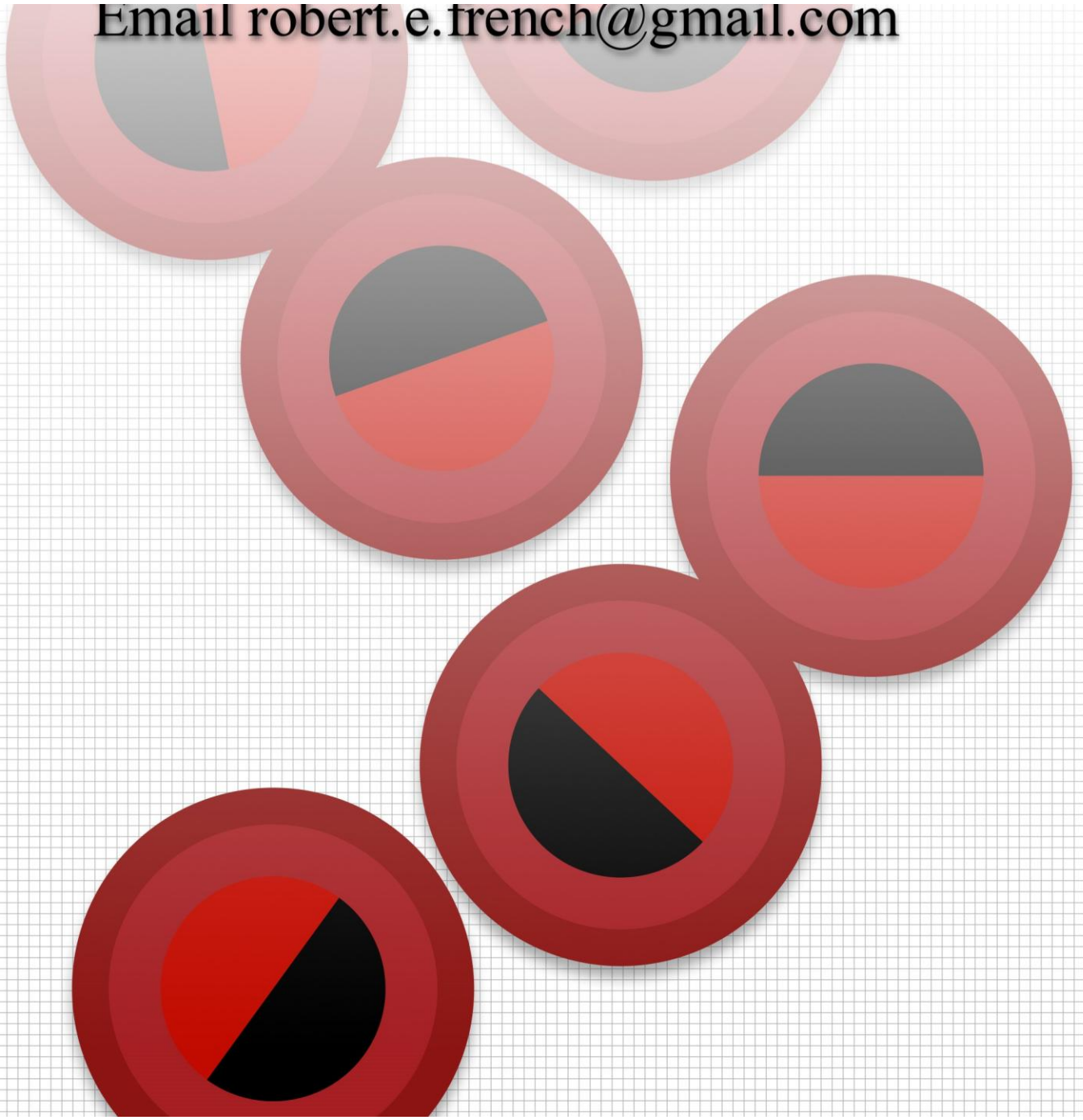


A Model of the Physical World
A Speculative Essay in Natural Philosophy

By
Robert E. French

Email robert.e.french@gmail.com



Copyright © 2022 by Robert E. French

ISBN 978-0-692-75314-9

Email robert.e.french@gmail.com

I would like to thank the following for having read the manuscript and for their valuable suggestions – John French, Harry French, Robin Siegel, Edward Dowdye, Yanhua Shih, Oleg Varnavski, Ernie Kent, David McGraw, Harvey Scribner, Ted Goodson, Hüseyin Yilmaz and Harry Ricker. I also wish to thank John Grimes for helping with production.

