at the midpoint between the rings $z = d/2$. Adding the field for the two coils we find that

$$B_z = (1.36 \pm 0.03) \times 10^{-3} I_1 \text{ Tesla}$$

where I_1 is the current in a single wire (ie. what is measured on the ammeter) in units of Amperes.

When you apply the field gradient, the current in one of the coils is reversed, thus changing the uniform field into a gradient field. To find the value of the gradient, differentiate B_z with respect to z . For our set of coils the resulting gradient is:

$$\frac{dB}{dz} = 1.69 \times 10^{-2} I \text{ Tesla/meter}$$

Air bearing

The air bearing is the spherical hollow in the cylindrical brass rod that is supported on the bottom coil form. The bearing has a narrow opening that allows air to be pumped into the spherical hollow. The ball sits in this hollow and floats on a cushion of air. This provides support without significant friction. The air enters the bearing through a vinyl hose that is connected to the building's compressed air system. Make sure not to restrict the air flow by accidentally kinking or compressing the vinyl hose.

Strobe light

The strobe light is located in an insulated housing on top of the upper coil. One can vary the frequency of the strobe flashes by a control located on the power supply front panel. The frequency of these flashes is automatically measured and read out to two significant figures on the power supply's front panel. This data is updated every 10 seconds.

Strobe light frequency display

Displays the frequency of the strobe light in Hz. Below the display is the frequency-adjust. Turning the knob clockwise will increase the frequency. After adjusting the frequency, one needs to wait for the instrument to count up to the actual frequency. Next to the frequency-adjust knob is the on-off switch for the strobe light.

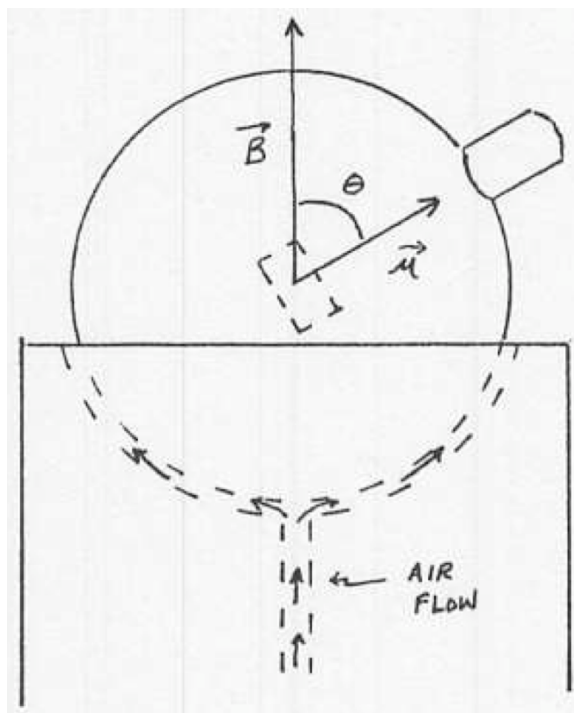


Figure 1 The air bearing provides a low friction rotational mount.

Air switch

Allows you to turn off the air pump when you are performing the magnetic force experiments.

Pilot light

Indicates when the AC power is on for the entire system inside the power supply case.

The Accessories

The accessories are those components of the Magnetic Torque instrument that are not permanently attached to the two main parts of the instrument. They are used to perform the various experiments. Let's examine them.

Cue ball

The cue ball is an aramith snooker ball that has at its center a small cylindrical permanent magnet with a large magnetic dipole moment. The magnetic dipole moment points in the direction of the ball's handle. The

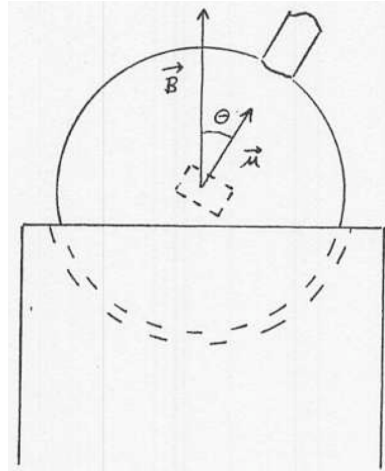


Figure 2 The snooker ball that houses the permanent dipole magnet.

handle allows you to spin the ball, measure its rotational frequency, and determine the direction of the magnetic dipole moment.

The handle on the balls has a small axial hole drilled in it. In this hole, a thin metal rod with an attached weight is placed as part of a static magnetic torque experiment. The aluminum rod has a steel tip at one end that holds fast to the magnetic inside of the ball. The movable weight is a small clear plastic cylinder with an O-ring inside. The O-ring allows the weight to be moved up and down the rod to vary the gravitational torque but prevents it from accidentally slipping once it has been positioned (Figure 3).

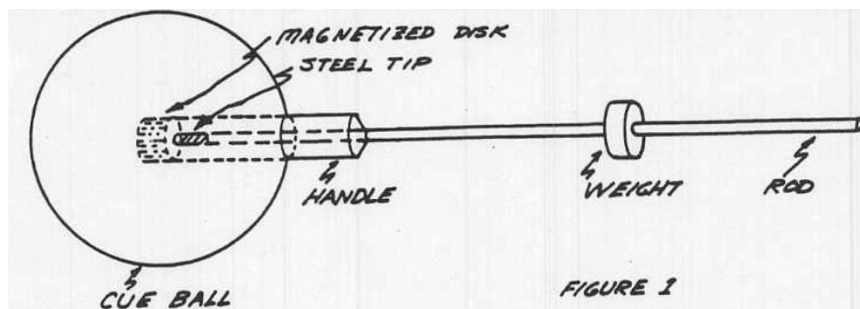


Figure 3 The weight and rod used to generate a variable gravitational torque.

Plastic tower

The clear plastic tube attached to a cylindrical base can be placed on top of the air bearing. This apparatus is used for the magnetic force experiment. A nylon cap placed on the top of the tube holds the rod that supports the spring. The other end of the spring is connected to the magnet. The position of the suspended magnet inside of the tube can be adjusted by moving the rod inside the cap. There is also a small screw-eye at

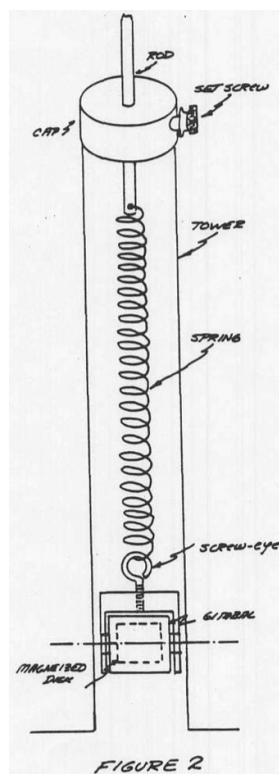


Figure 4 Tower with adjustable spring and weight for magnetic force experiment.

tached above the magnetic dipole that can be used to prevent the magnet from rotating on its gimbals (Figure 4). Ball bearings that weigh one gram each are provided for calibration of the spring.

Rotating magnetic field

The rotating magnetic field is simply a special configuration of permanent magnets and soft iron shims that provide a horizontal magnetic field. This uniform horizontal magnetic field (≈ 1.0 mT) can be manually rotated around the air bearing. It has a hole in its base that allows the air bearing to act as its rotation axis (Figure 5). This magnetic field is used to demonstrate nuclear magnetic resonance.

Bulls-eye level

If the air bearing is not level an additional torque due to unequal air flow can result. Such a torque can produce erroneous data. The bulls-eye level is simply a fluid-filled region that has a bubble in it. When the bubble is within the circle, the apparatus is reasonably level. The leveling can easily be accomplished by turning the screws in the four corners of the base to change the height of the feet. Make this adjustment once at the start of your experiment.

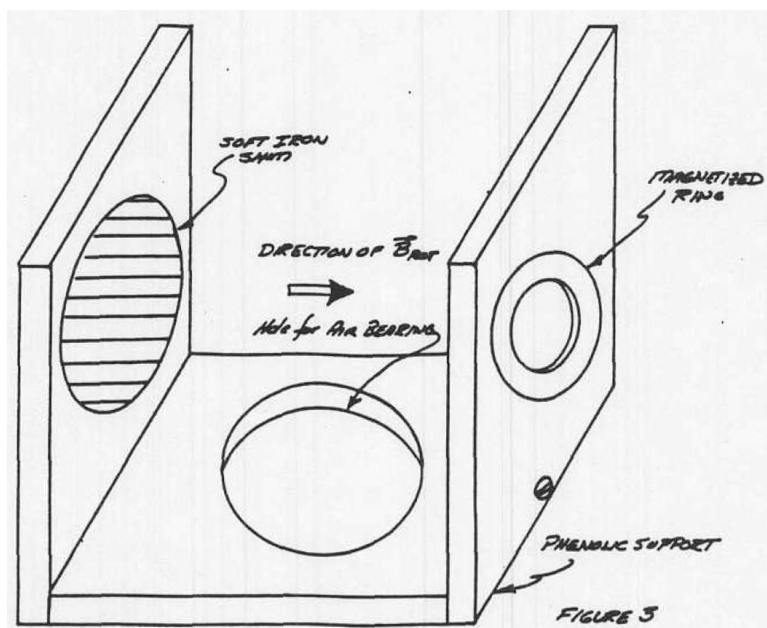


Figure 5 Horizontal magnet used to demonstrate NMR.

Power Supply

What we term the *power supply* contains components other than the power supply. On the front panel, starting from the left, is:

Analog ammeter

Reads the current passing through the coils (since the coils are connected in series). The knob below the meter is used to adjust the current. Turning the knob clockwise increases the current through the coils.

Field direction switch

Controls whether the magnetic field at the center is either up or down.

Field gradient switch

Controls whether there is a magnetic field gradient at the center of the coils, or a uniform magnetic field (gradient off).

On the back of the power supply are the following:

1. On-Off AC power switch for all of the components inside.
2. Power cord that plugs into the AC electrical socket.

3. Cinch Jones connector that connects the power supply to the magnet.
4. Male air hose connector that is no longer used for this experiment.

EXPERIMENT 1: MAGNETIC VS. GRAVITATIONAL TORQUE

Objectives

The main objective of this experiment is to measure the magnetic dipole moment, μ , of the magnet inside the cue ball. You will also verify the functional relationship: $\vec{\mu} \times \vec{B} = \vec{r} \times m\vec{g}$.

Equipment

Magnet, power supply, air bearing, cue ball, aluminum rod with a steel end, weight, ruler, balance, and calipers.

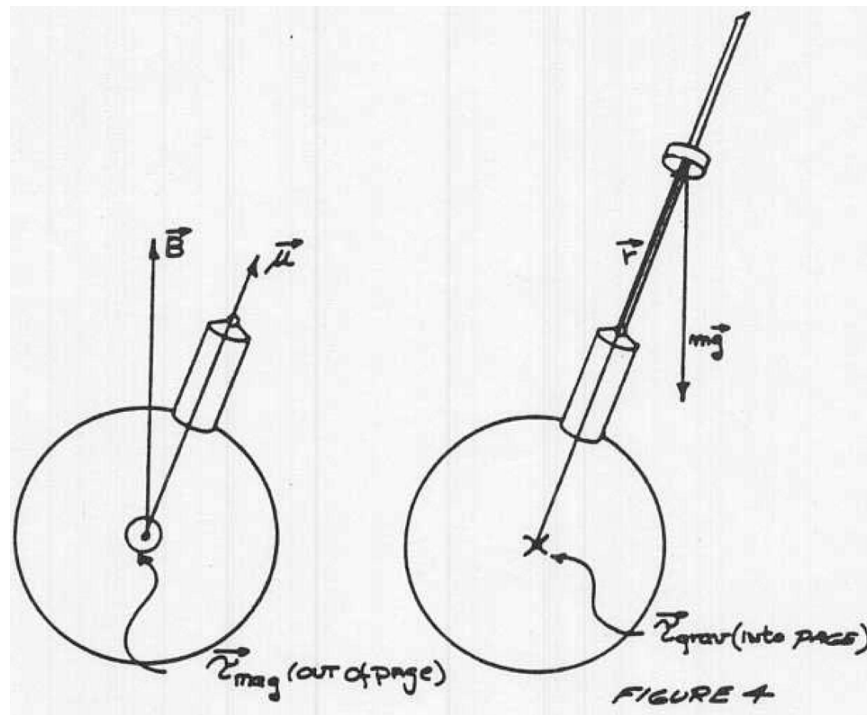


Figure 6 Set-up for the magnetic vs. gravitational torque experiment

Theory

From your electricity and magnetism course, you should be aware that a loop of continuous current will produce a magnetic dipole field and is sometimes referred to simple as a magnetic dipole. Inside the ball is

a neodymium iron boron magnetized disk (0.315 inch diameter by 0.25 inch thick) with the magnetic poles along the axis of the disk. The field from this magnet is almost purely a magnetic dipole field. In a uniform magnetic field (as exists at the center of the two Helmholtz coils), a magnetic dipole experiences a magnetic torque that is given by the expression:

$$\vec{\tau} = \vec{\mu} \times \vec{B}$$

Your magnet's dipole moment is aligned parallel to the handle on the ball; the magnetic field produced by the coils can be either up or down along the coils' axis. If the magnetic field points up, and the magnetic moment is aligned at some angle θ away from the direction of the magnetic field, the ball will experience a torque that will tend to rotate it so that the handle of the ball points upward. But if the aluminum rod is placed in the handle of the ball, there is now another torque due to the earth's gravitational field. The expression for this torque is:

$$\vec{\tau} = \vec{r} \times m\vec{g}$$

The gravitational torque tends to cause the ball to rotate so that the ball's handle points downward (Figure 7). Since a net torque causes a change in angular momentum, the ball will rotate if the gravitational torque is larger than the magnetic torque, or vice versa. But when the magnetic torque is equal to the gravitational torque, the ball will not rotate, since the net torque on the ball is zero. This configuration is mathematically represented by:

$$\begin{aligned} \mu B \sin \theta &= r m g \sin \theta \\ \mu B &= r m g \\ r &= \frac{\mu}{m g} B \end{aligned}$$

So if we measure r for various magnetic fields, the functional dependence of r to B should be a straight line with the slope being an expression that contains μ . In our experiment, B is the magnetic field at the center of the coils and can be calculated from the known current. There are two m 's and two r 's that we need to consider. The first m is the mass of the weight, with its corresponding r being the distance from the center of the ball to the center of the mass of the weight. The second m is the mass of the rod, with its corresponding r being the distance from the center of the ball to the center-of-mass of the rod. But remember that we're trying to measure μ using the slope of a line. *It turns out that the mass of the rod and its center-of-mass distance are combined in a constant in the graph of r versus B and do not affect the slope of the line.* So the only r that one needs to measure is the r of the weight, and the only m that must be measured is the mass of the weight. This is the advantage of a slope measurement of μ rather than a single point determination where the mass and center-of-mass of the rod would be essential.

So μ is the unknown in this experiment. The independent variable is the $\vec{B}$ -field at the center of the instrument, and the dependent variable is r , the displacement of the center of mass of the weight from the center of the ball.

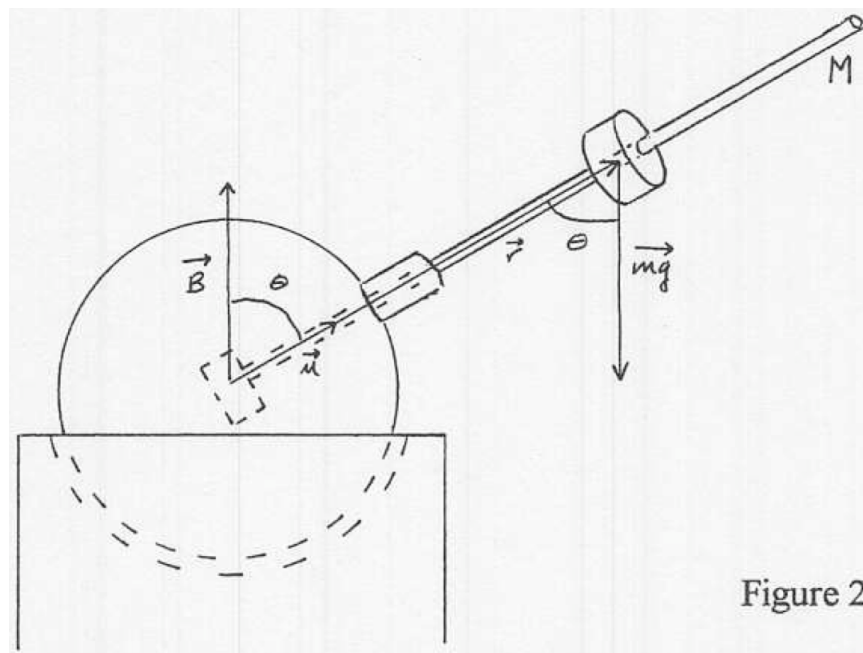


Figure 2

Figure 7 Equilibrium of the magnetic and gravitational torques

Procedure

1. First measure all of the constants that are involved in the experiment. Use a balance to determine the mass of the weight, calipers to measure the diameter of the ball, and a ruler to measure the length of the ball's handle. From these measurements, r of the weight can be determined. Make sure that you remember to keep all of your measurements in SI units, specifically keeping the magnetic field in Teslas, the length measurements in meters, and the mass measurements in kilograms.
2. Next measure r of the weight for various magnetic fields. Turn on the power supply and the air. Keep both the field gradient and the strobe light off. Set the direction of the magnetic field on up so that the handle of the ball points upward when the ball rests on the air bearing. For small currents the gravitational torque is greater than the magnetic torque, so the current to start at is about 2.5 amps. With that current, adjust the position of the weight until the ball and the tip remain stationary at about 90 degrees with respect to the vertical. You might have to steady the rod, weight, and ball with your hand, because the system tends to oscillate and drift due in part to the earth's magnetic field. When the gravitational torque equals the magnetic torque (remember to continuously check the current to make sure that it remains at the set value), take the ball off of the air bearing and turn the current down to zero. Measure the length from the end of the handle to the center of mass of the weight. Add this value to the length from the center of the ball to the end of the handle, and the resulting value is the r of the weight. Repeat these steps for at least six different currents up to 4 amps. The magnetic field can be calculated from the currents. A good way to organize the measurements

is with a data table with column headings of: I (amps), B (teslas), r (meters).

Report

In your report please include:

1. A table of the data,
2. A sample calculation of the magnetic field from the measured current,
3. A graph of r vs. B ,
4. The coefficients (slope and intercept) for the best fit line to the data,
5. The magnitude of the magnetic dipole moment, μ , calculated from the slope, and
6. The rod's center of mass, calculated from the intercept and compared with the measured value.

EXPERIMENT 2: HARMONIC OSCILLATION OF A SPHERICAL PENDULUM

Objectives

The primary objectives of this experiment are to determine the dipole moment of the magnet inside of the ball, and to study the behavior of a physical pendulum's small amplitude oscillation.

Equipment

Magnet, power supply, air bearing, cue ball, stopwatch, calipers, and balance.

Theory

This experiment involves *dynamics* principles. From classical mechanics, you are familiar that a net torque on an object causes a change in that object's angular momentum, given by the expression:

$$\sum \vec{\tau} = \frac{d\vec{L}}{dt}$$

For our particular system, if the cue ball is placed in the air bearing with a uniform magnetic field present, and if the intrinsic dipole moment of the ball is displaced some angle away from the direction of the magnetic field, the ball will experience a net torque and will change its angular

momentum. However, it's important to note the *direction* of the magnetic moment *relative* to the magnetic field. If the magnetic moment in the ball is displaced an angle θ from the axis of the coils (the direction the field), it experiences a restoring torque that acts against the angular displacement of μ . Thus, the differential equation that describes the motion of the ball having moment of inertia I is:

$$-\vec{\mu} \times \vec{B} = I \frac{d^2\vec{\theta}}{dt^2}$$

Where θ is the angular displacement from the direction of $\vec{B}$. The minus sign indicates that the torque is restoring in nature. In scalar form we have:

$$-\mu B \sin \theta = I \frac{d^2\theta}{dt^2}$$

But for small angle displacements, $\sin \theta \approx \theta$, so

$$-\mu B \theta = I \frac{d^2\theta}{dt^2}$$

We'll guess the solution of this equation to be $\theta(t) = A \cos \omega t$, where ω and A are constants. Substituting into the differential equation we have:

$$-\mu B A \cos \omega t = -I A \omega^2 \cos \omega t$$

For this equation to be true for all times t ,

$$\omega^2 = \frac{\mu}{I} B$$

where ω , is the angular frequency of oscillation. If T is the period of oscillation,

$$T = \frac{2\pi}{\omega}$$

Thus, the final expression is:

$$T^2 = \frac{4\pi^2 I}{\mu B}$$

Remember that this expression is only applicable for small angle displacements. The moment of inertia of a uniform solid sphere is given by:

$$I = \frac{2}{5} m r^2$$

where m is the mass of the ball and r is the ball's radius. B is again the independent variable, and T can be measured using a stopwatch. A graph of T vs. $1/B$ should yield a straight line whose slope includes the magnetic moment.

Procedure

1. First, determine the moment of inertia of the cue ball using its mass and its radius. The mass can be determined using a scale, and the radius can be determined using the calipers.
2. For this experiment, the field gradient should be off, the strobe light off, the air on; and the field *up*. Because the magnetic torque is the only torque involved in this experiment, the experiment can be performed at low currents (small B). Place the cue ball on the air bearing and set the current at or near one amp. Give the handle of the ball a small angular displacement from the vertical. Release the ball from rest, and it will oscillate. With a stopwatch measure the amount of time it takes the ball to complete twenty full cycles of motion. Make sure you include the counting of the first full cycle. This measured time divided by twenty will be the period of oscillation for the ball at that particular applied magnetic field. Repeat this for currents up to 4 amps. Because of the $1/B$ independent variable that will be graphed, it's a good idea to obtain data for quite a few different lower currents in order to obtain an even distribution of data points on the graph of T^2 vs. $1/B$. Plot your data in a table with the column headings: I (A), B (T), $t(20)$ (s), $T = t/20$ (s), $1/B$ (T^{-1})

Report

1. Compose a data table. Show sample calculations.
2. Graph the functional relationship of T^2 vs. $1/B$.
3. Calculate the magnitude of your magnetic moment from the slope of this graph.

EXPERIMENT 3: PRECESSION OF A SPINNING SPHERE

Objectives

The main objective is to measure the dipole moment of a permanent magnet inside the cue ball. A secondary objective is to observe and quantify the motion of a spinning sphere subject to an external torque.

Equipment

Magnet, power supply, air bearing, cue ball, strobe light, stopwatch, calipers, and balance.

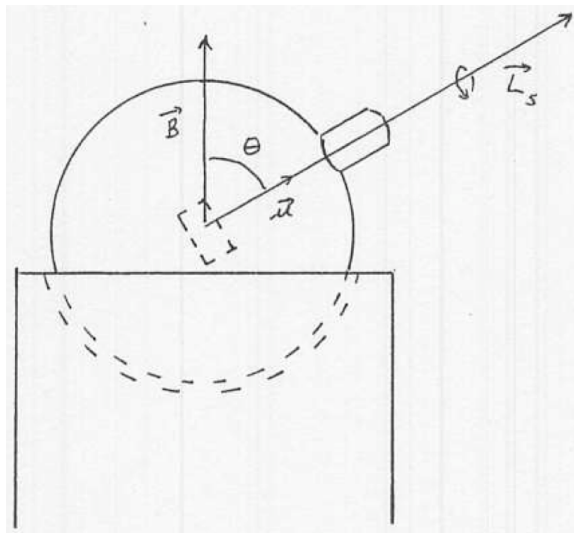


Figure 8 Angular momentum of the spinning ball.

Theory

When the magnetic moment is displaced some angle from the direction of the magnetic field, the magnetic dipole (and subsequently the ball) experiences a torque that causes a change in the ball's angular momentum in the direction of the torque. This is the central principle of this experiment. The ball is displaced from the vertical position and spun, with its spin-axis the axis that runs through the handle of the ball. This creates a large spin angular momentum. The spin axis will remain in a fixed position until the uniform magnetic field is turned on. When the magnetic field is turned on, the magnetic dipole will experience a torque. This torque will cause a change of angular momentum *in the direction of the torque*, but because the ball already has a large spin angular momentum, it will precess. This motion is similar to that of the spinning gyroscope in the earth's gravitational field. You may be familiar with gyroscopic motion from your mechanics course.

The differential equation for the motion of the ball is:

$$\vec{\mu} \times \vec{B} = \frac{d\vec{L}}{dt}$$

A side view of the angular momentum vector is shown below:

If we look from above the spinning ball at the change of the angular momentum for a short time Δt , the picture looks like the one below:

Since $s = r\theta$, we can show that:

$$\begin{aligned} \Delta L_S &= \Delta\theta L_S \sin\theta' \\ \frac{\Delta L}{\Delta t} &= \frac{\Delta\theta}{\Delta t} L \sin\theta' \end{aligned}$$

as $\Delta t \rightarrow 0$,

$$\frac{dL}{dt} = \Omega_p L \sin\theta'$$

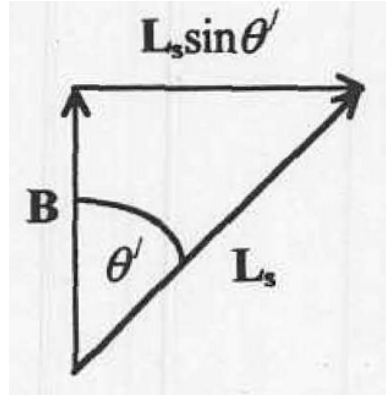


Figure 9 Side view of the angular momentum vector.

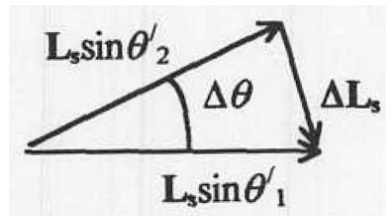


Figure 10 Top View of the change in the angular momentum vector

where Ω_p is the precessional angular velocity. Since μ is in the same direction as $\vec{L}$,

$$\frac{dL}{dt} = \mu B \sin \theta'$$

Thus

$$\begin{aligned} \mu B \sin \theta' &= \Omega_p L \sin \theta' \\ \mu B &= \Omega_p L \\ \Omega_p &= \frac{\mu}{L} B \end{aligned}$$

The precessional frequency (in radians/second) is the dependent variable. It can be determined by measuring the time needed for the handle of the ball to precess through 2π radians, and then dividing that time by 2π . The magnetic field is the independent variable.

The magnitude of the angular momentum L is a constant that can be measured using the strobe light. The handle of the ball has a white dot on its top. As the ball spins the strobe light reflects off of this white dot. When the strobe light is flashing at the same frequency at which the white dot is spinning, the dot will appear stationary. Thus, from the displayed strobe frequency and the measurement of the moment of inertia, the spin angular momentum of the ball at that time can be calculated. The graph of Ω_p vs. B will yield a straight line if L is held constant throughout the experiment. From the slope of this line, μ can be determined.

Procedure

1. First measure the constants. The moment of inertia of the ball is assumed to be that of a solid sphere. The mass of the ball needs to be measured with the scale; its radius, with the calipers. The other constant is the spin angular momentum of the ball. A constant value of the spin angular momentum of the ball can be accomplished by fixing the frequency of the strobe light. We recommend choosing a frequency somewhere between 4.5 and 6 Hertz. Set the strobe light at a frequency in that range and check it throughout the experiment to make sure that it remains constant. In order for the strobe light to illuminate the white dot, the room should be darkened. It is not necessary to make the room completely dark in order to use the strobe light effectively. You can maintain ample light for writing and working with the instrument.
2. Once the strobe light has been set to some particular frequency, record that number. Turn the air on, leave the field gradient off, and set the field direction *up*. Have a stopwatch ready, but don't turn on the current quite yet. Before you begin making measurements, practice spinning the ball. A good technique is to spin the ball (give it a good, hard spin!) and then use the tip of your fingernail to direct it to spin about the handle's axis with the handle in the strobe light. Notice that the frequency of the ball's spin *does change with time*. If you graph this change, it closely approximates an exponential decay. It's because of this decay that the range of frequency between 4.5 and 6 hertz is advised. In this range, the rotational frequency does not change significantly during the time it takes the ball to precess though one period. When you master the spin technique, begin the measurements. Leave the current off, spin the ball up, adjust it so that it's bathed in the strobe light, and then watch the white dot. You'll see it all around the handle at first, but as the ball slows down, the white dot will begin to have a regular rotation of its own. This rotation will slow down until the white dot stops. As soon as it stops, turn the current up to 1 amp, and time a period of the ball's precession. Record this time for that current, turn the current off, and spin the ball up again. Do the same procedure, but this time use 1.5 amps, and continue on in 0.5 amp intervals until you reach 4 amps. This should give you enough data. Record your data in a table with columns: I (A), B (T), T_p (s), $\Omega_p = 2\pi/T$ (Hz).

Report

1. Compose a data table. Show sample calculations.
2. Graph the relationship Ω vs. B .
3. Calculate the magnitude of the your magnetic moment from the slope of this graph.

EXPERIMENT 4: MAGNETIC FIELD GRADIENT FORCE

Objectives

There are three main objectives to this experiment. The first is to demonstrate that in a uniform magnetic field, *there is no net force on a magnetic dipole, only a net torque*. Secondly, this experiment demonstrates that there *is a net force on a magnetic dipole when it is in the presence of a magnetic field gradient*. The third objective is to measure the dipole moment, using the fact that the net force on a magnetic dipole in a magnetic field gradient is proportional to the magnetic moment.

Equipment

Magnet, power supply, plastic tower apparatus, calibrated spring that is supported inside the plastic tower, permanent magnet disk mounted in a gimbal, ball bearings that serve as weights, ruler, and scale.

Theory

The magnetic field of the disk magnet inside the cue ball is well described as a pure magnetic dipole field. A magnetic dipole field is the magnetic field generated by a circular current loop. Consider such a current loop placed in a uniform magnetic field that is aligned along the axis of the loop. The force on any infinitesimal section $d\vec{\ell}$ of the loop is given by:

$$d\vec{F} = I d\vec{\ell} \times \vec{B}$$

where I is the loop current. The same amount of current passes through each infinitesimal section of the current loop. Using the right-hand rule one can see that every $d\vec{F}$ adds to another $d\vec{F}$ of equal magnitude that points in the opposite direction. Therefore, there is no net force on a current loop in a uniform magnetic field. Thus there should be no net force on a magnetic dipole magnet in a uniform magnetic field.

However, if the magnetic disk is placed in a non-uniform magnetic field, that is, a spatial magnetic field gradient, then there *is* a net force on it. To show this, let's assume that the direction of the magnetic field is a direction parallel to the z-axis. This magnetic field changes only in the z-direction, where it changes its magnitude with increasing z. But Maxwell's second equation,

$$\oint_S \vec{B} \cdot \vec{n} da = 0$$

says that the flux of the magnetic field through a closed surface S must equal zero. If a closed surface was constructed around the field along the z-axis, there would be a net flux of the magnetic field through the surface, which would violate Maxwell's second equation. We can thus conclude that if there is a magnetic field that is changing in space, *it must change in more than one direction*. It follows that in our particular situation, the field lines are no longer parallel to the axis of the magnetic dipole, but are bowed, and curve away from the z-axis. Using the expression for the

force on an infinitesimal section of the current loop, it can be shown that there must be a net translational force on the magnetic dipole. It turns out that the expression for the magnitude of this force is:

$$F_z = \mu \frac{dB_z}{dz}$$

We can measure this force. Using the plastic tower apparatus and applying a field gradient, the suspended magnet experiences a translational magnetic force. For our apparatus, this force is nearly constant over a fairly wide range of z -values, since the magnetic field gradient is reasonably constant. However, there is another force on the suspended magnet, the force that the spring applies. The force due to the spring is given by Hooke's Law,

$$F = kz$$

where k is the spring constant, and z is the displacement of the magnet from its equilibrium position (the position where the magnet resides in the absence of a magnetic field gradient). The fact that there are two forces acting on the magnet means that the magnet will displace either upward or downward until the spring force equals the magnetic force. At that point the magnet will cease its acceleration, and if one stops its motion, it will come to rest. At this displacement z , with the forces in equilibrium, we obtain:

$$kz = \mu \frac{dB_z}{dz}$$

The magnetic field gradient can be calculated. First calculate the on-axis magnetic field produced by two coils with currents in the opposite directions. Then differentiate that expression with respect to z , with the axis of the coils being the z -axis. The spring constant k can be measured using the ball bearings as weights to calibrate the spring. The displacement z is then the dependent variable, and dB/dz is the independent variable, proportional to the current. The graph of z vs. dB/dz can be plotted, and from the slope of the expected line, the magnetic moment can be determined.

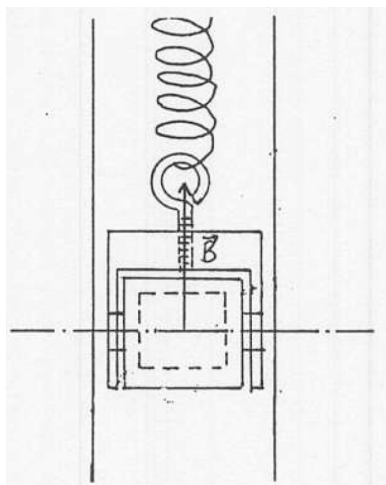


Figure 11 Dipole magnet on a spring used to measure the force in a gradient magnetic field.

Procedure

1. For this experiment, you'll need both the power supply and the magnet coils, but the air bearing is not needed, so the air should be turned off. Place the clear plastic tower on top of the air bearing. Take the cap for the tower, and insert the rod with the spring and suspended magnet into the hole in the cap. There is a screw in the cap that can be used to hold the rod in place. Place the cap on the top of the tower. Now adjust the length of the rod so that the suspended magnet is in the center of the coils, that is, where the center of the ball was when it rested on the air bearing. For the first part of the experiment, loosen the screw on the suspended magnet so the magnet is free to rotate about a horizontal axis. Set the field gradient off, and turn up a current. Now keep your eye on the magnet, and change the field direction. Observe what happens and write it down.

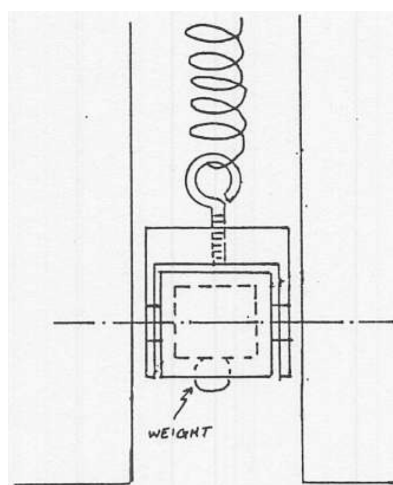


Figure 12 Weights used to measure the spring constant.

2. For the next portion of the experiment, the same equipment will be used, except that now the suspended magnet should be screwed down to prevent it from rotating. Tighten the screw eye so that it is in a position where the flat sides of the magnet are facing up and down.

Determine the spring constant k . To do this the magnetic field must be turned off. Hang the magnet down as far as possible, so that only a small portion of the rod is showing above its support. Note the position of the magnet on the side of the plastic tower using the graduated scale. You may also mark this initial position by placing masking tape on the tower. Measure the length of the rod above the plastic holder. Next, take the cap-rod-spring device out of the tower, and add one ball bearing to the magnet. Then place the cap-rod-spring device onto the tower and adjust the rod so that the magnet is again at the initial position. Measure the length of the rod above the holder. Subtract from this length, the initial length of the rod when only the magnet was in equilibrium, and the resulting length will be the displacement of the magnet from equilibrium due to the weight of one ball bearing. (Measure the mass of the ball bearings. It should be very close to one gram.) Continue these measurements by adding a second, third, and

fourth ball bearing. For each added weight, adjust the rod so that the magnet is in its initial position. Then measure the length of the rod. The rod length minus the initial rod length yields the displacement for that weight. Graph the force on the magnet vs. displacement. The slope of that line will be k .

- Now that the spring constant has been determined, you can proceed with the field gradient experiment. First, use the set screw to lock the dipole in place to prevent it from rotating. With all of the weights off of the magnetic dipole, adjust the position of the rod so that the magnetic dipole is at the center of the coils. Again, note the initial position of the magnet. You may place a piece of masking tape on the side of the plastic tube to make it easier to locate this position. Measure the length of the rod when the magnet is at this center position. With the field gradient switch in the *on* position, slowly turn the current knob to 0.5 amps. The dipole will displace some amount up or down, depending on which direction the field is relative to μ . It is probably best to switch the field direction so that the magnet displaces downward. When the magnet is at the center of the coils, there's more rod inside the tube than there is outside.

Once the magnetic dipole has displaced, return it to the center of the coils by adjusting the support rod. Then measure the length of the rod showing above the tower. The difference between this length and the initial rod length is the displacement of the magnetic dipole due to the magnetic field gradient. Record this displacement z for the current I . Proceed from an initial current of 0.5 amps to 4 amps, in 0.5 amp increments. Again, continuously check the ammeter at high currents to make sure that it is staying at the value that you need for your particular measurement.

The displacement z is the dependent variable for the experiment, while dB/dz is the independent variable that varies with different currents. The spring constant, k is a constant, so μ can be obtained from the slope of the graph of z vs. dB/dz .

Record the data in a table with the columns: I (A), dB/dz (T/m), z (m).

Report

- Write down your observations from part 1 of the procedure. Offer an explanation for these observations.
- Compose a data table for the second part of the experiment, showing sample calculations.
- Graph z vs. dB/dz .
- Calculate the magnitude of the magnetic dipole moment from the slope of this graph.

EXPERIMENT 5: DEMONSTRATING MAGNETIC RESONANCE

Objectives

Magnetic resonance can occur for objects that have a magnetic moment and a colinear angular momentum that are simultaneously subjected to a static magnetic field and a magnetic field in the orthogonal plane rotating at the Larmor precession frequency. Magnetic resonance can only be observed in particles that have an intrinsic angular momentum (spin) and magnetic moment, including protons, neutrons, electrons, and many nuclei.

In this experiment we can create the necessary conditions for magnetic resonance on a macroscopic scale. By spinning the billiard ball along the handle axis, we can give the billiard ball angular momentum that is aligned with its magnetic dipole moment. The Helmholtz coils provide a static magnetic field in the vertical direction. As observed in experiment 3, if the billiard ball has a colinear angular momentum and magnetic moment, then in a static magnetic field the ball will precess at the Larmor frequency about the axis of the static field. Now consider a horizontal magnetic field in the plane normal to the static field. Let this horizontal field also rotate at the Larmor frequency so that it is always perpendicular to the magnetic moment. In this case the horizontal field will generate a torque that will cause the object to rotate in the plane defined by the magnetic moment and the static field. On the classical (macroscopic) scale, we observe the magnetic dipole rotate. However, on the particle or quantum scale, this change in angular momentum causes a *spin flip*.

An group of particle whose magnetic moments are aligned will generate a small magnetic field. When these particles undergo a collective spin flip, they flip the direction of this field, a phenomenon that is quite detectable even for small samples. By measuring the Larmor frequency at which this spin flip occurs, one can determine the particles that are experiencing magnetic resonance. These are the fundamentals of methods used for nuclear magnetic resonance, also known as magnetic resonance imaging.

Theory

The ratio of the magnetic moment, μ , to the angular momentum, L , is called the *gyromagnetic ratio*, $\gamma = \mu/L$. If an object, which has a colinear magnetic moment and angular momentum, is located in a magnetic field, $\vec{B}$, (which may be time dependent) then it experiences a torque:

$$\vec{\tau} = \vec{\mu} \times \vec{B} = \frac{d\vec{L}}{dt} = \frac{1}{\gamma} \frac{d\vec{\mu}}{dt}$$

If $\vec{B}_0$ is a constant field in the z direction, the object will execute precessional motion with the direction of the magnetic moment sweeping out a circle in the x-y plane as shown in Figure 13.

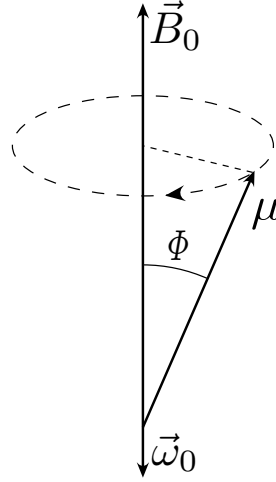


Figure 13 Precession of the angular momentum and magnetic moment in a static $\vec{B}$ field.

Consider a rotating horizontal magnetic field as that produced by the horizontal permanent magnet shown in Figure 5. This unit will produce a magnetic field of about 1.0 mT in the x-y plane perpendicular to the field produced by the coils, $\vec{B}_0$. This device sits around the air bearing and can be rotated by hand by the experimenter at approximately the Larmor precession frequency, ω_0 . If the experimenter is careful to rotate the horizontal field such that the field remains perpendicular to the plane defined by $\vec{\mu}$ and $\vec{B}_0$, then $\vec{\mu}$ will rotate in that plane about the axis of the horizontal field. In other words, Φ will become time dependent, with

$$\Omega = \frac{d\Phi}{dt} = \text{constant}$$

When Φ has changed by 180° then we have produced the macroscopic analog of a spin flip.

These dynamics can also be understood by examining the system in a noninertial rotating reference frame. Imagine moving in a reference frame that is rotating about the z-axis at the Larmor precession frequency. In this frame the ball still has angular momentum, $\vec{L}$, and magnetic moment, $\vec{\mu}$ but it is not moving and Φ is constant. The reason that there is no precession is that there is no $\vec{B}$ field. In the noninertial rotating frame the effective magnetic field is

$$\vec{B}_{\text{effective}} = \vec{B}_0 - \frac{\vec{\omega}}{\gamma} = 0$$

If we now add the rotating horizontal magnetic field, $\vec{B}_{\text{rot}}$, then in the rotating reference frame this field becomes a static field, ie. $\vec{B}_{\text{eff}} = \vec{B}_{\text{rot}}$. Thus in the rotating frame we see a precession about $\vec{B}_{\text{rot}}$ at the precession frequency $\Omega = \gamma B_{\text{rot}}$. It is this precession that results in the *spin flip*.

Procedure

Our goal is to see the classical *spin flip*; there is no data to collect or analyze. However the technique can be a bit challenging and it may take you several attempts to see the effect.

First, you must be able to spin the ball rapidly with as little nutation (wobble) as possible. You will start with the handle of the ball at an angle that is almost horizontal. Be careful not to allow the handle to touch the air bearing. Use the power drill to spin up the ball and practice releasing the ball smoothly so that it does not wobble.

Second, with the power supply off, turn the current on the coils to its maximum setting. Now spin the ball up with the drill and immediately after the release turn on the power supply. Note the direction and frequency of the precession of the ball.

Third, place the horizontal magnetic field device over the air bearing. Practice turning it by hand in the direction and speed of the ball's precession. It is easiest to turn it from above. If the device does not turn smoothly, remove it and clean all the contact surfaces.

Finally, position the horizontal field so that it is perpendicular to the direction of the magnetic moment of the ball. In other words the handle of the ball should be in the gap of the horizontal field. Now spin the ball with the drill. Immediately upon release turn on the power supply (ie. the coil power) and begin to turn the horizontal field. The horizontal field should be rotated so that the magnetic moment remains in a plane normal to the axis of the horizontal field. You should see the magnetic moment precess about the axis of the horizontal field. If you see the ball precess down so that the handle hits the air bearing, rotate the horizontal field 180° and try again. It may take a few tries, but you should be able to produce the effect. And when you do, you will be watching *the world's lowest frequency magnetic resonance experiment*.

Pulsed NMR

originally by ?? & Eric Black, edited by S. Penn

INTRODUCTION AND THEORY

In 1946 nuclear magnetic resonance (NMR) in condensed matter was discovered simultaneously by Edward Purcell at Harvard and Felix Bloch at Stanford using different instrumentation and techniques. Both groups, however, observed the response of magnetic nuclei, placed in a uniform magnetic field, to a continuous (cw) radio frequency magnetic field as the field was tuned through resonance. This discovery opened up a new form of spectroscopy which has become one of the most important tools for physicists, chemists, geologists, and biologists.