Table of Contents for
A Model of the Physical World
A Speculative Essay in Natural Philosophy

Introduction	1
Ch. 1 The Concept of a Field	7
1.1 Force Fields versus Energy Density Fields	8
1.2 Ideal Liquids	12
1.3 Parallel Subspaces	16
1.4 Thin Shells	21
1.5 Wave Particle Unity	29
Ch. 2 Electromagnetism	34
2.1 The Electric Field	35
2.2 The Magnetic Field	42
Ch. 3 Light	57
3.1 An Electromagnetic Account of Light	57
3.2 Feynman Path Integrals	72
3.3 Entanglement	79
3.4 Delayed-Choice and Quantum Eraser Experiments	86
Ch. 4 Atomic Physics	89
4.1 Electron Orbitals	92
4.2 Spin	99
4.3 The Chemical Bond	107
4.4 Reduction of Wave Packets	109
Ch. 5 Gravity	117
5.1 Mass and Charge	118
5.2 Gravity as a Residual Effect of Electromagnetism	122
5.3 Discussion of Experimental Evidence	133
Appendix A Are Complex Numbers Essential to Quantum Mechanics?	135
Appendix B Nuclear Structures	140
Bibliography	144

Introduction

In this book, I present a model of the physical universe apart from our experiences of that universe. Thus, the book is an essay on what has traditionally been termed "natural philosophy.". Both because my approach is novel, and since one feature of overall truth is at least consistency with established empirical facts (across physics and not just in one domain) I cover more than is traditional in a book like this. In particular, I cover a wide variety of topics ranging from electromagnetism and gravity to optics and to atomic physics. This coverage is only within the confines of my approach though, and thus what I say on any of these topics is by no means comprehensive. Also, I do not treat such fields as particle physics (including how there can be the fundamentally stochastic processes associated with radioactivity) and quantum field theory or even how there can be interference effects among individual atoms or resonance effects between different molecular structures. Of course, if there is something for my approach these subjects would have to be eventually treated but I do not attempt to do so in the book. In any event, the aim of the book is to make a contribution to our understanding of the nature of the physical universe with respect to atomic physics, electromagnetism, optics, and gravity, and even these treatments are incomplete. Hopefully someone else can build on what I say to make the treatments more complete and to address other issues.

My approach invokes special reference frames (those of source particles) and thus is not based on relativity since this holds that the basic laws of physics are invariant across all inertial reference frames. To at least partially motivate the need for something like my approach it should be emphasized that even the special theory of relativity remains controversial among a sizable number of physicists (see Herbert Dingle, 1972, and Petr Beckmann, 1987) in spite of having been given a dominant status as the preferred theory in both mainstream textbooks and pedagogy. In view of this ongoing controversy, in my opinion alternative theories, which can also account for the empirical evidence cited in support of special

relativity, at least deserve to be explored and I explore one such alternative theory in this book.

Obviously, the previously-mentioned modeling process cannot be done a priori (independently of experience), but nevertheless experience can at least be used as a constraint on hypotheses. Thus, the methodology used in the book is frankly speculative, although checked by experimental results. In effect then, I am using a version of the hypothetico-deductive method. Unlike traditional usages of this method though, where the method is used in the context of testing purely mathematical models of the physical world, I am using it in the context of testing ontological hypotheses about the nature of that world. Still, as with traditional usages of the hypothetico-deductive method in traditional science, with my usage of the method, hypotheses are tested against experience and certainty is never gained in the process. Clearly, the fallacy of affirming the consequent is committed if it is claimed that the hypotheses are shown to be true because their empirical consequences are verified. This is because it is always at least logically possible that another hypothesis would have the same empirical consequences. Still, it can be held that the hypothesis is in some sense, which Karl Popper (1934/1959, sec. 83, app.*ix) at least attempts to clarify in his *Logik der Forschung (The Logic of Scientific Discovery)*, “corroborated” when it is consistent with observations.

Another methodological point concerns the issue of to what extent there is a carry-over between standard laws of physics (e. g., various conservation principles) and the internal model. One point to make in this regard is that the carry-over need not be complete. For example, there would be obvious regress considerations if it were postulated that atoms are comprised of atoms. Similarly, regress issues clearly arise if force fields are explicated in terms of structures which themselves contain force fields. My approach with respect to these methodological issues is frankly pragmatic; I include whatever carry-over concerning the standard laws of physics which is required in order to make the model work but not beyond this.