In 1950 Erwin Hahn, a young postdoctoral fellow at the University of Illinois, explored the response of magnetic nuclei in condensed matter to pulse bursts of these same radio frequency (rf) magnetic fields. Hahn was interested in observing transient effects on the magnetic nuclei after the rf bursts. During these experiments he observed a spin echo signal; that is, a signal from the magnetic nuclei that occurred after a two pulse sequence at a time equal to the delay time between the two pulses. This discovery, and his brilliant analysis of the experiments, gave birth to a new technique for studying magnetic resonance. This pulse method originally had only a few practitioners, but now it is the method of choice for most laboratories. For the first twenty years after its discovery, continuous wave (cw) magnetic resonance apparatus was used in almost every research chemistry laboratory, and no commercial pulsed NMR instruments were available. However, since 1966 when Ernst and Anderson showed that high resolution NMR spectroscopy can be achieved using Fourier transforms of the transient response, and cheap fast computers made this calculation practical, pulsed NMR has become the dominant commercial instrumentation for most research applications.

This technology has also found its way into medicine. MRI (magnetic resonance imaging; the word "nuclear" being removed to relieve the fears of the scientifically illiterate public) scans are revolutionizing radiology. This imaging technique seems to be completely noninvasive, produces remarkable three dimensional images, and has the potential to give physicians detailed information about the inner working of living systems. For example, preliminary work has already shown that blood flow patterns in both the brain and the heart can be studied without dangerous catheterization or the injection of radioactive isotopes. Someday,

MRI scans may be able to pinpoint malignant tissue without biopsies. MRI is in its infancy, and we will see many more applications of this diagnostic tool in the coming years.

You have purchased the first pulsed NMR spectrometer designed specifically for teaching. The PS1-A is a complete spectrometer, including the magnet, the pulse generator, the oscillator, pulse amplifier, sensitive receiver, linear detector, and sample probe. You need only supply the oscilloscope and the substances you wish to study. Now you are ready to learn the fundamentals of pulsed nuclear magnetic resonance spectroscopy.

Nuclear magnetic resonance is a vast subject. Tens of thousands of research papers and hundreds of books have been published on NMR. We will not attempt to explain or even to summarize this literature. Some of you may only wish to do a few simple experiments with the apparatus and achieve a basic conceptual understanding, while others may aim to understand the details of the density matrix formulation of relaxation processes and do some original research. The likelihood is that the majority of students will work somewhere in between these two extremes. In this section we will provide a brief theoretical introduction to many important ideas of PNMR. This will help you get started and can be referred to later. These remarks will be brief, not completely worked out from first principles, and not intended as a substitute for a careful study of the literature and published texts. An extensive annotated bibliography of important papers and books on the subject is provided at the end of this section.

Magnetic resonance is observed in systems where the magnetic constituents have **both a magnetic moment and an angular momentum**. Many, but not all, of the stable nuclei of ordinary matter have this property. In classical physics terms, magnetic nuclei act like a small spinning bar magnet. For this instrument, we will only be concerned with one nucleus, the nucleus of hydrogen, which is a single proton. The proton can be thought of as a small spinning bar magnet with a magnetic moment $\vec{\mu}$ and an angular momentum $\vec{J}$, which are related by the vector equation:

$$\vec{\mu} = \gamma \vec{J} \quad (1)$$

where γ is called the *gyromagnetic ratio*. The nuclear angular momentum is quantized in units of $\hbar$ as:

$$\vec{J} = \hbar \vec{S} \quad (2)$$

where $\vec{S}$ is the *spin* of the nucleus. The magnetic energy U of the nucleus in an external magnetic field is

$$U = -\vec{\mu} \cdot \vec{B} \quad (3)$$

If the magnetic field is in the z-direction, then the magnetic energy is

$$U = -\mu_z B_0 = -\gamma \hbar S_z B_0 \quad (4)$$

Quantum mechanics requires that the allowed values of S_z , m_s , be quantized as

$$m_s = S, S - 1, S - 2, \dots - S \quad (5)$$

For the proton, with spin one half ($S = 1/2$), the allowed values of S_z , are simply

$$m_s = \pm \frac{1}{2} \quad (6)$$

which means there are only two magnetic energy states for a proton

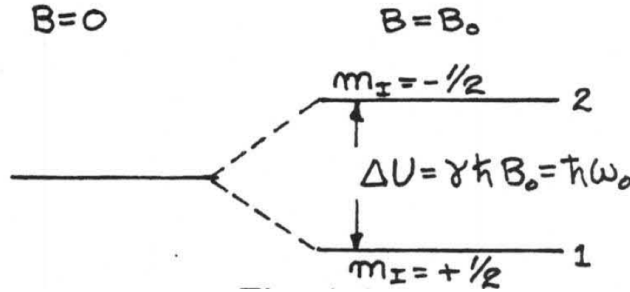


Figure 1 Splitting of the proton energy state in a static magnetic field.

residing in a constant magnetic field B_0 . These are shown in figure 1. The energy separation between the two states ΔU can be written as

$$\Delta U = \hbar\omega_0 = \gamma\hbar B_0 \quad (7)$$

$$\omega_0 = \gamma B_0 \quad (8)$$

This is the fundamental resonance condition. For the proton

$$\gamma_{\text{proton}} = 2.675 \times 10^4 \frac{\text{rad/s}}{\text{gauss}} \quad (9)$$

Thus the resonant frequency of the proton is proportional to the magnetic field

$$f_0 (\text{MHz}) = 4.258 B_0 (\text{kilogauss}) \quad (10)$$

This scale factor will be used frequently and is a number worth remembering.

If a one milliliter (ml) sample of water (containing about 7×10^{19} protons) is placed in a magnetic field in the z-direction, a nuclear magnetization in the z-direction eventually becomes established. This nuclear magnetization occurs because of unequal population of the two possible quantum states. If N_1 and N_2 are the number of spins per unit volume in the respective states (see Fig. 1), then the population ratio (N_2/N_1), in thermal equilibrium, is given by the Boltzmann factor as

$$\left[\frac{N_2}{N_1} \right] = e^{-\frac{\Delta U}{k_B T}} = e^{-\frac{\hbar\omega_0}{k_B T}} \quad (11)$$

and the magnetization is

$$\vec{M}_z = (N_1 - N_2) \mu \quad (12)$$

The thermal equilibrium magnetization per unit volume for N magnetic moments is

$$M_0 = N\mu \tanh\left(\frac{\mu B}{k_B T}\right) \approx N \frac{\mu^2 B}{k_B T} \quad (13)$$

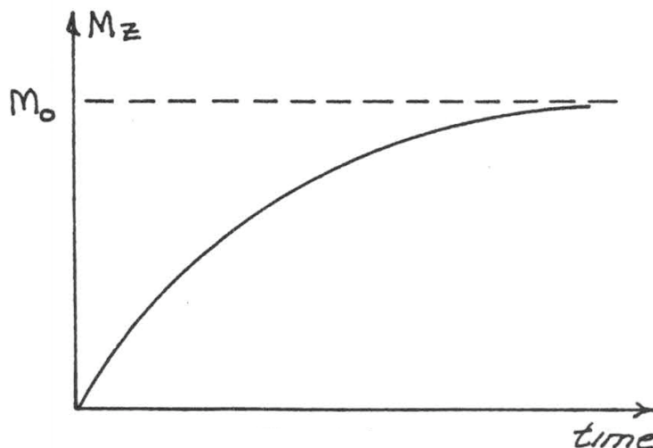


Figure 2 Increase in magnetization after establishing a magnetic field.

where $N = N_1 + N_2$. This magnetization does **not** appear instantaneously when the sample is placed in the magnetic field. It takes a finite time for the magnetization to build up to its equilibrium value along the direction of the magnetic field (which we define as the z-axis). For most systems, the z-component of the magnetization is observed to grow exponentially as depicted in Figure 2. The differential equation that describes such a process assumes the rate of approach to equilibrium is proportional to the separation from equilibrium :

$$\frac{dM_z}{dt} = \frac{M_0 - M_z}{T_1} \quad (14)$$

where T_1 , is called the spin-lattice relaxation time. If the unmagnetized sample is placed in a magnetic field, so that at $t = 0$, $M_z = 0$, then direct integration of Eqn. 14, with these initial conditions, gives

$$M_z(t) = M_0 \left(1 - e^{-t/T_1} \right) \quad (15)$$

The rate at which the magnetization approaches its thermal equilibrium value is characteristic of the particular sample. Typical values range from microseconds to seconds. What makes one material take $10 \mu\text{s}$ to reach equilibrium while another material (also with protons as the nuclear magnets) takes 3 seconds? Obviously, some processes in the material make the protons *relax* towards equilibrium at different rates. The study of these processes is one of the major topics in magnetic resonance.

Although we will not attempt to discuss these processes in detail, a few ideas are worth noting. In thermal equilibrium more protons are in the lower energy state than the upper. When the unmagnetized sample was first put in the magnet, the protons occupied the two states equally, that is $N_1 = N_2$. During the magnetization process, energy must flow from the nuclei to the surroundings, since the magnetic energy from the spins is reduced. The surroundings which absorb this energy is referred to as *the lattice*, even for liquids or gases. Thus, the name *spin-lattice* relaxation time for the characteristic time of this energy flow.

However, there is more than energy flow that occurs in this process of magnetization. Each proton has angular momentum (as well as a magnetic moment) and the angular momentum must also be transferred from the spins to the lattice during magnetization. In quantum mechanical terms, the lattice must have angular momentum states available when a spin goes from $m_S = -1/2$ to $m_S = +1/2$. In classical physics terms, the spins must experience a torque capable of changing their angular momentum. The existence of such states is usually the critical determining factor in explaining the enormous differences in T_1 for various materials. Pulsed NMR is ideally suited for making precise measurements of this important relaxation time. The pulse technique gives a direct unambiguous measurement, whereas cw spectrometers use a difficult, indirect, and imprecise technique to measure the same quantity.

What about magnetization in the x-y plane? In thermal equilibrium the only net magnetization of the sample is M_z , the magnetization along the external constant magnetic field. This can be understood from a simple classical model of the system. Think of placing a collection of tiny current loops in a magnetic field. The torque on the loop is given by:

$$\vec{\tau} = \vec{\mu} \times \vec{B} \quad (16)$$

$$\vec{\tau} = \frac{d\vec{J}}{dt} \quad (17)$$

$$\vec{\mu} \times \vec{B} = \frac{1}{\gamma} \frac{d\vec{\mu}}{dt} \quad (18)$$

Eqn. 18 is the classical equation describing the time variation of the mag-

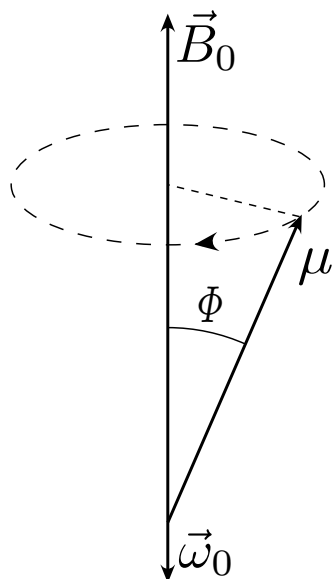


Figure 3 Precession of the colinear magnetic moment and angular momentum in a static magnetic field.

netic moment of the proton in a magnetic field. It can be shown from Eqn. 18 that the magnetic moment will execute precessional motion, depicted in Figure 3. The precessional frequency $\omega_0 = \gamma B_0$ is just the resonant frequency in Eqn. 8.

If we add up all the magnetization for the 10^{20} protons in our sample in thermal equilibrium, the μ_z components sum to M_z , but the x and y components of the individual magnetic moments add to zero. For the x-components of every proton to add up to some M_x , there must be a definite phase relationship among all the precessing spins. For example, we might start the precessional motion with the x-component of the spins lined up along the x-axis. But that is not the case for a sample simply placed in a magnet. In thermal equilibrium the spin components in the x-y plane are randomly positioned. Thus, in thermal equilibrium there is no transverse (x and y) component of the net magnetization of the sample. However, as we shall soon see, there is a way to **create** such a transverse magnetization using radio frequency pulsed magnetic fields. The idea is to rotate the thermal equilibrium magnetization M_z into the x-y plane and thus create a temporary M_x and M_y . Let's see how this is done.

Equation 18 can be generalized to describe the classical motion of the net magnetization of the entire sample. It is

$$\frac{d\vec{M}}{dt} = \gamma \vec{M} \times \vec{B} \quad (19)$$

where $\vec{B}$ is any magnetic field including time dependent rotating fields. Suppose we apply not only a constant magnetic field $B_0 \hat{k}$, but a rotating (circularly polarized) magnetic field of frequency ω in the x-y plane so the total field is written as

$$\vec{B}(t) = B_1 \cos(\omega t) \hat{i} + B_1 \sin(\omega t) \hat{j} + B_0 \hat{k} \quad (20)$$

The analysis of the magnetization in this complicated time dependent magnetic field can best be carried out in a noninertial rotating coordinate system. The coordinate system of choice is rotating at the same angular frequency as the rotating magnetic field with its axis in the direction of the static magnetic field. In this rotating coordinate system (denoted by $\hat{x}^*$, $\hat{y}^*$, $\hat{z}^*$), the rotating magnetic field appears to be stationary and aligned along the $\hat{x}^*$ (Fig. 4). However, from the point of view of the rotating coordinate system, B_0 , and B_1 , are not the only magnetic fields. An **effective field** along the $\hat{z}^*$, of magnitude $\frac{-\omega}{\gamma} \hat{k}$ must also be included. Let's justify this new effective magnetic field with the following physical argument.

Equations 18 and 19 predict the precessional motion of a magnetization (= net magnetic dipole moment) in a constant magnetic field $B_0 \hat{k}$. Suppose one observes this precessional motion from a rotating coordinate system which rotates at the precessional frequency. In this frame of reference the magnetization appears stationary, in some fixed position. The only way a magnetization can remain fixed in space is if there is no torque on it. If the magnetic field is zero in the reference frame, then the torque on $\vec{M}$ is always zero no matter what direction $\vec{M}$ is oriented. The magnetic field is zero (in the rotating frame) if we add the effective field $\frac{-\omega}{\gamma} \hat{k}$ which is equal to $B_0 \hat{k}^*$.

Transforming the magnetic field expression in equation (20) into such a rotating coordinate system, the total magnet field in the rotating

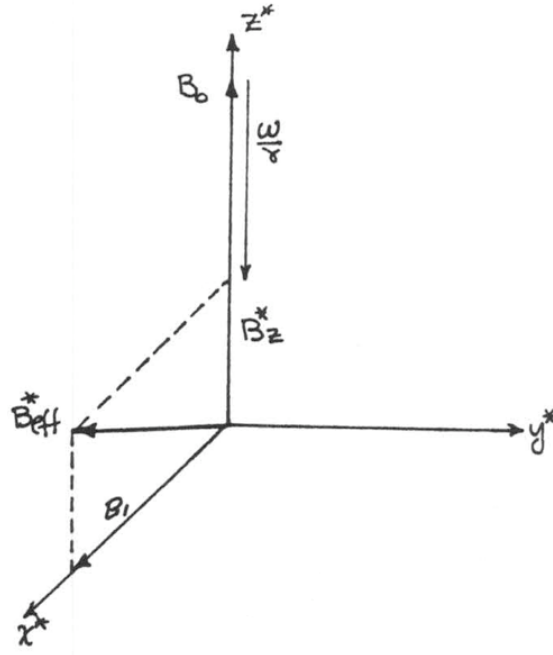


Figure 4 The effective B field in a rotating frame.

frame B^* is

$$\vec{B}_{\text{eff}}^* = B_1 \hat{i}^* + \left(B_0 - \frac{\omega}{\gamma} \right) \hat{k}^* \quad (21)$$

shown in Figure 4. The classical equation of motion of the magnetization as observed in the rotating frame is then

$$\frac{d\vec{M}}{dt} = \gamma \vec{M} \times \vec{B}_{\text{eff}}^* \quad (22)$$

which shows that $\vec{M}$ will precess about $\vec{B}_{\text{eff}}^*$ in the rotating frame.

Suppose now, we create a rotating magnetic field at a frequency ω_0 such that

$$\omega = \omega_0 = \gamma B_0 \quad (23)$$

In that case, $\vec{B}_{\text{eff}}^* = B_1 \hat{i}^*$, a constant magnetic field in the x^* -direction. Then the magnetization M_z , begins to precess about this magnetic field at a rate $\Omega = \gamma B_1$ (in the rotating frame). If we **turn off** the B_1 field at the instant the magnetization reaches the x - y plane, we will have created a transient (non-thermal equilibrium) situation where there **is** a net magnetization in the x - y plane. If this rotating field is applied for twice the time the transient magnetization will be $-\vec{M}_z$, and if it is left on four times as long the magnetization will be back where it started, with $\vec{M}_z$, along the z^* -axis. These are called:

- 90° or $\pi/2$ pulse: $M_z \rightarrow M_y$
- 180° or π pulse: $M_z \rightarrow -M_z$

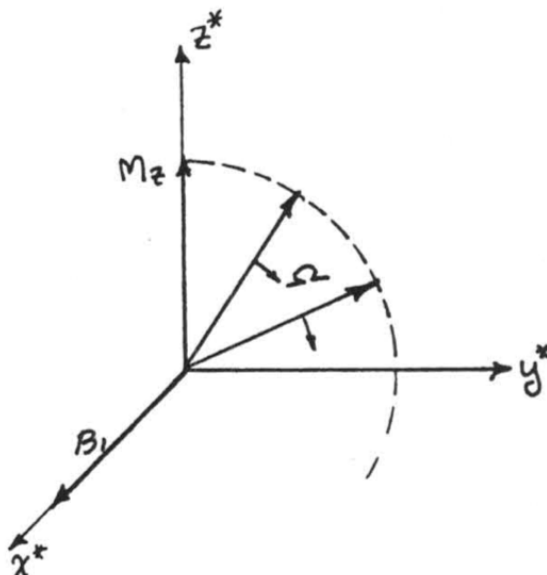


Figure 5 Precession about the horizontal B field as seen in the rotating frame.

- 360° or 2π pulse: $M_z \rightarrow M_z$

In the laboratory (or rest) frame where the experiment is actually carried out, the magnetization not only precesses about B_1 but rotates about $\hat{k}$ during the pulse. It is not possible, however, to observe the magnetization **during** the pulse. Pulsed NMR signals are observed **AFTER THE TRANSMITTER PULSE IS OVER**. But, what is there to observe **AFTER** the transmitter pulse is over? The spectrometer detects the **net magnetization precessing about the constant magnetic field $B_0\hat{k}$ in the x-y plane. Nothing Else!**

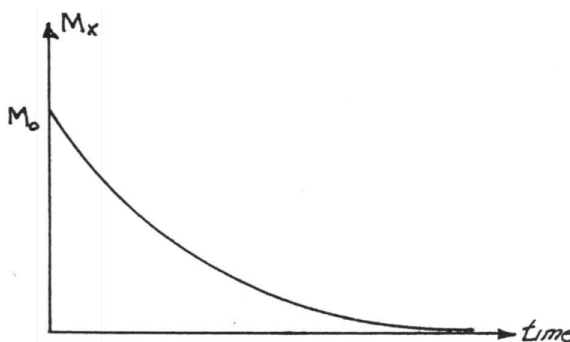


Figure 6 The decay of magnetization as the proton's return to thermal equilibrium.

Suppose a 90° ($\pi/2$) pulse is imposed on a sample in thermal equilibrium. The net equilibrium magnetization will be rotated into the x-y plane where it will precess about $B_0\hat{k}$. But the x-y magnetization will not last forever. For most systems, this magnetization decays exponen-

tially as shown in Figure 6. The differential equations which describe the decay in the rotating coordinate system are:

$$\frac{dM_{x^*}}{dt} = -\frac{M_{x^*}}{T_2} \quad \text{and} \quad \frac{dM_{y^*}}{dt} = -\frac{M_{y^*}}{T_2} \quad (24)$$

whose solutions are

$$M_{x^*}(t) = M_0 e^{-t/T_2} \quad \text{and} \quad M_{y^*}(t) = M_0 e^{-t/T_2} \quad (25)$$

where the characteristic decay time T_2 is called the Spin-Spin Relaxation Time. One simple way to understand this relaxation process from the classical perspective, is to recall that each proton is itself a magnet and produces a magnetic field that is felt by its neighbors. Therefore for a given distribution of these protons there must also be a distribution of local fields at the various proton sites. Thus, the protons precess about $B_0 \hat{k}$ with a distribution of frequencies, not a single frequency ω_0 . Even if all the protons begin in phase (after the 90° pulse) they will soon get out of phase and the net x-y magnetization will eventually go to zero. A measurement of T_2 , the decay constant of the x-y magnetization, gives information about the distribution of local fields at the nuclear sites.

From this analysis it would appear that the spin-spin relaxation time T_2 can simply be determined by plotting the decay of M_x (or M_y) after a 90° pulse. This signal is called the **free precession or free induction decay** (FID). If the magnet's field were perfectly uniform over the entire sample volume, then the time constant associated with the free induction decay would be T_2 . But in most cases it is the **magnet's nonuniformity** that is responsible for the observed decay constant of the FID. The magnet in our apparatus has a small region of sufficient uniformity, a *sweet spot*, to produce at least a .3 millisecond decay time. Thus, for a sample whose $T_2 < 0.3$ ms the free induction decay constant is also the T_2 of the sample. But what if T_2 is actually 0.4 msec or longer? The observed decay will still be about 0.3 ms. We won't be observing T_2 but the decay due to the nonuniformity of our magnet. Here is where the genius of Erwin Hahn's discovery of the spin echo plays its crucial role.

Before the invention of pulsed NMR, the only way to measure the real T_2 , was to improve the magnets homogeneity and make the sample smaller. But, Pulsed NMR (PNMR) changed this. Suppose we use a two pulse sequence, the first one 90° and the second one, turned on a time τ later, a 180° pulse. What happens? Figure 7 shows pulse sequence and Figure 8 shows the progression of the magnetization in the rotating frame.

Study these diagrams carefully. The 180° pulse allows the x-y magnetization to rephase to the value it would have had with a perfect magnet. This is analogous to an egalitarian foot race for the kindergarten class; the race that makes everyone in the class a winner. Suppose you made the following rules. Each kid would run in a straight line as fast as he or she could and when the teacher blows the whistle, every child would turn around and run back to the finish line, again as fast as he or she can run. The faster runners go farther, but must return a greater distance and the slower ones go less distance, but all reach the finish line at the same time. The 180° pulse is like that whistle. During the first interval τ the spins in the larger field get out of phase by $+\Delta\theta$ with respect to

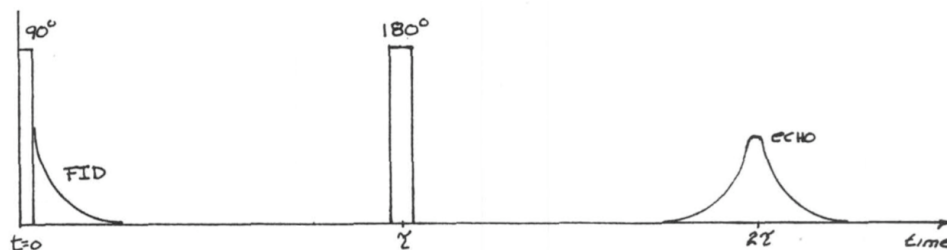


Figure 7 Sequence of pulses of observed magnetization using spin echo to detect T_2

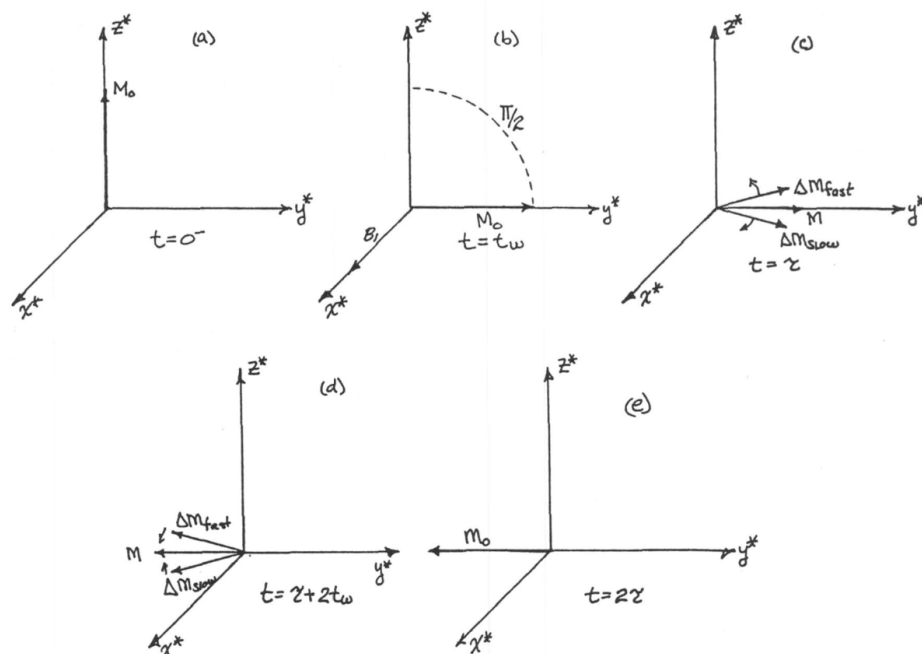


Figure 8 a) Thermal equilibrium magnetization along the z axis before the 90° radio frequency (rf) pulse. b) M_0 rotated to the y -axis after the 90° pulse. c) The magnetization in the x - y plane is decreasing because some spins Δm_{fast} are in a higher field (ω faster), and some Δm_{slow} a lower field (ω slower). d) Spins are rotated 180° (flip the entire x - y plane about $\hat{x}^*$ like a pancake on the griddle) by a 180° pulse of the rf magnetic field. e) The rephasing the faster and slower magnetization groups to form an echo at $t = 2\tau$.

the average. The spins in the smaller field get out of phase by $-\Delta\theta$ with respect to the average. After the 180° pulse, they continue to precess at the same rate but their relative order is reversed. After a second interval τ the faster spins have caught up to the slower spins. The groups have come back into phase at 2τ .

Yet some loss of $M_{x,y}$ magnetization has occurred and the maximum height of the echo is not the same as the maximum height of the FID. This loss of transverse magnetization occurs because of stochastic fluctuation in the local fields at the nuclear sites which is not rephasable by the 180° pulse. These are the real T_2 processes that we are interested in measuring.

A series of 90° - τ - 180° pulse experiments, varying τ , and plotting the echo height as a function of time between the FID and the echo, will give us the *real* T_2 .

The transverse magnetization as measured by the maximum echo height is written as:

$$M_{x,y} = M_0 e^{-2\tau/T_2} \quad (26)$$

That's enough theory for now. Let's summarize:

1. Magnetic resonance is observed in systems whose constituent particles have both a magnetic moment and angular momentum.
2. The resonant frequency of the system depends on the applied magnetic field in accordance with the relationship $\omega = \gamma B_0$ where

$$\begin{aligned} \gamma_{\text{proton}} &= 2.675 \times 10^4 \frac{\text{rad/s}}{\text{gauss}} \\ f_0 &= 4.258 \text{ MHz/kilogauss} \end{aligned}$$

3. The thermal equilibrium magnetization is parallel to the applied magnetic field, and approaches equilibrium following an exponential rise characterized by the constant T_1 , the spin-lattice relaxation time.
4. Classically, the magnetization obeys the differential equation

$$\frac{d\vec{M}}{dt} = \gamma \vec{M} \times \vec{B}$$

where $\vec{B}$ may be a time dependent field.

5. Pulsed NMR employs a rotating radio frequency magnetic field described by

$$\vec{B}(t) = B_1 \cos(\omega t) \hat{i} + B_1 \sin(\omega t) \hat{j} + B_0 \hat{k}$$

6. The easiest way to analyze the motion of the magnetization during and after the rf pulsed magnetic field is to transform into a rotating coordinate system. If the system is rotating at an angular frequency ω , along the direction of the magnetic field, a fictitious magnetic field must be added to the real fields such that the total effective magnetic field in the rotating frame is:

$$\vec{B}_{\text{eff}}^* = B_1 \hat{i}^* + \left(B_0 - \frac{\omega}{\gamma}\right) \hat{k}^*$$

7. On resonance $\omega = \omega_0 = \vec{B}_{\text{eff}}^*$ and $\vec{B}_{\text{eff}}^* = B_1 \hat{i}^*$. In the rotating frame during the pulse the spins precess around B_1^*

8. A 90° pulse is one where the pulse is left on just long enough (t_w) for the equilibrium magnetization M_0 , to rotate to the x-y plane. That is:

$$\begin{aligned}\omega_1 t_w &= \pi/2 \text{ radians} \\ t_w &= \frac{\pi}{2\omega_1}\end{aligned}$$

Since B_1 is the only field in the rotating frame on resonance

$$\omega_1 = \gamma B_1$$

Thus during the 90° pulse

$$t_w = \frac{\pi}{2\gamma B_1}$$

9. The spin-spin relaxation time, T_2 , is the characteristic decay time for the nuclear magnetization in the x-y (or transverse) plane.
10. The spin-echo experiments allow the measurement of T_2 in the presence of a nonuniform static magnetic field. For those cases where the free induction decay time constant, (sometimes written T_2^*) is shorter than the real T_2 , the decay of the echo envelope's maximum heights for various times τ , gives the real T_2 .

EXPERIMENTAL SET-UP

The technique of using NMR to analyze materials has become so well established that any modern research lab using NMR would simply buy a research-grade NMR instrument, much like one buys a cryostat or an oscilloscope. Research-grade NMR spectrometers usually include strong, superconducting magnets, computer control, digital signal processing, and sometimes cryogenic sample spaces for studying NMR at low temperatures.

The instrument used in this lab is a TeachSpin PS1-A, a pulsed NMR apparatus designed for teaching the physics behind NMR measurements. It focuses on the use of 90° and 180° pulses for measuring T_1 , T_2 , and T_2^* . The pulse and measurement process is made transparent by making all of the signal paths accessible from the front panel of the electronics module.

The purpose of this lab is not to study the details of the spin relaxation mechanisms, but to learn how the measurements are made. The zero-crossing method used to measure T_1 and spin-echo technique employed for T_2 are exactly the same methods used in a research NMR lab. Even MRI, Magnetic Resonance Imaging, used for generating beautiful maps of the interiors of living things, is just a repeated application of these measurement techniques. Spatial information in these images comes from a strong magnetic field gradient, which allows us to correlate the positions of spins by their precession frequency, or the strength of the field felt.

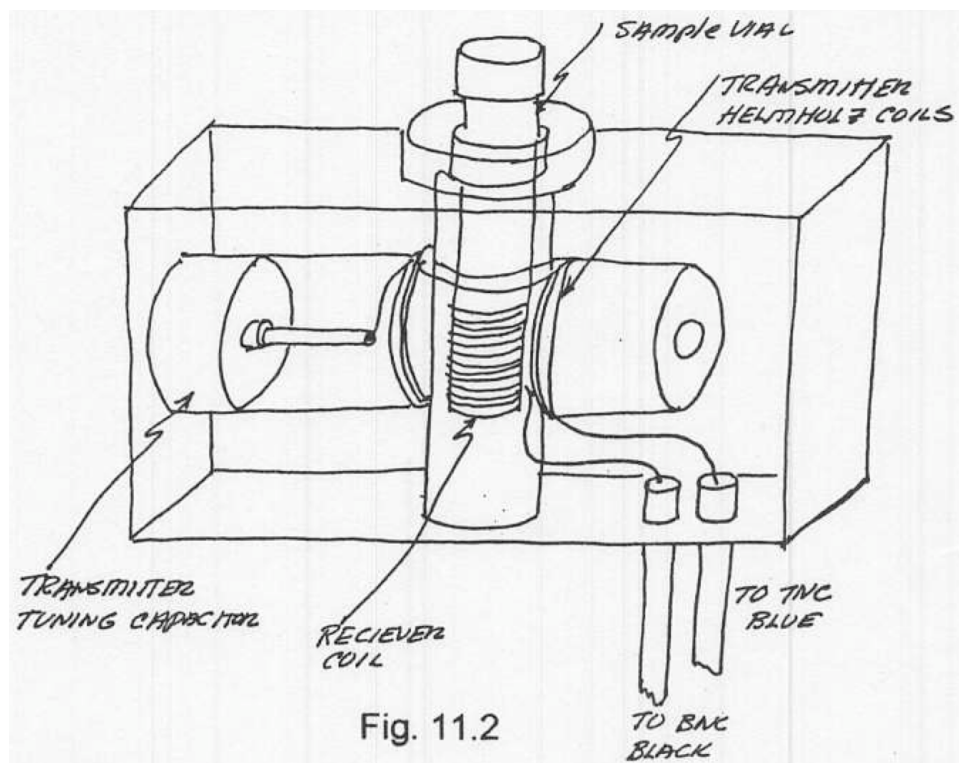


Figure 9 The magnets and coils used in the NMR.

Magnet Apparatus

The magnet apparatus, shown in Figure 9, consists of the static and oscillating magnet fields needed to rotate the sample's magnetization, and the detector coil used to measure the magnetization. The magnets are surrounded by a thick, and rather heavy, flux return used to contain the powerful magnetic field within the apparatus. The top and bottom surfaces are covered with plexiglass to prevent magnetic objects from being pulled into the strong field. Do not place near the field any magnetic or magnetically susceptible objects, including computer disks, cards with magnetic strips, or any ferrous material. **Keep the top cover closed at all times except when changing the sample.**

A permanent magnet supplies the static field B_0 along the z-axis. (In our coordinate system, $\hat{z}$ points horizontally out toward an observer facing the front of the apparatus, $\hat{y}$ points vertically up, and $\hat{x}$ points to the right.) The brass box between the poles of the permanent magnet contains the rotating coils, the sample holder and the detector coil. The location of the box in the X-Y plane may be adjusted by two knobs. The X position is adjusted using the knob on the right side of the apparatus and the position is given on a slide indicator on the front. The Y position is adjusted using the knob on the front of the apparatus and the position is shown on the dial. As is shown later in this lab, one can measure the uniformity of the magnetic field by measuring the precession frequency of a sample as a function of the X-Y position. A rough map of the permanent magnet is given in Figure 10.

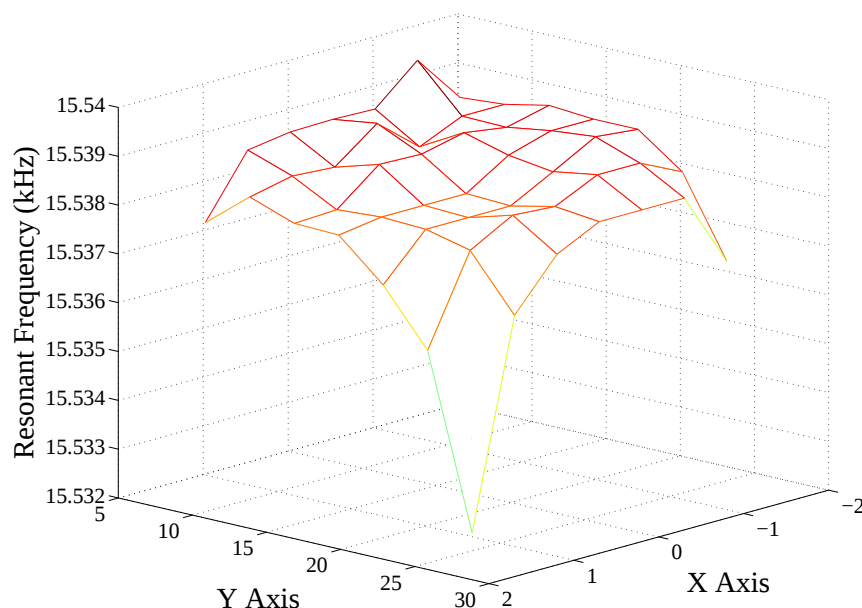


Figure 10 A coarse map of the permanent magnet. The *sweet spot* is at (0.5,18.5)

In the introduction we discussed how the spin flip (or rotation of the classical angular momentum) was produced by precession about a horizontal field that rotated at the Larmour precession frequency. In our apparatus, the rotating magnetic field is supplied by a Helmholtz coil with its axis along $\hat{x}$. The field this coil supplies is not actually rotating, rather it oscillates sinusoidally at the precession frequency. This field can be thought of as the sum of two half-amplitude, counter-rotating fields in the X-Y plane. If we transform into the rotating frame, the field that rotates in the direction of the precessing protons will appear static in the $\hat{x}'$ direction. However the other field will appear to rotate at twice the precession frequency. Like any oscillatory system, when driven at frequencies above the resonant frequency the response is quite diminished. In our case the counter-rotating field produces a small oscillation in the azimuthal angle of the magnetic dipole that can be neglected.

The permanent magnet's field strength is temperature dependent, changing about 4 Gauss/°C. This change corresponds to a shift in the proton's resonant frequency of 17 kHz/°C. Therefore, keep the magnet away from direct heat sources such as incandescent lamps, direct sunlight, or large power supplies. **Daily drifts in the temperature will require that the oscillator frequency be retuned each day.**

Electronics Module

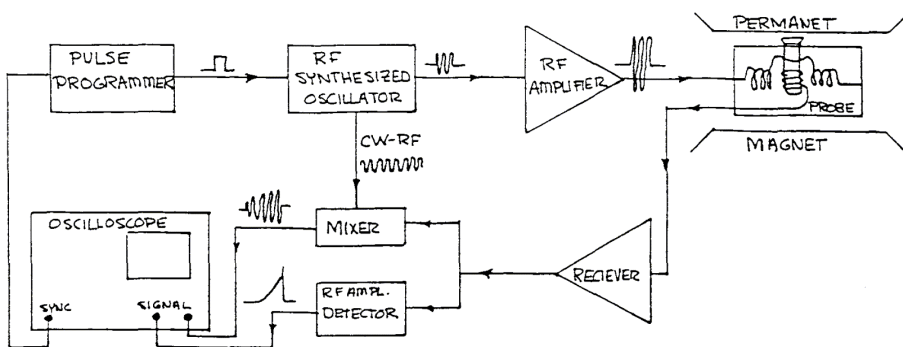


Figure 11 A block diagram of the pulsed-NMR apparatus.

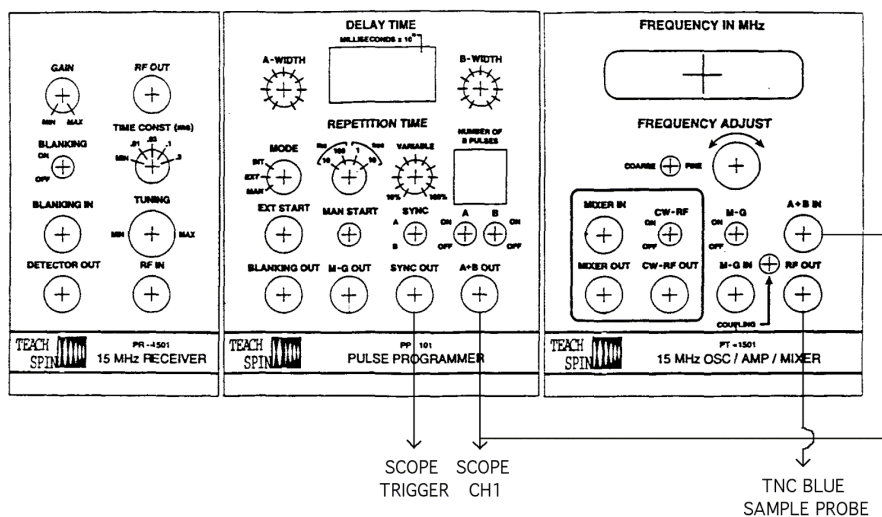


Figure 12 Wiring diagram for observing pulse gating signal and measuring B .

A block diagram of the TeachSpin apparatus is shown in Figure 11. This many components in this diagram are all housed in three devices: the oscilloscope, the magnet apparatus, and the electronics module. The electronics module contains the receiver, the pulse programmer, and a third component containing the oscillator, the power amplifier, and the mixer.

An event begins with the pulse programmer that produces a series of logic pulses to control when the horizontal coils are on or off. The pulses range in width from $1\text{--}30\ \mu\text{s}$. There can be multiple pulses, but the first pulse (the A pulse) is controlled separately. All subsequent B pulses are identical. There is a variable delay between pulses ranging from $10\ \mu\text{s}$ to $10\ \text{s}$. The pulse programmer also issues a logic pulse, a sync signal, that can be used to trigger the oscilloscope. The sync can be coincident with

either the A or the first B pulse.

The pulses are sent to the oscillator which outputs to the power amplifier a sine wave signal only when the pulse is high (nonzero). The power amplifier uses that signal to drive the horizontal coils.

If the proton spins are aligned in the X-Y plane then they will precess about the static field, $B_0 \hat{z}$. The rotating magnetic field produced by the protons induces an emf in the probe coil. Note carefully the orientation of the coils around the sample! As shown in this figure, the permanent magnetic field is applied perpendicular to the page; the applied magnetic field B is orthogonal to the axis of the sample tube; and the axis of the receiver coil is parallel to the test tube's axis.

The probe signal is sent to the receiver which amplifies the signal and filters out noise outside of the frequency range of interest. The output of the receiver goes both to a mixer (analog multiplier) and to an RF amplitude detector (rectifier).

The mixer multiplies inputs from the receiver and from the oscillator. This multiplication allows for the detection of beating between the two signals. By eliminating the beating one ensures that the driving frequency is equal to the precession frequency.

The amplitude detector produces an envelope of the oscillating signal from which the time constant of the decay can be easily observed.

PRELAB EXERCISES, PART 1

These exercises provide the mathematical underpinnings of the principles that are the basis of this lab. Please review the questions now. You should collectively work out the solutions and provide them in the appendix of your report.

1. Consider a collection of independent protons at room temperature in a magnetic field of 4kG. Approximately what fraction of these protons will be aligned parallel with the magnetic field, as opposed to antiparallel?
2. If we disturb this distribution, say with a 180° pulse, we may expect the system to decay back to thermal equilibrium with some characteristic time T_1 . Phenomenologically, we can modify the precession equation for the net magnetization $\vec{M}$ to include this decay term.

$$\begin{aligned}\dot{M}_x &= \gamma (\vec{M} \times \vec{B})_x \\ \dot{M}_y &= \gamma (\vec{M} \times \vec{B})_y \\ \dot{M}_z &= \gamma (\vec{M} \times \vec{B})_z + (M_0 - M_z) / T_1\end{aligned}$$

where M_0 is the equilibrium magnetization along $\vec{B} = B_0 \hat{k}$.

Consider a system in thermal equilibrium for time $t < 0$. Let a 180° pulse be applied just before $t = 0$ so that, at $t = 0$, $\vec{M} = -M_0 \hat{k}$. Solve

for $\vec{M}(t)$ for time $t \geq 0$, and plot $M_z(t)$ from $t = 0$ to $t = 3T_1$. For what value of t is the net magnetization zero?

3. For this exercise and the next one, consider the case of a non-uniform magnetic field. Here, each spin feels a slightly different value of B_0 and thus precesses at a slightly different frequency. In this case, it will be useful for us to consider the net magnetization as the sum of the individual magnetic moments of each proton.

$$\vec{M} = \sum_{i=1}^N \vec{\mu}_i,$$

where each magnetic moment $\vec{\mu}_i$ is a classical vector of constant magnitude $e\hbar/2Mc$, obeying the precession equation

$$\dot{\vec{\mu}}_i = \gamma \vec{\mu}_i \times \vec{B}_i.$$

Now model the magnetic field at each proton as

$$\vec{B}_i = (B_0 + b_i) \hat{\mathbf{k}},$$

where B_0 is the average field value, and the individual variations b_i have a mean of zero. In the rotating reference frame, the precession equations then become

$$\dot{\vec{\mu}}_i = \gamma \vec{\mu}_i \times (b_i \hat{\mathbf{k}}_r).$$

Let us consider, as in the last exercise, a system in thermal equilibrium for times $t < 0$, with $\vec{M} = M_0 \hat{\mathbf{k}}$. This time, however, let's apply a 90° pulse at $t = 0$, instead of a 180° pulse, with rotation being about $\hat{i}_r$ in the rotating reference frame. Just after the 90° pulse, then, the magnetic moment of each proton becomes, in this rotating frame,

$$\vec{\mu}_i(0) = \frac{e\hbar}{2Mc} \hat{j}_r,$$

and hence the magnetization is $\vec{M}(0) = M_0 \hat{j}_r$.

Solve the precession equations for the individual moments $\vec{\mu}_i$ in the rotating frame for times $t \ll T_1$, then evaluate the sum to get the net magnetization $\vec{M}(t)$. Show that this net magnetization decays due to dephasing with a characteristic timescale

$$\frac{1}{T_2^*} = \gamma \sqrt{\langle b_i^2 \rangle}.$$

Hint: What you will get here is a sum of phases, of the form

$$\sum_{n=1}^N e^{i\phi_n}$$

To evaluate this sum, expand each phase as a series, then sum each order in ϕ separately.

$$\begin{aligned}\sum_{n=1}^N e^{i\phi_n} &= \sum_{n=1}^N \left\{ 1 + i\phi_n + \frac{1}{2}(i\phi_n)^2 + \dots \right\} \\ &= N + \sum_{n=1}^N i\phi_n + \frac{1}{2} \sum_{n=1}^N (i\phi_n)^2 + \dots\end{aligned}$$

The series you wind up with will have an (approximate) exponential representation, with a characteristic decay time equal to T_2^* .

4. Now apply a 180° pulse at a much later time $t = \tau$ when the initial free-induction-decay signal has long since died away, *i.e.* $\tau > T_2$. Show that the magnetization will return to its full initial value at time $t = 2\tau$, then decay again with the same time constant T_2^* . This phenomenon is known as *spin echo*.

Hint: Replace the x-y plane in the rotating frame with the real and imaginary axes of the complex plane, write the vector $\vec{\mu}_i$ as a complex number, and express the precession equation in this complex notation. Remember that the 180° pulse essentially rotates all of the spins around the x-axis, in the rotating reference frame, and this will be particularly easy to represent if the complex plane is used to represent the x-y plane. This is the starting point for a lot of NMR dephasing analysis and is widely used in the NMR literature.

If there are several, uncorrelated mechanisms that cause the variation in $\vec{B}_i$, then the dephasing rates for each just add to give a total dephasing rate T_2^* .

$$\frac{1}{T_2^*} = \frac{1}{T_2^A} + \frac{1}{T_2^B} + \dots$$

Some examples of sources of nonuniformity are

1. Inhomogeneities in the field of the permanent magnet used to supply the field $\vec{B}$, and
2. Interactions between individual spins.

If you were doing research, you would probably want to look at the interactions between the spins and be uninterested in the inhomogeneities of the magnet. Because the inhomogeneities in the magnet do not change with time, any dephasing they cause can be recovered using this spin-echo technique. The spin-spin interactions, however, change over time, and the dephasing they cause cannot be recovered by spin-echo.

We may expect the net magnetization, then, to decay immediately following the initial 90° pulse as

$$M(t) \sim M(0)e^{-(t/T_2^*)^2}.$$

This initial decay is known as the *free-induction decay*, or FID. Spin-spin interactions, and other effects, should keep the spin echo signal from returning to its full initial height $M(0)$, and the height of the spin echo should decay as well, with a time constant of $T_2 \gg T_2^*$.

This provides us with a way of separating the irreversible dephasing time T_2 from the total dephasing time T_2^* . Because T_2 is intrinsic to the sample and T_2^* depends on your apparatus, it is usually T_2 that you are interested in.

Applied RF field and 90° pulses

The first thing we are going to do is look at the applied magnetic field B . Supplied with the instrument are two vials with small loops of wire embedded in epoxy and coaxial cables connected to these loops. We will observe B and measure its amplitude by placing one of these coils in the sample space and looking at the voltage induced by the oscillating B field.

Notice that the two coils have different orientations. One is designed for measuring B , and the other, which we will not use in this lab, is designed for generating a false signal in the pickup coil. For the orientation of coils shown in Figure 11, which coil should you use to measure B ? What is the relationship between the voltage generated across the coil and B ? What property of the coil do you need to measure to convert between this voltage and B ? (This will just be a rough estimate, so don't spend too much time on precision here. 20% or so will be fine.) Clearly describe your procedure for measuring B in your lab book, along with the formula you use for converting your observed voltage to B and how you arrived at it.

Note: When you attach the cable connected to your measurement coil to your oscilloscope, set the input impedance of that channel to be 50Ω .

If there are any cables connected to the front panel of the TeachSpin electronics bin, disconnect them and turn the instrument on. (The switch is in the back.) Locate the A+B OUT port, in the PULSE PROGRAMMER module. Using a tee, connect this to the A+B IN port on the 15 MHz OSC/AMP/MIXER module and then to the other channel of your oscilloscope. Trigger the scope off of the SYNC OUT signal, provided by the pulse programmer module. The blue cable attached to the sample holder with the TNC connector on the end of it supplies current to the Helmholtz coil. Plug this into the RF OUT port on the 15 MHz OSC/AMP/MIXER module.

Because this instrument is only designed for pulsed NMR, it does not supply a continuous RF signal to the Helmholtz coils. Instead, the RF oscillator is gated by the pulse programmer. When the voltage going into the A+B IN port is high, a signal is applied to the coils. When it is low (zero), no current flows through the coils, and no magnetic field is applied. The length of time over which the field is applied is set by the A-WIDTH knob on the pulse programmer module. On the pulse programmer module, make sure that MODE is set to INT, the SYNC switch is set to A, the A switch is ON, and the B switch is OFF.

You should now see both the gating signal and an RF signal, as detected by the coil you placed in the sample space, on your oscilloscope screen. Make sure your coil is optimally aligned, measure the amplitude of the RF signal you observe, and use this to estimate the magnitude of B .

Estimate how long this signal should be applied to produce a 90° pulse, and set the A-WIDTH to this value. Take a screenshot, print it out, and include it in your lab book.

Single-pulse free-induction decay

Now remove the test coil from the sample space and disconnect it from your oscilloscope. The RF pickup coil, wrapped around the sample space inside the sample holder assembly, is connected to a thin, black coaxial cable with an ordinary BNC connector at the other end. Connect this cable to the RF IN port in the 15 MHz RECEIVER module. Send the RF OUT signal to your scope, and send the BLANKING OUT signal to the BLANKING IN port. Start with the BLANKING switch OFF.

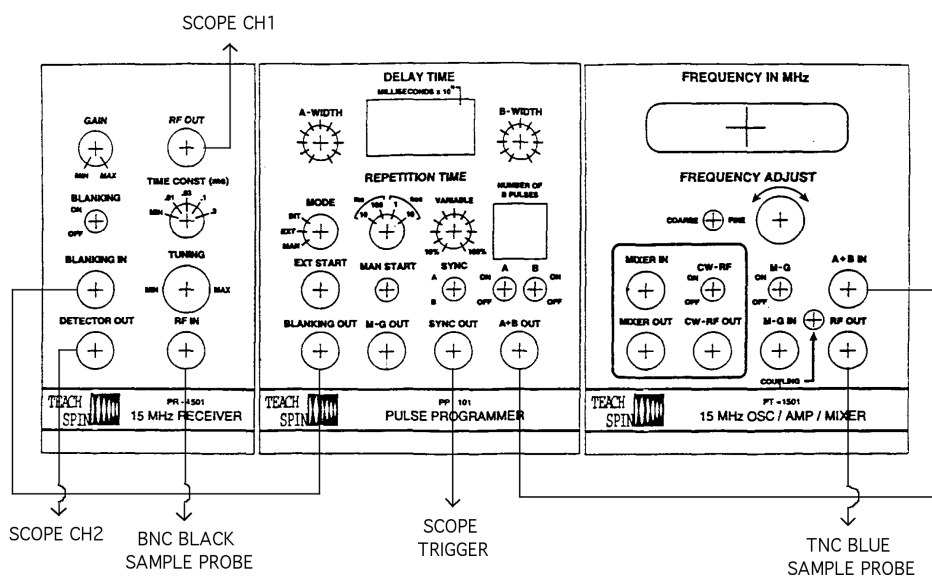


Figure 13 Wiring diagram for observing free-induction decay (FID).

There is, as you should observe, considerable crosstalk between the Helmholtz coils and the receiver coil! The blanking signal here is essentially the opposite of the A+B gating signal you looked at in the last section. With blanking on, whenever an RF signal is applied to the Helmholtz coils, the RF receiver is turned off. Look at the RF OUT and A+B OUT signals on your scope with blanking on and off. Take screenshots of both cases, and include them in your lab book. Now switch blanking on and leave it there.

Pick out a sample, either mineral oil or glycerin, and insert it into the sample space. Adjust the scale on your scope until you can see the free induction decay signal as it appears after the A pulse ends. Play with the GAIN and TUNING knobs to get a good signal, then take screenshots

with the sample in and out to demonstrate that the signal you see really does come from the sample. (Adjust the tuning to maximize your signal and the gain to keep it from saturating.)

Now we are at a point where we no longer need to look at the gating signal. Disconnect the A+B signal from your scope, and replace it with the DETECTOR OUT signal. (This is just a rectified version of the RF OUT signal.) In the last section, you set the A-WIDTH to approximate a 90° pulse. Now you should fine-tune the A-WIDTH to a more precise 90° pulse by maximizing your observed free-induction decay. Measure T_2^* , and record your result in your lab book.

Play around with the TIME CONSTANT knob to see what that does. Record your observations. (*Hint:* If you are at all familiar with passive filters, it does exactly what you'd expect.)

Mixer

We have adjusted the receiver circuit to optimize our signal, but we have not yet paid any attention to our oscillator. For resonance to occur, recall that the applied RF magnetic field must be at the same frequency as the natural proton precession frequency in the constant magnetic field B_0 . If we are close, we will still be able to stimulate some precession, but for accurate measurements of T_1 and T_2 we will need to meet the resonance condition with a fair degree of precision.

We can compare our oscillator's frequency with the protons' precession frequency by using a *mixer*, contained in the 15 MHz OSC/AMP/MIXER module. A mixer is a nonlinear, passive electronic device that essentially multiplies two signals. Recall that the product of two sine waves at different frequencies has the form

$$\sin(\omega_1 t) \sin(\omega_2 t) = \frac{1}{2} \{ \cos[(\omega_1 - \omega_2)t] - \cos[(\omega_1 + \omega_2)t] \}.$$

If the two inputs of a mixer contain signals at different frequencies, the output contains signals at both the sum $(\omega_1 + \omega_2)$ and difference $(\omega_1 - \omega_2)$ frequencies.

In the TeachSpin apparatus, one of the mixer's inputs is internal and not accessible from the front panel. If the CW-RF switch is ON, this internal input receives a signal from the 15 MHz oscillator, which also appears on the CW-RF OUT port. The other input is accessible from the front panel and is labeled MIXER IN. Connect the RF OUT port (on the receiver module) to the MIXER IN port, and send the MIXER OUT signal to your oscilloscope. Note that you will only see the difference-frequency signal, as the sum-frequency signal, at ~ 30 MHz, is filtered out inside the module. During a free-induction decay, you should see a beat signal, on the output of the mixer, between the oscillator and the FID signal, which oscillates at the precession frequency. Tune the oscillator frequency to get it as close as you can to the natural precession frequency. Record the precession frequency, and use it to calculate the value of B_0 . Record a screenshot of the optimized mixer-out signal in your lab book.

(This tuning procedure will be familiar to anyone who plays a stringed musical instrument. When tuning a guitar, for example, you compare some reference tone to that produced by a string, adjusting the tension on that string to eliminate audible beats between the two tones.)

Once you have tuned the oscillator frequency, verify that your A-WIDTH, TUNING, and GAIN settings are still optimized. Correct them if they are not. From now on, you should monitor the MIXER OUT signal any time you are performing a measurement, to make sure your oscillator is still tuned to resonance. The permanent magnet is very sensitive to temperature drifts, and the precession frequency is likely to drift over time.

MULTIPLE PULSE SEQUENCES

Now that you have learned how to tune the system to produce clean 90° pulses and observed a single-pulse free-induction decay, we are ready to start investigating properties of the sample. In particular, we will explore techniques for measuring T_1 and T_2 . For this you will need to apply sequences of pulses, and the PS1-A has the capability to do that.

Look at the A+B OUT signal on your scope again, only this time switch B to ON. Start with the following settings:

DELAY TIME	1.00×10^1 ms
MODE	INT
REPETITION TIME	1 s
NUMBER OF B PULSES	02
SYNC	A
A: ON	B: ON

You should see three pulses on your scope screen. Play around with the settings to get a feel for what they do. Change the A and B widths. (Don't forget to record the A width first so that you can go back to a good 90° pulse later.) Change the delay time, sync source, A and B switches on and off, and the repetition time.

Zero-crossing measurement of T_1

In your prelab exercises (# 4) you saw how the magnetization would decay back to equilibrium after a 180° pulse, and how the zero-crossing point of this decay provides a measure of T_1 . Set your system up to produce only two pulses, one A pulse and one B pulse. Set the A-width to produce a 180° pulse and the B-width to produce a 90° pulse. The first, 180° pulse inverts the sign of the magnetization along the z-axis, and the second pulse is used to measure the net magnetization. Do this for several values of the DELAY TIME, the time between the A and B pulses, and plot the net magnetization versus the time after a 180° pulse. Fit your theory from prelab Exercise 4 to your data, and use this to measure T_1 .

Notes:

- Tune the A and B pulses separately to make sure they are good 180° and 90° pulses, respectively. We saw earlier how to tune the pulse width to produce a good 90° pulse by maximizing the free-induction decay. What should the free-induction decay be like for a perfect 180° pulse? Why?
- Note that we do not need to convert the amplitude of the initial FID following the B pulse to an actual magnetization. As you saw in the prelab exercises, the zero-crossing technique gives T_1 without requiring us to know the actual value of $M(t)$, only its shape.
- The magnetization decay curve you observe has one distinct difference from the one you predicted in the prelabs. Why is what you observe different from what you predicted?

Spin-echo and T_2

Now reverse the order of the pulses. Set the A pulse to be a 90° pulse and the B width for a 180° pulse. Again, tune each separately, using the A and B switches to isolate each one. See if you can observe a single spin echo. (It may be useful to plot, on your scope, A+B OUT and DETECTOR OUT simultaneously. A good delay time to start with here is $100\mu s$.) Take a screenshot of your spin-echo signal, and record it in your lab book.

We could now just vary the delay time and plot the height of the spin echo to get T_2 , but there is a less tedious and more accurate way to do this measurement. Increase the number of B pulses from one to twenty or more, and make sure the M-G switch is set to OFF. Observe all of the subsequent spin-echo pulses, and from this infer an initial estimate of T_2 .

One thing to be careful about here is that, if the B-widths are not perfect, the B-pulses will not produce exactly 180° of rotation. If the B-pulses produce, say, only 175° of rotation, then after ten pulses an error of 50° has built up. This will substantially attenuate the spin echo and contaminate your measurement of T_2 ! To see how sensitive the spin echo is to the B width, slightly adjust the B width to try and maximize the amplitude of the spin-echo train. As you can see, your measurement of T_2 by this method is not very reliable. Take a screenshot of your optimized spin-echo train, and record it in your lab book.

One way around this dilemma is to alternate the sign of the 180° pulses for every other B pulse. Now, if one B pulse only produces 175° of rotation, the next pulse will rotate the system by -175° , for a net accumulated error of zero. This alternating-sign technique is named after its inventors, Meiboom and Gill, and it is applied in the PS1-A by connecting the M-G OUT port on the pulse programmer to the M-G IN port on the oscillator module, and switching M-G to ON. Do this, and use the Meiboom-Gill technique to get a refined measurement of T_2 .

(The first method, that of applying a train of what-you-hope-are- 180° pulses, is known as the Carr-Purcell method, after the first discoverers of NMR.)

PRELAB II

5. In the first prelab, you modeled the free induction decay due to the individual spins' dephasing in an inhomogeneous magnetic field. There, we assumed that each spin stayed where it was, which made it possible to recover the magnetization by the spin-echo technique. In the first lab session, you saw that the spin echo did not fully recover all of its original magnetization, and that the height of the echos decayed with a time constant we called T_2 . There are many possible causes for this irreversible dephasing, and in this second lab session you will explore one of the most common ones, diffusion.

Molecules in a liquid sample wander around, a process we call diffusion, and their average motion is easily quantified by a partial differential equation known, not surprisingly, as the diffusion equation. Diffusion is covered in most texts on differential equations, and there is an excellent chapter in the Feynman lectures on the subject (Volume I, Chapter 41, on Brownian motion). We won't go into the physics of diffusion here, but we will use one result from that theory. The average distance L a molecule diffuses over a time t from its starting point is

$$\langle L^2 \rangle = Dt$$

Here, the constant D is known as the diffusion constant, and it depends on the properties of the liquid and the molecule doing the diffusing.

Recall how, in the first prelab, you modeled dephasing of an ensemble of protons evenly distributed in a spatially-varying magnetic field. In that case, you assumed that the protons precessed but did not diffuse, and thus the magnetic field they felt did not change over time. Now consider a collection of protons that all start out near the origin and, at time $t = 0$, feel the same initial magnetic field. This time, however, allow them to diffuse into regions where the field strength is different, and calculate the expected net magnetization.

Hint: Taylor-expand the magnetic field about the origin to approximate the spatial variation, like so

$$b_i(x, y, z) \approx \frac{\partial B}{\partial x} \Delta x + \frac{\partial B}{\partial y} \Delta y + \frac{\partial B}{\partial z} \Delta z,$$

then assume

$$\left| \frac{\partial B}{\partial x} \right| \approx \left| \frac{\partial B}{\partial y} \right| \approx \left| \frac{\partial B}{\partial z} \right|$$

Also, assume isotropic diffusion, which is perfectly reasonable for a liquid sample.

$$\langle \Delta x(t)^2 \rangle = \langle \Delta y(t)^2 \rangle = \langle \Delta z(t)^2 \rangle = \frac{1}{3}Dt$$

Finally, assume that how far a molecule diffuses in one direction is completely unrelated to how far it diffuses in any other direction.

$$\langle \Delta x(t) \Delta y(t) \rangle = 0, \text{ etc.}$$

You should find that the magnetization decays as

$$M(t) \approx M(0)e^{-(\frac{t}{T_2})^3},$$

and you will derive an expression for T_2 . You have made some pretty drastic approximations, the most egregious being all of the protons starting out at the same place. A complete and careful derivation requires a lot more work than this, but the result is not too different. The actual answer is

$$M(t) = M(0)e^{-\frac{1}{6}(\frac{t}{T_2})^3},$$

using the expression for T_2 you obtain.

LAB II

Self diffusion and viscosity

With both field inhomogeneity and diffusion contributing to the magnetization decay, there will be an initial free-induction decay, and the heights of subsequent spin-echos should decay with an envelope given by the $\exp(-t^3/T_2^3)$ law you derived in the prelab. Using the Meiboom-Gill method, observe the spin-echo decay, and check to see if it obeys the power law you expect.

The T_2 you derived in the prelab should depend on the field gradient $\partial B/\partial z$. The PS1-A's sample holder can be moved along both the x- and z-axes, using the knobs on the front of the magnet assembly. Use this degree of freedom to measure the field gradient and to check the approximation $|\partial B/\partial z| \approx |\partial B/\partial x|$. (Remember, the value of the field is got from the precession frequency.) Using either a glycerine or mineral oil sample, measure the field contours in both the x- and y-axes, then determine the diffusion constant D of your sample from your data.

REFERENCES

1. Ramamurti Shankar, *Principles of Quantum Mechanics*, Plenum Press, New York (1980).
2. Brian Cowan, *Nuclear Magnetic Resonance and Relaxation*, Cambridge University Press, Cambridge UK (1997).
3. Feynman, Leighton, and Sands, *The Feynman Lectures on Physics*, Addison-Wesley, Reading Mass. (1963).

Muon Physics

originally by T. E. Coan and J. Ye, edited by S. Penn

INTRODUCTION

The muon is one of nature's fundamental building blocks of matter and acts in many ways as if it were an unstable heavy electron, for reasons no one fully understands. Discovered in 1937 by C.W. Anderson and S.H. Neddermeyer when they exposed a cloud chamber to cosmic rays, its finite lifetime was first demonstrated in 1941 by F. Rasetti. The instrument described in this manual permits you to measure the charge averaged mean muon lifetime in plastic scintillator, to measure the relative flux of muons as a function of height above sea-level and to demonstrate the time dilation effect of special relativity. The instrument also provides a source of genuinely random numbers that can be used for experimental tests of standard probability distributions.

THE MUON SOURCE

The top of earth's atmosphere is bombarded by a flux of high energy charged particles produced in other parts of the universe by mechanisms that are not yet fully understood. The composition of these "primary cosmic rays" is somewhat energy dependent but a useful approximation is that 98% of these particles are protons or heavier nuclei and 2% are electrons. Of the protons and nuclei, about 87% are protons, 12% helium nuclei and the balance are still heavier nuclei that are the end products of stellar nucleosynthesis. See Simpson in the reference section for more details.

The primary cosmic rays collide with the nuclei of air molecules and produce a shower of particles that include protons, neutrons, pions (both charged and neutral), kaons, photons, electrons and positrons. These secondary particles then undergo electromagnetic and nuclear interactions to produce yet additional particles in a cascade process. Figure 1 indicates the general idea. Of particular interest is the fate of the charged pions produced in the cascade. Some of these will interact via the strong force with air molecule nuclei but others will spontaneously decay (indicated by the arrow) via the weak force into a muon plus a neutrino or antineutrino:

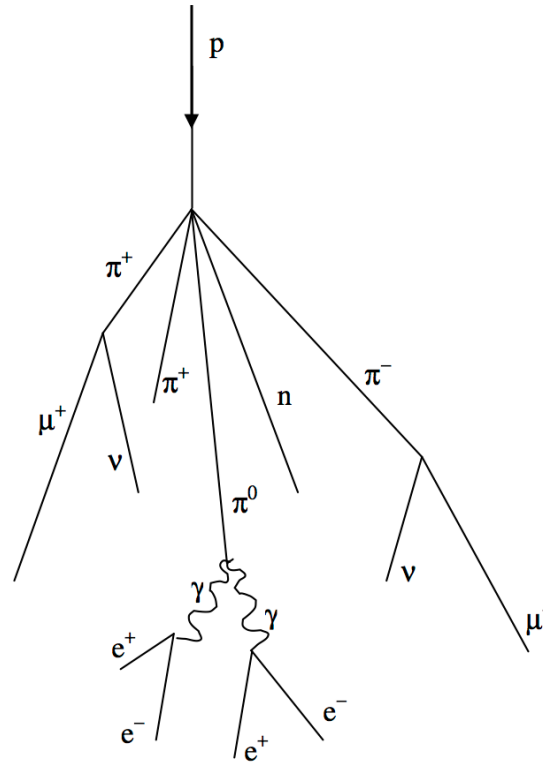


Figure 1 Cosmic ray cascade induced by a cosmic ray proton striking an air molecule nucleus.

$$\begin{aligned}\pi^+ &\rightarrow \mu^+ \nu_{\mu} \\ \pi^- &\rightarrow \mu^- \bar{\nu}_{\mu}\end{aligned}$$

The muon does not interact with matter via the strong force but only through the weak and electromagnetic forces. It travels a relatively long instance while losing its kinetic energy and decays by the weak force into an electron plus a neutrino and antineutrino. We will detect the decays of some of the muons produced in the cascade. (Our detection efficiency for the neutrinos and antineutrinos is utterly negligible.)

Not all of the particles produced in the cascade in the upper atmosphere survive down to sea-level due to their interaction with atmospheric nuclei and their own spontaneous decay. The flux of sea-level muons is approximately 1 per minute per cm² (see <http://pdg.lbl.gov> for more precise numbers) with a mean kinetic energy of about 4 GeV.

Careful study [<http://pdg.lbl.gov>] shows that the mean production height in the atmosphere of the muons detected at sea-level is approximately 15 km. Travelling at the speed of light, the transit time from production point to sea-level is then 50 μ sec. Since the lifetime of at-rest muons is more than a factor of 20 smaller, the appearance of an appreciable sea-level muon flux is qualitative evidence for the time dilation effect of special relativity.

MUON DECAY TIME DISTRIBUTION

The decay times for muons are easily described mathematically. Suppose at some time t we have $N(t)$ muons. If the probability that a muon decays in some small time interval dt is λdt , where λ is a constant *decay rate* that characterizes how rapidly a muon decays, then the change dN in our population of muons is just $dN = -N(t) \lambda dt$, or $dN/N(t) = -\lambda dt$. Integrating, we have $N(t) = N_0 e^{-\lambda t}$, where $N(t)$ is the number of surviving muons at some time t and N_0 is the number of muons at $t = 0$. The *lifetime* τ of a muon is the reciprocal of λ , $\tau = 1/\lambda$. This simple exponential relation is typical of radioactive decay.

Now, we do not have a single clump of muons whose surviving number we can easily measure. Instead, we detect muon decays from muons that enter our detector at essentially random times, typically one at a time. It is still the case that their decay time distribution has a simple exponential form of the type described above. By decay time distribution $D(t)$, we mean that the time-dependent probability that a muon decays in the time interval between t and $t + dt$ is given by $D(t) dt$. If we had started with N_0 muons, then the fraction $-dN/N_0$ that would on average decay in the time interval between t and $t + dt$ is just given by differentiating the above relation:

$$\begin{aligned} -dN &= N_0 \lambda e^{-\lambda t} dt \\ -dN/N_0 &= \lambda e^{-\lambda t} dt \end{aligned}$$

The left-hand side of the last equation is nothing more than the decay probability we seek, so $D(t) = \lambda e^{-\lambda t}$. This is true regardless of the starting value of N_0 . That is, the distribution of decay times, for new muons entering our detector, is also exponential with the very same exponent used to describe the surviving population of muons. Again, what we call the muon lifetime is $\tau = 1/\lambda$.

Because the muon decay time is exponentially distributed, it does not matter that the muons whose decays we detect are not born in the detector but somewhere above us in the atmosphere. An exponential function always *looks the same* in the sense that whether you examine it at early times or late times, its e-folding time (time to decrease by a factor e) is the same.

DETECTOR PHYSICS

The active volume of the detector is a plastic scintillator in the shape of a right circular cylinder of 15 cm diameter and 12.5 cm height placed at the bottom of the black anodized aluminum alloy tube. Plastic scintillator is transparent organic material made by mixing together one or more fluors with a solid plastic solvent that has an aromatic ring structure. A charged particle passing through the scintillator will lose some of its kinetic energy by ionization and atomic excitation of the solvent molecules.

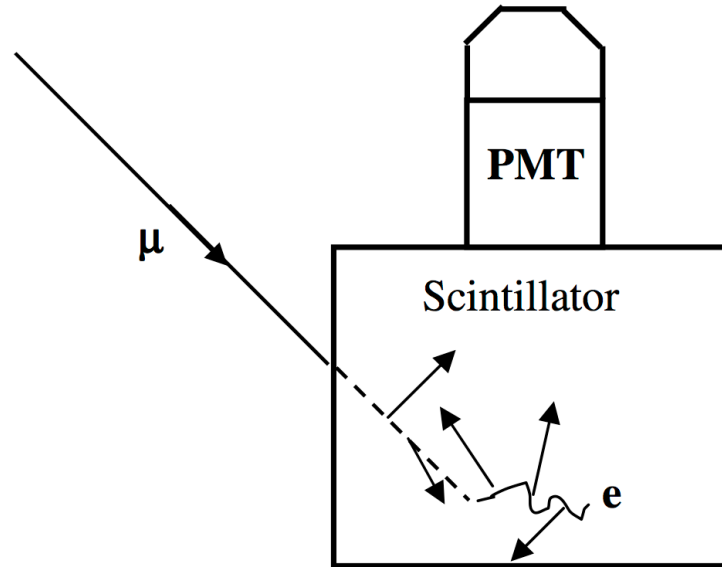


Figure 2 Schematic showing the generation of the two light pulses (short arrows) used in determining the muon lifetime. One light pulse is from the slowing muon (dotted line) and the other is from its decay into an electron or positron (wavey line).

Some of this deposited energy is then transferred to the fluor molecules whose electrons are then promoted to excited states. Upon radiative de-excitation, light in the blue and near-UV portion of the electromagnetic spectrum is emitted with a typical decay time of a few nanoseconds. A typical photon yield for a plastic scintillator is 1 optical photon emitted per 100 eV of deposited energy. The properties of the polyvinyltoluene-based scintillator used in the muon lifetime instrument are summarized in table 1.

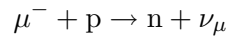
To measure the muon's lifetime, we are interested in only those muons that enter, slow, stop and then decay inside the plastic scintillator. Figure 2 summarizes this process. Such muons have a total energy of only about 160 MeV as they enter the tube. As a muon slows to a stop, the excited scintillator emits light that is detected by a photomultiplier tube (PMT), eventually producing a logic signal that triggers a timing clock. (See the electronics section below for more detail.) A stopped muon, after a bit, decays into an electron, a neutrino and an anti-neutrino. (See the next section for an important qualification of this statement.) Since the electron mass is so much smaller than the muon mass, $m_\mu/m_e \sim 210$, the electron tends to be very energetic and to produce scintillator light essentially all along its pathlength. The neutrino and anti-neutrino also share some of the muon's total energy but they entirely escape detection. This second burst of scintillator light is also seen by the PMT and used to trigger the timing clock. The distribution of time intervals between successive clock triggers for a set of muon decays is the physically interesting quantity used to measure the muon lifetime.

INTERACTION OF μ^- 'S WITH MATTER

Mass Density	1.032 g/cm ³
Refractive index	1.58
Base material	Polyvinyltoluene
Rise time	0.9 ns
Fall time	2.4 ns
$\lambda_{\max}$	423 nm

Table 1 General Scintillator Properties

The muons whose lifetime we measure necessarily interact with matter. Negative muons that stop in the scintillator can bind to the scintillator's carbon and hydrogen nuclei in much the same way as electrons do. Since the muon is not an electron, the Pauli exclusion principle does not prevent it from occupying an atomic orbital already filled with electrons. Such bound negative muons can then interact with protons



before they spontaneously decay. Since there are now two ways for a negative muon to disappear, the effective lifetime of negative muons in matter is somewhat less than the lifetime of positively charged muons, which do not have this second interaction mechanism. Experimental evidence for this effect is shown in figure 3 where *disintegration* curves for positive and negative muons in aluminum are shown. (See Rossi, 1952) The abscissa is the time interval t between the arrival of a muon in the aluminum target and its decay. The ordinate, plotted logarithmically, is the number of muons greater than the corresponding abscissa. These curves have the same meaning as curves representing the survival population of radioactive substances. The slope of the curve is a measure of the effective lifetime of the decaying substance. The muon lifetime we measure with this instrument is an average over both charge species so the mean lifetime of the detected muons will be somewhat less than the free space value $\tau = 2.19703 \pm 0.00004 \mu\text{sec}$.

The probability for nuclear absorption of a stopped negative muon by one of the scintillator nuclei is proportional to Z^4 , where Z is the atomic number of the nucleus [Rossi, 1952]. A stopped muon captured in an atomic orbital will make transitions down to the K-shell on a time scale short compared to its time for spontaneous decay [Wheeler]. Its Bohr radius is roughly 200 times smaller than that for an electron due to its much larger mass, increasing its probability for being found in the nucleus. From our knowledge of hydrogenic wavefunctions, the probability density for the bound muon to be found inside the nucleus is proportional to Z^3 . Once inside the nucleus, a muon's probability for encountering a proton is proportional to the number of protons there and so scales like Z . The net effect is for the overall absorption probability to scale like Z^4 . Again, this effect is relevant only for negatively charged muons.

μ^+/μ^- CHARGE RATIO AT GROUND LEVEL

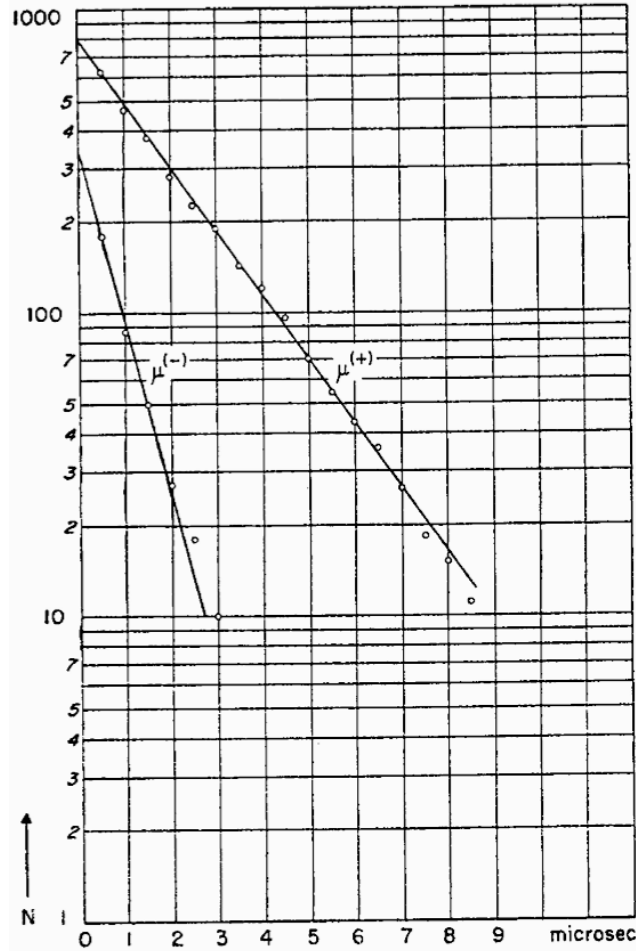


Figure 3 Disintegration curves for positive and negative muons in aluminum. The ordinates at $t = 0$ can be used to determine the relative numbers of negative and positive muons that have undergone spontaneous decay. The slopes can be used to determine the decay time of each charge species. (From Rossi, p168.)

Our measurement of the muon lifetime in plastic scintillator is an average over both negatively and positively charged muons. We have already seen that μ^- s have a lifetime somewhat smaller than positively charged muons because of weak interactions between negative muons and protons in the scintillator nuclei. This interaction probability is proportional to Z^4 , where Z is the atomic number of the nuclei, so the lifetime of negative muons in scintillator and carbon should be very nearly equal. This latter lifetime τ_C is measured to be $\tau_C = 2.0430.003\mu\text{sec}$. [Reiter, 1960]

It is easy to determine the expected average lifetime τ_{obs} of positive and negative muons in plastic scintillator. Let λ^- be the decay rate per negative muon in plastic scintillator and let λ^+ be the corresponding quantity for positively charged muons. If we then let N^- and N^+ represent the number of negative and positive muons incident on the scintillator per unit time, respectively, the observed average decay rate λ_{obs} is given by

$$\begin{aligned}
\lambda_{\text{obs}} &= \frac{N^+ \lambda^+ + N^- \lambda^-}{N^+ + N^-} \\
&= \frac{\rho \lambda^+ + \lambda^-}{1 + \rho}
\end{aligned}$$

where $\rho \equiv N^+/N^-$ is the ratio of positive muons to negative muons. We can define the lifetime of negative muons on the scintillator as $\tau^- \equiv 1/\lambda^-$ and the lifetime of positive muons as $\tau^+ \equiv 1/\lambda^+$. Then the observed average lifetime τ_{obs} is given by

$$\begin{aligned}
\tau_{\text{obs}} &= \lambda_{\text{obs}}^{-1} \\
&= \frac{1 + \rho}{\rho \lambda^+ + \lambda^-} \\
&= (1 + \rho) \left(\frac{1}{\tau^-} + \frac{\rho}{\tau^+} \right)^{-1} \\
&= (1 + \rho) \frac{\tau^- \tau^+}{\tau^+ + \rho \tau^-}
\end{aligned}$$

The lifetime of negative muons in plastic scintillator is approximately the same as the lifetime in carbon, $\tau^- = \tau_C$. On the other hand, positive muons are not captured by the scintillator nuclei. Therefore the lifetime of positive muons, τ^+ , is equal to the lifetime of muons in free space, τ_μ . Setting $\rho = 1$ allows us to estimate the average muon lifetime we expect to observe in the scintillator.

We can *measure* ρ for the momentum range of muons that stop in the scintillator by rearranging the above equation:

$$\rho = -\frac{\tau^+}{\tau^-} \left(\frac{\tau^- - \tau_{\text{obs}}}{\tau^+ - \tau_{\text{obs}}} \right)$$

BACKGROUNDS

The detector responds to any particle that produces enough scintillation light to trigger its readout electronics. These particles can be either charged, like electrons or muons, or neutral, like photons, that produce charged particles when they interact inside the scintillator. Now, the detector has no knowledge of whether a penetrating particle stops or not inside the scintillator and so has no way of distinguishing between light produced by muons that stop and decay inside the detector, from light produced by a pair of through-going muons that occur one right after the other. This important source of background events can be dealt with in two ways. First, we can restrict the time interval during which we look for the two successive flashes of scintillator light characteristic of muon decay events. Secondly, we can estimate the background level by looking at large times in the decay time histogram where we expect few events from genuine muon decay.

FERMI COUPLING CONSTANT G_F

Muons decay via the weak force and the Fermi coupling constant G_F is a measure of the strength of the weak force. To a good approximation, the relationship between the muon lifetime τ and G_F is particularly simple:

$$\tau = \frac{192 \pi^3 \hbar^7}{G_F^2 m^5 c^4}$$

where m is the mass of the muon and the other symbols have their standard meanings. Measuring τ with this instrument and then taking m from, say, the Particle Data Group (<http://www.pdg.lbl.gov>) produces a value for G_F .

TIME DILATION EFFECT

A measurement of the muon stopping rate at two different altitudes can be used to demonstrate the time dilation effect of special relativity. Although the detector configuration is not optimal for demonstrating time dilation, a useful measurement can still be preformed without additional scintillators or lead absorbers. Due to the finite size of the detector, only muons with a typical total energy of about 160 MeV will stop inside the plastic scintillator. The stopping rate is measured from the total number of observed muon decays recorded by the instrument in some time interval. This rate in turn is proportional to the flux of muons with total energy of about 160 MeV and this flux decreases with diminishing altitude as the muons descend and decay in the atmosphere. After measuring the muon stopping rate at one altitude, predictions for the stopping rate at another altitude can be made with and without accounting for the time dilation effect of special relativity. A second measurement at the new altitude distinguishes between competing predictions.

A comparison of the muon stopping rate at two different altitudes should account for the muons energy loss as it descends into the atmosphere, variations with energy in the shape of the muon energy spectrum, and the varying zenith angles of the muons that stop in the detector. Since the detector stops only low energy muons, the stopped muons detected by the low altitude detector will, at the elevation of the higher altitude detector, necessarily have greater energy. This energy difference $E(h)$ will clearly depend on the pathlength between the two detector positions.

Vertically travelling muons at the position of the higher altitude detector that are ultimately detected by the lower detector have an energy larger than those stopped and detected by the upper detector by an amount equal to $E(h)$. If the shape of the muon energy spectrum changes significantly with energy, then the relative muon stopping rates at the two different altitudes will reflect this difference in spectrum shape at the two different energies. (This is easy to see if you suppose muons do not decay at all.) This variation in the spectrum shape can be corrected for by calibrating the detector in a manner described below.

Like all charged particles, a muon loses energy through coulombic interactions with the matter it traverses. The average energy loss rate in matter for singly charged particles traveling close to the speed of light is approximately 2 MeV/g/cm^2 , where we measure the thickness s of the matter in units of g/cm^2 . Here, $s = \rho x$, where ρ is the mass density of the material through which the particle is passing, measured in g/cm^3 , and x is the particles pathlength, measured in cm. (This way of measuring material thickness in units of g/cm^2 allows us to compare effective thicknesses of two materials that might have very different mass densities.) A more accurate value for energy loss can be determined from the Bethe-Bloch equation.

$$\begin{aligned} \frac{dE}{dx} &= -\frac{4\pi N_e^4}{mc^2 \beta^2} z^2 \left(\ln \frac{2mc^2 \beta^2 \gamma^2}{I} - \beta^2 \right) \\ &= 0.3071 \left(\frac{\text{MeV}}{\text{g/cm}^2} \right) \frac{Z}{A} \rho \frac{1}{\beta^2} z^2 \left(\ln \frac{2mc^2 \beta^2 \gamma^2}{I} - \beta^2 \right) \end{aligned}$$

Here N is the number of electrons in the stopping medium per cm^3 , e is the electronic charge, z is the atomic number of the projectile, Z and A are the atomic number and weight, respectively, of the stopping medium. The projectile has a velocity, β , expressed in units of the speed of light, c . Its corresponding Lorentz factor is γ . The mean excitation energy of the stopping medium atoms, I , can be approximates as $I \approx AZ$, where $A \cong 13 \text{ eV}$. More accurate values for I , as well as corrections to the Bethe-Bloch equation, can be found in [Leo, p26].

A simple estimate of the energy lost, ΔE , by a muon as it travels a vertical distance H is

$$\Delta E = 2H \rho_{\text{air}} \frac{\text{MeV}}{\text{gm/cm}^2}$$

where ρ_{air} is the density of air, possibly averaged over H using the density of air according to the *standard atmosphere*. The atmosphere is assumed to be isothermal and the reduced air pressure $p = \bar{P}/g$ at some height h above sea level is given by $p = p_0 e^{-h/h_0}$, where $p_0 = 1030 \text{ g/cm}^2$ is the total thickness of the atmosphere and $h_0 = 8.4 \text{ km}$. The units of reduced pressure are g/cm^2 , which you can see if you recall from hydrostatics that the pressure at the base of a stationary fluid is $P = \rho gh$. Thus the reduced pressure, $p \equiv \bar{P}/g = \rho h$ will have units of g/cm^2 . Then the air density ρ , in familiar units of g/cm^3 , can be expressed as $\rho = -dp/dh$.

In the lab frame, we let t denote the transit time for a particle to travel vertically from some height H down to sea level. In the rest frame of the particle, the corresponding time, t' , is given by

$$t' = \int_H^0 \frac{dh}{c \beta(h) \gamma(h)}$$

Here β and γ have their usual relativistic meanings for the projectile and are measured in the lab frame. Since relativistic muons lose energy at essentially a constant rate when travelling through a medium

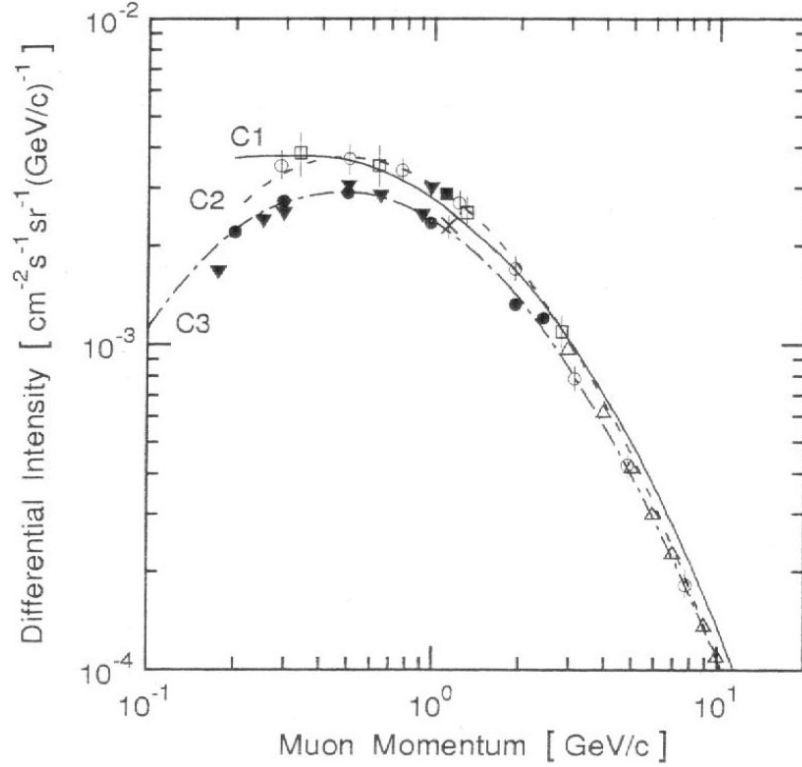


Figure 4 Muon momentum spectrum at sea level. The curves are fits to various data sets. Figure from Greider, p399.

of mass density ρ , $dE/ds = C_0$, so we have $dE = \rho C_0 dh$, with $C_0 = 2$ MeV/(g/cm²). Also, from the Einstein relation, $E = \gamma mc^2$, $dE = mc^2 d\gamma$, so $dh = (mc^2/\rho C_0) d\gamma$. Hence,

$$\begin{aligned} t' &= \frac{mc}{\rho C_0} \int_{\gamma_1}^{\gamma_2} \frac{d\gamma}{\beta\gamma} \\ &= \frac{mc}{\rho C_0} \int_{\gamma_1}^{\gamma_2} \frac{d\gamma}{\sqrt{\gamma^2 - 1}} \end{aligned}$$

Here γ_1 is the muons gamma factor at height H and γ_2 is its gamma factor just before it enters the scintillator. We can take $\gamma_2 = 1.5$ since we want muons that stop in the scintillator and assume that on average stopped muons travel halfway into the scintillator, corresponding to a distance $s = 10$ g/cm². The entrance muon momentum is then taken from range-momentum graphs at the Particle Data Group WWW site and the corresponding γ_2 computed. The lower limit of integration is given by $\gamma_1 = E_1/mc^2$, where $E_1 = E_2 + \Delta E$, with $E_2 = 160$ MeV. The integral can be evaluated numerically. (See, for example, http://people.hofstra.edu/faculty/Stefan_Waner/RealWorld/integral/integral.html)

Hence, the ratio R of muon stopping rates for the same detector at two different positions separated by a vertical distance H , and ignoring

for the moment any variations in the shape of the energy spectrum of muons, is just $R = e^{-t'/\tau}$, where τ is the muon proper lifetime.

When comparing the muon stopping rates for the detector at two different elevations, we must remember that muons that stop in the lower detector have, at the position of the upper detector, a larger energy. If, say, the relative muon abundance grows dramatically with energy, then we would expect a relatively large stopping rate at the lower detector simply because the starting flux at the position of the upper detector was so large, and not because of any relativistic effects. Indeed, the muon momentum spectrum does peak, at around $p = 500 \text{ MeV}/c$ or so, although the precise shape is not known with high accuracy. See figure 4.

We therefore need a way to correct for variations in the shape of the muon energy spectrum in the region from about 160 – 800 MeV. (Corresponding to momentum range $p = 120 - 790 \text{ MeV}/c$.) We do this by first measuring the muon stopping rate at two different elevations ($\Delta h = 3008$ meters between Taos, NM and Dallas, TX) and then computing the ratio R_{raw} of raw stopping rates. ($R_{\text{raw}} = \text{Dallas}/\text{Taos} = 0.41 \pm 0.05$) Next, using the above expression for the transit time between the two elevations, we compute the transit time in the muons rest frame ($t' = 1.32\tau$) for vertically travelling muons and calculate the corresponding theoretical stopping rate ratio $R = e^{-t'/\tau} = 0.267$. We then compute the double ratio $R_0 = R_{\text{raw}}/R = 1.5 \pm 0.2$ of the measured stopping rate ratio to this theoretical rate ratio and interpret this as a correction factor to account for the increase in muon flux between about $E = 160 \text{ MeV}$ and $E = 600 \text{ MeV}$. This correction is to be used in all subsequent measurements for any pair of elevations.

To verify that the correction scheme works, we take a new stopping rate measurement at a different elevation ($h = 2133$ meters a.s.l. at Los Alamos, NM), and compare a new stopping rate ratio measurement with our new, corrected theoretical prediction for the stopping rate ratio $R_\mu = R_0 R = 1.6e^{-t'/\tau}$. We find $t' = 1.06\tau$ and $R_\mu = 0.52 \pm 0.06$. The raw measurements yield $R_{\text{raw}} = 0.56 \pm 0.01$, showing good agreement.

For your own time dilation experiment, you could first measure the raw muon stopping rate at an upper and lower elevation. Accounting for energy loss between the two elevations, you first calculate the transit time t' in the muons rest frame and then a naïve theoretical lower elevation stopping rate. This naïve rate should then be multiplied by the muon spectrum correction factor 1.5 ± 0.2 before comparing it to the measured rate at the lower elevation. Alternatively, you could measure the lower elevation stopping rate, divide by the correction factor, and then account for energy loss before predicting what the upper elevation stopping rate should be. You would then compare your prediction against a measurement.

ELECTRONICS

A block diagram of the readout electronics is shown in figure 5. The logic of the signal processing is simple. Scintillation light is detected by a photomultiplier tube (PMT) whose output signal feeds a two-stage amplifier.

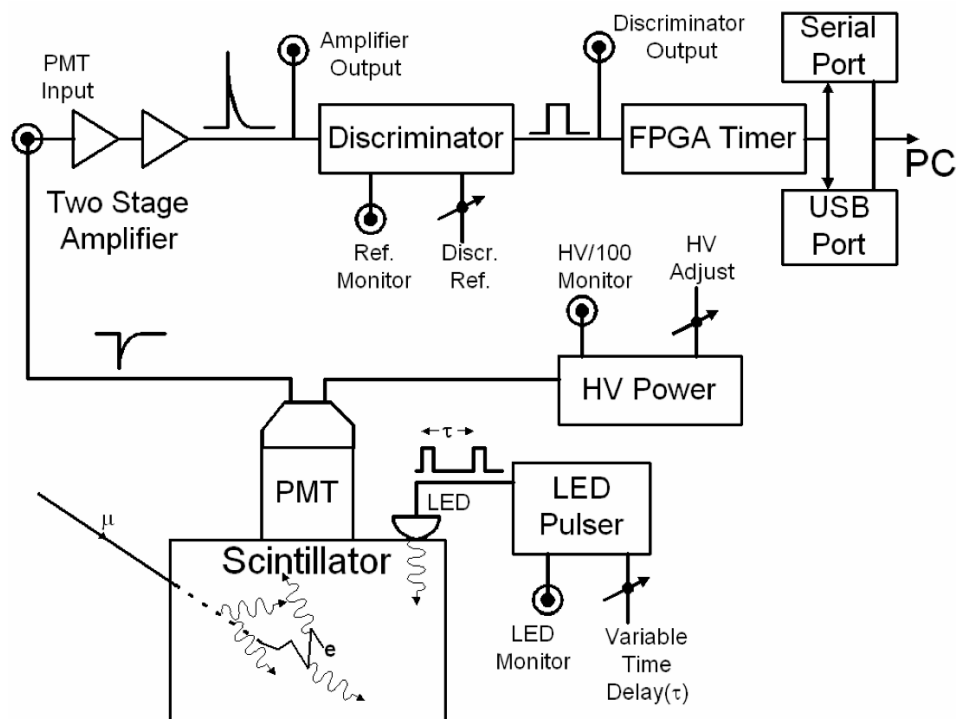


Figure 5 Block diagram of the readout electronics. The amplifier and discriminator outputs are available on the front panel of the electronics box. The HV supply is inside the detector tube.

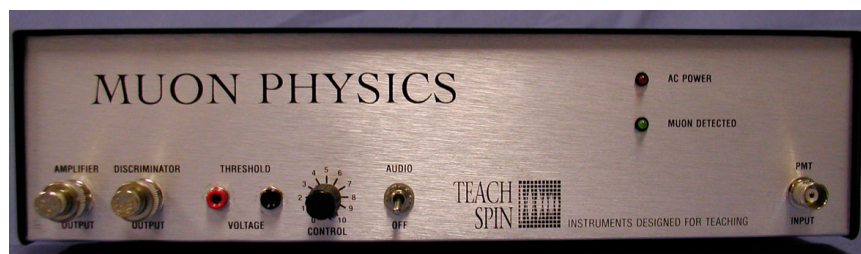


Figure 6 Front of the electronics box.

The output from the amplifier is then input into a voltage discriminator with an adjustable threshold. If the input signal is above the discriminator threshold, then the discriminator outputs a TTL (5 V) logic pulse that triggers the timing circuit of the FPGA. If, within a fixed time interval, the FPGA receives a second TTL pulse, then the timing circuit will stop. The time interval between the start and stop timing pulses is temporarily stored in a data buffer and then the timing circuit is reset. (The reset takes about 1 msec during which the detector is disabled.) The data stored in the buffer is sent to the PC via the communications module. This data is used to determine the muon lifetime. If a second TTL pulse does not arrive within the fixed time interval, the timing circuit stores the value of this time interval in the data buffer and then automatically resets the circuit for the next measurement.



Figure 7 Rear of electronics box. The communications ports are on the left. Use only one.

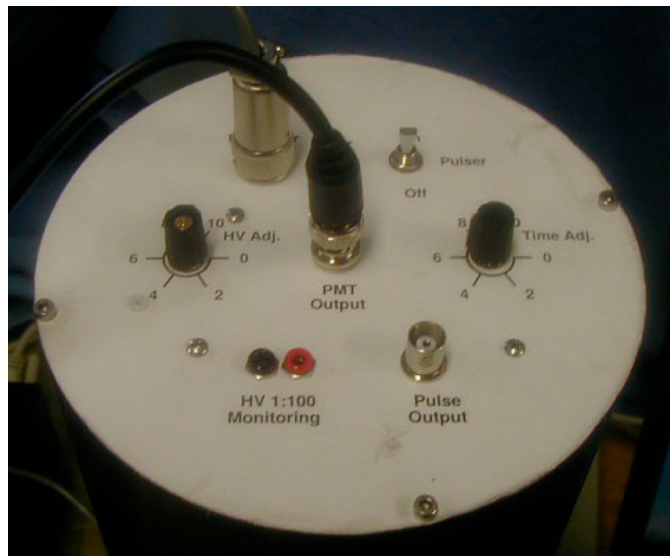


Figure 8 Rear of electronics box. The communications ports are on the left. Use only one.

The back panel of the electronics box is shown in figure 7. An extra fuse is stored inside the power switch.

The front panel of the electronics box is shown in figure 6. The amplifier output is accessible via the BNC connector labeled *Amplifier output*. Similarly, the discriminator output is accessible via the connector labeled *Discriminator output*. The voltage level against which the amplifier output is compared to determine whether the discriminator triggers can be adjusted using the *Threshold control* knob. The threshold voltage is monitored by using the red and black connectors that accept standard multi-meter probe leads. The toggle switch controls a beeper that sounds when an amplifier signal is above the discriminator threshold. The beeper can be turned off.

Figure 8 shows the top of the detector cylinder. DC power to the electronics inside the detector tube is supplied from the electronics box through the connector *DC Power*. The high voltage (HV) to the PMT can be adjusted by turning the potentiometer located at the top of the detector tube. The HV level can be measured by using the pair of red and



Figure 9 Output pulse directly from PMT into a $50\ \Omega$ load. Horizontal scale is 20 ns/div and vertical scale is 100 mV/div.

black connectors that accept standard multimeter probes. The HV monitor output is 1/100 times the HV applied to the PMT.

A pulser inside the detector tube can drive a light emitting diode (LED) imbedded in the scintillator. It is turned on by the toggle switch at the tube top. The pulser produces pulse pairs at a fixed repetition rate of 100 Hz while the time between the two pulses comprising a pair is adjusted by the knob labeled *Time Adj.* The pulser output voltage is accessible at the connector labeled *Pulse Output*.

For reference, figure 9 shows the output directly from the PMT into a $50\ \Omega$ load. Figure 10 shows the corresponding amplifier and discriminator output pulses.

DATA ACQUISITION

To acquire the data, you will use the program *Muon*. A link to the program can be found on the desktop of the computer located on the lab table. This program is rather basic. It collects the data, validates it, and saves it to a file. *Muon* can provide an initial estimate of the time constant (muon lifetime) of the collected data. However, to perform a proper calculation of the lifetime and its uncertainty, you will need to analyze the data in MatLab.

When you launch *muon* you will bring up the main window as shown in Figure 11. Within the window are five sections:

- Control
- Muon Decay Time Histogram

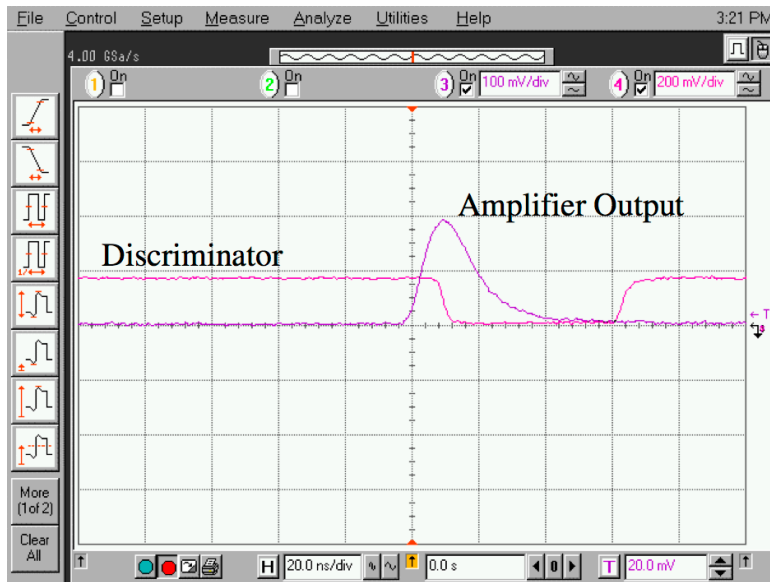


Figure 10 Amplifier output pulse from the input signal from figure 9 and the resulting discriminator output pulse. Horizontal scale is 20 ns/div and the vertical scale is 100 mV/div (amplifier output) and 200 mV/div (discriminator output).

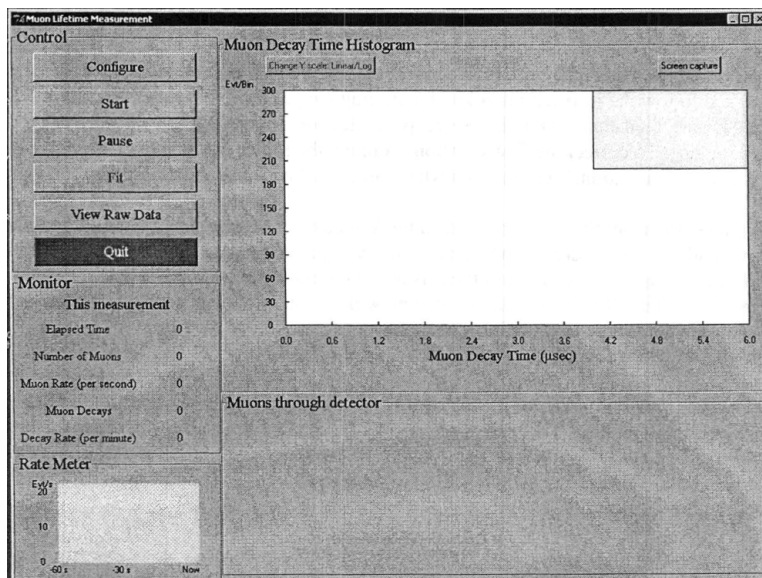


Figure 11 Main window of the data acquisition program.

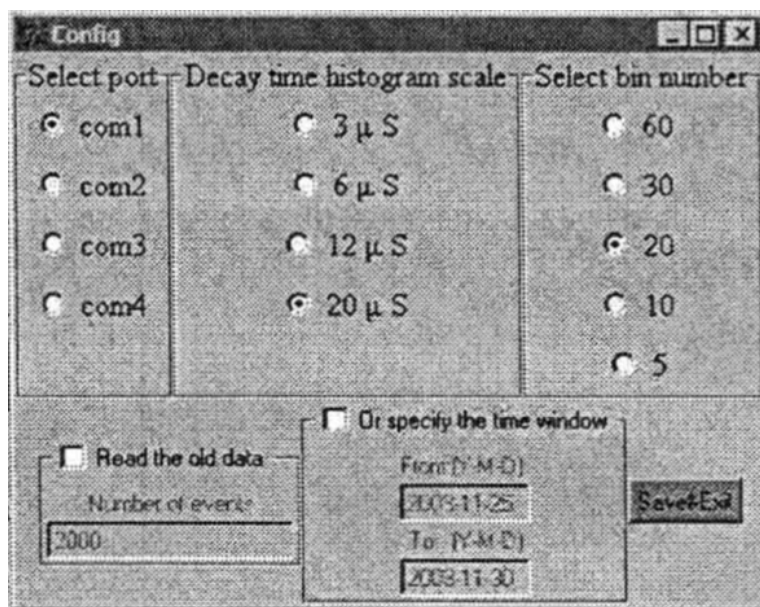


Figure 12 Muon's Configure Pane.

- Monitor
- Rate Meter
- Muons Through Detector

Control

Before data collection can begin, the system must be configured to know on which port the data will be received and how it should be histogrammed. Unfortunately these settings are not stored and must to reset with each run. (At this point, you may better understand the description *basic* that I applied to the program.)

By clicking the *Configure* button, one opens the *Configure* pane as shown in Figure 12. Select the port, which in our case, since we use USB, is com2. The program does not properly store this data and sometimes resets the port to com1. If, during the experiment, you are not receiving data, check that the correct port is selected.

Select the range and number of bins in the histogram. During experimental set-up, good values to choose are $6\ \mu\text{s}$ and 10 bins. However, for long data collection runs, choose $20\ \mu\text{s}$ and 60 bins. The greater resolution provides a better check on the accuracy of your data. Close the Configure pane, the other sections of the pane are not relevant to data acquisition.

The *Start* and *Pause/Resume* buttons do exactly what their names imply.

The *Fit* button provides a crude fit to the data. Use this only when the data acquisition is paused, since the collection of new data will erase

the fit. (N.B. The fit routine use in *Muon* has had stability problems. It's results are not sufficient for your lab report.)

The *View Raw Data* button opens a window that allows you to display the timing data for a user selected number of events, with the most recent events read in first. Here an event is any signal above the discriminator threshold so it includes data from both through going muons as well as signals from muons that stop and decay inside the detector. Each raw data record contains two fields of information. The first is a time, indicating the year, month, day, hour, minute and second, reading left to right, in which the data was recorded. The second field is an integer that encodes two kinds of information. If the integer is less than 40000, it is the time between two successive flashes, in units of nanoseconds. If the integer is greater than or equal to 40000, then the units position indicates the number of "time outs," (instances where a second scintillator flash did not occur within the preset timing window opened by the first flash). See the data file format below for more information. Typically, viewing raw data is a diagnostic operation and is not needed for normal data taking.

The *Quit* button stops the measurement and asks you whether you want to save the data. The questions seems rather straightforward, but the resulting action is as counterintuitive as pressing START to shutdown a computer. Answering *NO* writes the data to a file that is named after the date and time the measurement was originally started, i.e., 03-07-13-17-26.data. Answering *YES* appends the data to the file muon.data. The file muon.data is intended as the main data file. I recommend performing a test run to see what happens in both cases. As I prefer having the datafile name include the data, I often answer *NO*. Others prefer choosing *YES*, then renaming the filename.

Data File Format

Timing information about each signal above threshold is written to the data file. The file is in ascii format and can be read by any text editor. The first data field provides information on the time between successive pulses. If that integer is less than 40000, then the value corresponds to the time, to the nearest nanosecond, between successive signals. If the value is greater than or equal to 40000 then the timing circuit exceeded its maximum number of clock cycles before receiving a second signal. The circuit *timed-out*. The number in the units place indicates the number of times the circuit timed-out before receiving a signal. For example, 40005 corresponds to a circuit that was triggered but then timed-out 5 times before a subsequent signal arrived.

The second field is the computer time in seconds. This value is typically the number of seconds since 1 January 1970, a date convention used in C and in Unix.

Monitor

This panel shows rate-related information for the current measurement. The elapsed time of the current measurement is shown along with the

accumulated number of times from the start of the measurement that the readout electronics was triggered (Number of muons). The Muon Rate is the number of times the readout electronics was triggered in the previous second. The number of pairs of successive signals, where the time interval between successive signals is less than the maximum number of clock cycles of the timing circuit, is labeled Muon Decays, even though some of these events may be background events and not real muon decays. Finally, the number of muon decays per minute is displayed as Decay Rate.

Rate Meter

This continuously updated graph plots the number of signals above discriminator threshold versus time. It is useful for monitoring the overall trigger rate.

Muons through Detector

This graph shows the time history of the number of signals above threshold. Its time scale is automatically adjusted and is intended to show time scales much longer than the rate meter. This graph is useful for long term monitoring of the trigger rate. Strictly speaking, it includes signals from not only through going muons but any source that might produce a trigger. The horizontal axis is time, indicated down to the second. The scale is sliding so that the far left-hand side always corresponds to the start of the measurement session. The bin width is indicated in the upper left-hand portion of the plot.

Muon Decay Time Histogram

This plot is probably the most interesting one. It is a histogram of the time difference between successive triggers and is the plot used to measure the muon lifetime. The horizontal scale is the time difference between successive triggers in units of microseconds. Its maximum displayed value is set by the *Configure* menu. (All time differences less than 20 μsec are entered into the histogram but may not actually be displayed due to menu choices.) You can also set the number of horizontal bins using the same menu. The vertical scale is the number of times this time difference occurred and is adjusted automatically as data is accumulated. A button (*Change y scale Linear/Log*) allows you to plot the data in either a linear-linear or log-linear fashion. The horizontal error bars for the data points span the width of each timing bin and the vertical error bars are the square root of the number of entries for each bin. The upper right hand portion of the plot shows the number of data points in the histogram. Again, due to menu selections not all points may be displayed. If you have

selected the *Fit* button then information about the fit to the data is displayed. The muon lifetime is returned, assuming muon decay times are exponentially distributed, along with the chi-squared per degree of

freedom ratio, a standard measure of the quality of the fit. (See Bevington for more details.)

A *Screen capture* button allows you to produce a plot of the display. Select the button and then open the *Paint* utility (in Windows) and paste the graphic into a blank document window.

PROCEDURE

Getting Started

The black aluminum cylinder (detector) is placed in the wooden pedestal for convenience; it will work in any orientation. Connect the power cable and signal cable between the electronics box and the detector. Connect the USB cable between the back of the electronics box and the PC. If USB is not available then use the serial cable. **DO NOT USE BOTH.**

Turn on power to the electronics box. The red LED indicates that the power is on. The green LED indicates an event in the PMT.

Set the HV between 1100 and 1200 Volts using the knob at the top of the detector tube. The exact setting is not critical and the voltage can be monitored by using the multimeter probe connectors at the top of the detector tube.

If you are curious, you can look directly at the output of the PMT using the *PMT Output* on the detector tube and an oscilloscope. (A digital scope works best.) **Be certain to terminate the scope input at 50 Ω or you signal will be distorted.** You should see a signal that looks like Figure 9. The figure shows details like scope settings and trigger levels.

Connect the BNC cable between *PMT Output* on the detector and *PMT Input* on the box. Adjust the discriminator setting on the electronics box so that it is in the range 180–220 mV. The green LED on the box front panel should now be flashing.

You can look at the amplifier output by using the *Amplifier Output* on the box front panel and an oscilloscope. The scope input impedance must be 50 Ω . Similarly, you can examine the output of the discriminator using the *Discriminator Output* connector. Again, the scope needs to be terminated at 50 Ω . Figure 10 shows typical signals for both the amplifier and discriminator outputs on the same plot. Details about scope time settings and trigger thresholds are on the plot.

Launch the program *Muon* which is located on the desktop of the PC. Configure the program using the *Configure* button to set the data port and the histogram dimensions.

Click on *Start*. You should see the rate meter at the lower left-hand side of your computer screen immediately start to display the raw trigger rate for events that trigger the readout electronics. The mean rate should be about 6 Hz or so.

Lab Exercises

1. Measure the gain of the 2-stage amplifier using a sine wave. Apply a 100kHz 100mV peak-to-peak sine wave to the input of the electronics box input. Measure the amplifier output and take the ratio V_{out}/V_{in} .
Increase the frequency. How good is the frequency response of the amp?
Estimate the maximum decay rate you could observe with the instrument.
2. Measure the saturation output voltage of the amp. Increase the magnitude of the input sine wave and monitor the amplifier output. Does a saturated amp output change the timing of the FPGA? What are the implications for the size of the light signals from the scintillator?
3. Examine the behavior of the discriminator by feeding a sine wave to the box input and adjusting the discriminator threshold. Monitor the discriminator output and describe its shape.
4. Measure the timing properties of the FPGA:
 - (a) Using the pulser on the detector, measure the time between successive rising edges on an oscilloscope. Compare this number with the number from software display.
 - (b) Measure the linearity of the FPGA: Alter the time between rising edges and plot scope results v. FPGA results; Can use time between $1\ \mu s$ and $20\ \mu s$ in steps of $2\ \mu s$.
 - (c) Determine the timeout interval of the FPGA by gradually increasing the time between successive rising edges of a double-pulse and determine when the FPGA no longer records results. What does this imply about the maximum time between signal pulses?
 - (d) Decrease the time interval between successive pulses and try to determine/bound the FPGA internal timing bin width. What does this imply about the binning of the data? What does this imply about the minimum decay time you can observe?
5. Adjust (or misadjust) discriminator threshold. Increase the discriminator output rate as measured by the scope or some other means. Observe the raw muon count rate and the spectrum of "decay" times. (This exercise needs a digital scope and some patience since the counting rate is slowish.)
6. What HV should you run at? Adjust/misadjust HV and observe amp output. (We know that good signals need to be at about 200 mV or so before discriminator, so set discriminator before hand.) With fixed threshold, alter the HV and watch raw muon count rate and decay spectrum.
7. Connect the output of the detector can to the input of the electronics box. Look at the amplifier output using a scope. (A digital scope works best.) Be sure that the scope input is terminated at $50\ \Omega$. What do you see? Now examine the discriminator output simultaneously. Again, be certain to terminate the scope input at $50\ \Omega$. What do you see?

8. Set up the instrument for a muon lifetime measurement. Start and observe the decay time spectrum. The muons whose decays we observe are born outside the detector and therefore spend some (unknown) portion of their lifetime outside the detector. So, we never measure the actual lifetime of any muon. Yet, we claim we are measuring the lifetime of muons. How can this be?
9. Write a MatLab routine to fit the decay time histogram and the background.
10. From your measurement of the muon lifetime and a value of the muon mass from some trusted source, calculate the value of Fermi coupling constant G_F . Compare your value with that from a trusted source.
11. Using the approach outlined in the text, measure the charge ratio ρ of positive to negative muons at ground level or at some other altitude.

Gamma Ray Attenuation in Matter

by S. Penn. Experiment designed by L. Campbell

INTRODUCTION

Observing the absorption and scatter of radiation as it passes through matter is a commonly used technique to study the structure of a material.

- To observe a defect in glass or other transparent material, a person holds the object up to the light to see where the light is scattered.
- By measuring the changes in the polarization of transmitted light one can detect stress in a transparent material.
- To determine the composition or density of a material one measures the absorption of transmitted light as a function of frequency.

This principle is the basis for observations as seemingly disparate as perceiving the color of glass, to analyzing the composition of a gaseous nebula, to detecting a bone fracture through X-ray photography.

Typically observations of the passage of radiation through matter allow us to examine the structure of the matter by noting changes in the transmitted radiation. The underlying premise of this technique is that the radiation passes through the material in a manner that is uniform both spatially and within the observed frequency range of the light. Thus any change to the nature of the light (i.e. loss of intensity, shifts in polarization or direction) are due to the material. However, as one looks across the entire electromagnetic frequency range the processes by which radiation interacts with matter can change. Materials that are transparent at one frequency range are opaque in another. One learns about the fundamental interactions of radiation with matter by studying these changes in the attenuation of light as a function of frequency.

One of the most interesting frequency regimes for this type of study is the low energy gamma ray spectrum. In that region light transitions through three principle mechanisms for interacting with matter: the photoelectric effect, Compton scattering, and particle-antiparticle pair production. Moreover, this transition in scattering mechanisms also marks a

change from processes that depended on the structure of the atom (photoelectric), to ones that depend only on the density of electrons (Compton) or of nuclei (pair production).

Attenuation

Before we discuss these mechanisms in detail we must derive a method for characterizing the attenuation of radiation by matter. As radiation passes through matter, it can be absorbed, scattered, or transmitted. Light that is attenuated is light that is either absorbed or scattered. The amount of attenuation is dependent on the target's composition and the amount of target material. Light of different frequencies will interact with objects of varying size and binding energy. Radio waves are attenuated by molecules; visible light is attenuated by atoms; and high energy gamma rays by nuclei. Nevertheless, the fundamentals are the same. The rate of attenuation depends on the total number of possible absorption and scattering sites for photons of a given frequency.

We observe the attenuation by sending a beam of radiation through the material and noting the rate of photons transmitted. For a uniform density target material, the amount of attenuation will depend on the thickness of the target. Let us consider the change in the rate of photons passing through a thin slice of the target material. In this case:

$$dR = -R P_A ds$$

where R is the rate of incident photons, P_A is the probability of attenuation per thickness of material, ds . The change in photon rate, dR , is negative because the photons are being attenuated. Solving for the photon rate as a function of target thickness, we get

$$R(s) = R_0 e^{-P_A s} = R_0 e^{-ks} \quad (1)$$

where R_0 is the incident photon rate and s is the target thickness. The probability of attenuation per distance, P_A is commonly labeled k , the *linear attenuation coefficient*. As noted above, k will depend on the target composition and density.

Cross Section

There are three mechanisms by which low energy gamma rays are attenuated: photoelectric effect, Compton scattering, and pair production. In the photoelectric effect, the gamma ray is absorbed by an electron that then transitions to a free state. In other words, the gamma ray ejects the electron from the atom. Since the transition is to a free state, it is not strongly energy dependent like the transition to a bound state. For photons with energies greater than the electron's binding energy, the probability depends only on the number of electrons.

In Compton scattering, first described by A. H. Compton in 1922, the photon scatters from a free, or nearly free, electron. The photon imparts some of its energy to the electron and as a result the photon wavelength and direction both shift in such a way as to conserve energy and

momentum. To think in the manner of quantum electrodynamics, the photon is absorbed by the electron which momentarily transitions to an excited virtual state before radiating a photon and decaying to a free particle state. As the energy of gamma rays increases, the binding energy of atomic electrons becomes comparatively smaller and all atomic electrons appear essentially free. Therefore, in that energy regime, the probability of Compton scattering will depend primarily on the number of electrons in the target.

Moving up in gamma ray energy, we cross the threshold of 1.022 MeV, twice the rest energy of the electron. At this energy the gamma ray can decay into an electron-positron pair, thus transforming the photon energy into the rest mass and kinetic energy of the particle pair. However, in order to conserve momentum, this reaction must occur near other particles with which there must be an exchange of the excess momentum. At these energies the reaction typically occurs within the potential well of the nucleus and is therefore dependent on the number of protons. Thus all three processes depend on the atomic number Z , which is the number of protons in a given element (or the number of electrons if the atom is neutral, as in our case).

We defined $P_A s$ as the probability of attenuation of a beam of photons passing through a target of thickness s . However we can also write this probability as

$$P_A s = k s = P_N N$$

where N is the number of possible scattering or absorption sites and P_N is the probability of attenuation for a given site. If we define the number density, $\mathcal{N} = N/V$, as the number of attenuation sites per volume, then

$$k s = P_N N = P_N \mathcal{N} s \mathcal{A}$$

where $\mathcal{A}$ is the area that the beam samples. But when we are considering the possible attenuation of a photon through matter what is $\mathcal{A}$? What area does a single photon sample? What if we were to make our input beam more diffuse by increasing the beam area without changing the number of photons? The area of the beam would have increased, but the number of incident photons and the number of transmitted photons would remain the same. Thus $\mathcal{A}$ is not the beam area but the effective interaction area of the photon. It seems clear that $\mathcal{A}$ depends on the photon parameters and possibly the target's composition, but its exact expression is left undefined. Instead we combine this effective area with the probability of interaction into a single concept called the cross section, $\sigma = P_N \mathcal{A}$. The cross section therefore can be thought of as the effective area of attenuation for a given incident particle interacting with a given target material. A useful way of visualizing the cross section is to imagine an opaque area the size of the cross section centered at each attenuation site; if the photon hits one of these areas, it is attenuated. Indeed, the concept of cross section came into use when physicists were first performed nuclear scattering measurements. The target was uranium, whose nucleus was so large that scattering off of it was said to be as easy as "hitting the broad side of a barn". Hence the unit used for cross section is the *barn* (b), which is 10^{-28} m^2 , the approximate cross sectional area of the uranium nucleus.

Finally the number density, $\mathcal{N}$ can be calculated from the target pa-

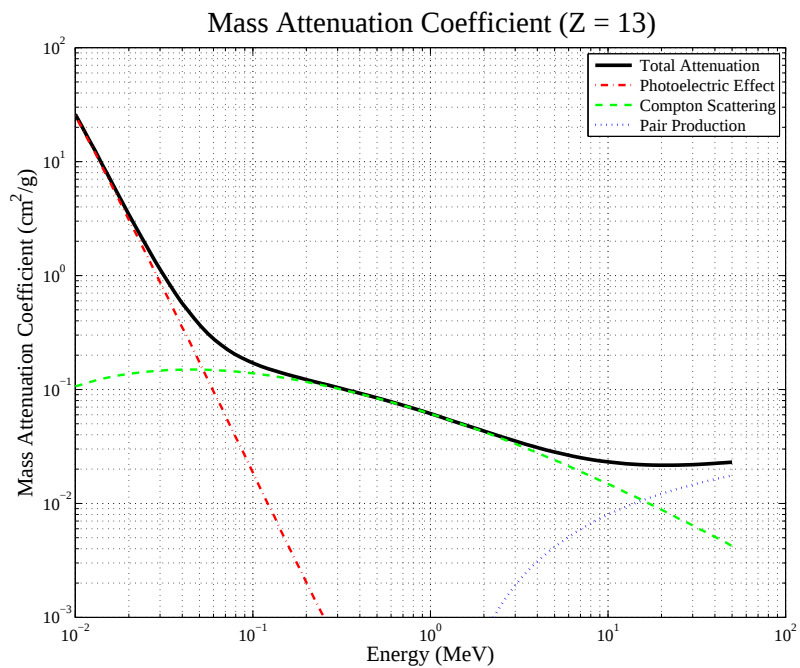


Figure 1 Mass Attenuation coefficient for Aluminum, $Z = 13$

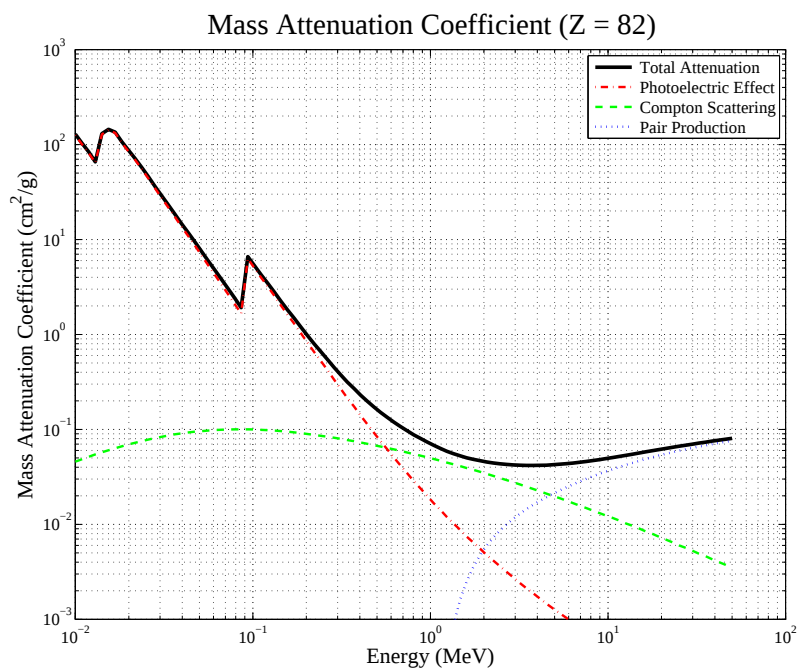


Figure 2 Mass Attenuation coefficient for Lead, $Z = 82$

rameters

$$\begin{aligned}\mathcal{N} &= (\text{mass density})(\text{atoms per unit of mass})(\text{atomic number}) \\ &= \rho(N_A/A)Z \\ &= \rho N_A Z/A\end{aligned}$$

where N_A is Avagadro's number, A is the atomic weight, and Z is the atomic number.

Therefore the attenuation probability becomes

$$\begin{aligned}ks &= \mathcal{N}s\sigma \\ &= (\rho N_A Z/A)\sigma s \\ &= (\sigma N_A Z/A)(\rho s) \\ &= k's'\end{aligned}$$

where $s' = s\rho$ is the areal mass density. Finally we define the *mass attenuation coefficient*, k' , as

$$k' = k/\rho = \sigma N_A Z/A$$

The mass attenuation coefficient is a very useful concept because it is an easily calculated value that parameterizes the attenuation but is approximately constant across target materials.

While the cross section, σ depends on the photon energy it has little dependence on the target material. The attenuation is primarily a photon-particle interaction that is independent of the target's atomic structure. Likewise Z/A varies little (0.4 – 0.5) for most elements, and N_A is a constant. Thus k' approximately scales with the cross section and is target independent.

The goal of this experiment is to show that the measured mass attenuation coefficient describes the energy dependence of the gamma ray reaction mechanisms and that k' is relatively uniform across target materials. Compare the similarities between the mass attenuation coefficients of aluminum, $Z = 13$, and lead, $Z = 82$, as shown in Figures 1 & 2.

To measure k' one starts with gamma ray sources at a variety of energies. Then the transmission rate is measured through a set of absorbers of various areal densities. Using the photon rate equation

$$R(s') = R_0 e^{-k's'} \quad (2)$$

one can fit the decaying exponential to get k' for a given photon energy.

EXPERIMENTAL SET-UP

To measure the mass attenuation coefficient as a function of energy one needs a source of gamma rays at known energies, a set of calibrated absorbers, and a detector for measuring the transmitted gamma rays.



Figure 3 High Purity Germanium Detector (Oxford Instruments #CNVDS30-20195)

The gamma rays are counted using a high purity germanium (HPGe) detector from Oxford Instruments (#CNVDS30-20195), which is a very low noise detector with an energy sensitivity range of 3 keV – 10 MeV. (See Fig. 3) The detector consists of a piece of single crystal, high purity germanium that is machined into the shape of an inverted cup. Over the inner and outer surfaces of the cup are coated thin conducting electrodes and a voltage bias of -3.6 kV is placed across the crystal.

The largest section of the detector is the liquid nitrogen (LN₂) dewar. The liquid nitrogen is used to keep the germanium crystal at 77 K. At this temperature all of the semiconductor's electrons are in the valence band and the probability for thermal excitation of an electron into the conduction band is quite low. If the crystal were at room temperature, the dark current of thermally activated electrons would greatly reduce the detector sensitivity.

When a gamma ray enters the material it undergoes the three reaction mechanisms (photoelectric, Compton scattering, and pair production) converting all its energy into the kinetic energy of electrons. After

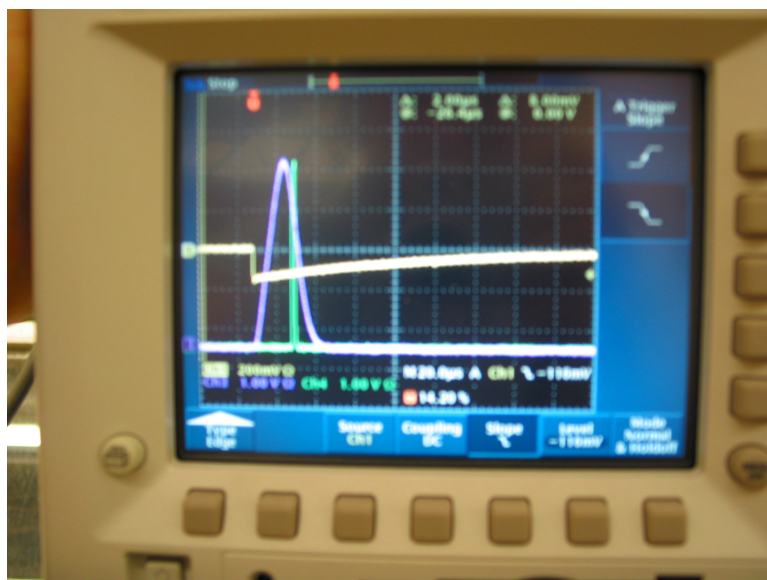


Figure 4 Oscilloscope traces of Preamp output (yellow), Unipolar shaped output (pink) and Pile-up Rejected output (green)

a few collisions these electrons have *thermalized* at the conduction band energy. Thus the energy deposited by the gamma ray is converted into a group of uniform energy electrons and the number of electrons scales linearly with the gamma ray energy. The electrons collected on the detector plates are a current pulse that the detector preamp converts to a voltage pulse. A typical output from the preamp is shown as the yellow trace in Figure 4.

The integral of the voltage pulse is a measure of the gamma ray energy. To produce a spectrum of the gamma ray energies we employ a multi-channel analyzer (MCA). This device integrates the incoming pulses and produces a histogram of the number of events for each energy. MCAs were originally dedicated, stand-alone instruments that were the size of a large desktop computer. However, as electronics have gotten smaller and faster, MCAs have greatly reduced in size while increasing their speed and resolution. In our set-up the MCA is located on a card inside computer and the energy spectrum is displayed in a window on the screen.

(Note that some MCAs integrate the pulse and some measure the pulse height. If the pulses produced by the detector are all of similar shape, these methods are equivalent. Our MCA records the pulse height.)

The critical component in an MCA is the Analog-to-Digital Converter (ADC) which converts the pulse into a series of digital values to represent the pulse. The problem with ADC's is their speed and accuracy in performing this conversion. Slow ADCs lose information while non-linear ADCs distort the information. Unfortunately when our MCA was produced in the mid-1990's ADCs were of lower quality than they are today. To get around this problem, one places between the preamp and the MCA a shaping amplifier. This device amplifies the signal and it uses a capacitive circuit with selectable time constant to reshape the pulse. The



Figure 5 High Voltage Power Supply (TC 950A) and the Analog Shaping Amplifier (TC 244)

amplifier output is shown as the pink trace in Figure 4.

When two gamma rays that arrive in the detector within one pulse width of each other, it is called a pile-up event. The signals from the two events will overlap and will be difficult to disentangle. Fortunately the Amplifier has a *pile-up rejection* (PUR) circuit that senses these events. The circuit produces an pulse output that can be sent to the MCA which is free of such events. The PUR output is shown as the green trace in Figure 4.

In Figure 6 is shown the gamma source holder and the absorber stack, both of which are housed in the small white plastic house that sits above the detector. At the top is a small metal cylinder which has a recessed hole into which one places the gamma sources. There are eight gamma sources. Each source produces gamma rays at from two to several different energies. On the inside cover of the source container is a list of the sources and their respective gamma ray energies. You should choose a few sources so that the spectrum from all the sources sample the range of energies up to about 1.1 – 1.3 MeV.

The absorbers can be found in two boxes stored on shelves adjacent to the sources. Each box contains a list of the areal density of each sources. The sources from the metal box are labeled with a nominal value close to the actual areal density. **DO NOT USE THIS NOMINAL VALUE IN YOUR CALCULATIONS.** The actual areal densities are provided. When using the absorbers feel free to combine absorbers in the four slots of the



Figure 6 Source holder (top cylinder) and the Absorber stack

absorber stack in order to get the total areal density that you desire.

PROCEDURE

First, please be aware that the cryostat must not be allowed to exhaust its liquid nitrogen until the experiment is completely over. If the nitrogen runs out, the detector must be allowed to come to room temperature and remain there for 48 hours before cooling again. We will fill the dewar at least once per week. Those fills will likely be on Monday or Wednesday. Rick will need to turn off the high voltage when he fills. Therefore, please leave a sign on the computer whenever you are taking data stating when the run will complete.

Second, plan out your data runs. This experiment consists of lots of long data taking runs. Draw up a schedule to ensure that you will finish in time.

Third, the detector operates at -3.6 kV. Whenever you raise this voltage you must do so at a rate of LESS THAN 100 V/s. Fast voltage increases damage the detector. When raising the voltage look at the preamp output on the oscilloscope. On the 500mV/div scale you should be able to keep the trace on the oscilloscope while you raise the voltage.

Fourth, Your first run should be a background run without sources. This should be a long run (40 ks of live time). You will subtract a scaled

version of the background from all other measurements so this spectrum should be done well.

Fifth, When taking you data take all data with a given source at once. Avoid swapping out samples. You data will be much cleaner if the source is always in the same position.

Finally use the same absorber set for each source.

When the data is collected it should be stored on you M: drive and shipped upstairs for analysis, which we will cover in the next section.