Still another issue which concerns both methodology and theory corroboration involves how "naturally" various parts of the theory fit in which each other; i. e., to what extent various portions are or are not *ad hoc* with respect to each other. Lack of *ad hocness* of one aspect of the model with respect to another aspect can be fleshed out in terms of either one aspect logically entailing the other aspect or at raising the epistemic probability of its being the case. I will leave it to readers to judge the naturalness of the internal fit of the model as just elucidated.

In view of the foregoing methodological points, I certainly make no claims for the truth of the model or for the uniqueness of the model with respect to known empirical data. Still, the aim is truth. The concept of 'truth' here can be fleshed out in accordance with the correspondence theory of truth in the sense of accurately depicting the manner in which the physical world exists. Also, I at least attempt to make what I say be consistent with known empirical facts. Still, I am somewhat more confident with some aspects of the model, such as those of the electric field and of light, than with other aspects such as gravity and atomic physics. In fact, I may well have erred on the side of including too much highly-speculative material, although I have not included anything which I know is false. The particle physicist Sabine Hossenfelder (2018) points out the non-efficacy of the standard model of physics in making new predictions based on such criteria as the beauty or symmetry of the mathematics involved, and thus there is something to be said for a new approach such as the one of this book.

I agree with Popper (1982) and David Bohm (Bohm and Hiley, 1993) in defending physical realism and in decrying such intrusions of subjectivity into quantum physics as epistemic interpretations of Werner Heisenberg's uncertainty relations, appeals to "distinguishability" and "indistinguishability" in accounting for interference phenomena, claims that physical quantities are only meaningful when they involve "observables," and in claims that either the measurement process per se (such as when instruments are involved) or even the registration of this process in consciousness reduces wave packets. I also decry Niels Bohr's embrace

of contradictions with talk of dualities (such as the wave particle duality) and complementary properties as being basic. Noteworthy, Popper (1982, p. 126) also rejects claims that these concepts are basic. A necessary, although not sufficient, mark of truth is internal consistency, and thus a minimal condition for a physical model, even on the quantum level, is logical consistency.

In spite of the agreements with Popper which I just cited, I disagree with Popper with respect to his claim that quantum uncertainty is the result of scatter relations. I also disagree with Bohm's appeal to physical pilot waves for guiding the trajectories of particles, along with such other attempts at objective interpretations as the stochastic approach of Edward Nelson (1985) and the many worlds approach of Hugh Everett (1957). Instead I appeal to a physically realist interpretation of Richard Feynman's (Feynman and Hibbs, 1965) path integral approach whereby it is claimed that physical particles take all physically-possible paths between sources and absorbers.

Like René Descartes (1644/1983, Part 2, par. 28) in *Principia Philosophiae* (*Principles of Philosophy*), I do not believe in the existence of physical action at a distance, and thus hold that all physical causation involves contact forces. In particular, in the book I outline an account of the physical world which is based on the existence of fields filling space. Thus, I make a distinction between a space and what "fills" it. I in turn flesh out the subject matter of what "fills" the space of the field in terms of properties of "ideal liquids" filling three-dimensional subspaces. These three-dimensional subspaces in turn are construed as being located parallel to each other in a four-dimensional overall space. The postulation of these subspaces is also a necessary component of both my explanations of superpositions (where particles exist in more than one state simultaneously) and of interference effects of light.

It should be emphasized that, in effect, the existence of something filling space, constitutes a special reference frame for the space. In fact, as I noted at the beginning of the introduction, my approach is based throughout on special

reference frames and thus I do not hold that the basic laws of physics are invariant among different reference frames. For example, I appeal to privileged reference frames both for determining the speed of light (which I take to be relative to that of the source particle) and for reductions for wave packets in quantum theory, which I hold involve the absorption process and thus also determine a special reference frame. In fact, I utilize special reference frames when I give an account of polarization entanglement where I account for correlations of angles of polarization at a distance in terms of properties of the electromagnetic field and the manner in which "photons" are absorbed from it. While this account utilizes special reference frames, it does not also appeal to the concept of action at a distance.

As noted, since the model postulates special reference frames it is not based on the claim of special relativity that all reference frames are equivalent with respect to their basic physics. For a discussion of the role of special reference frames, even in the context of special relativity see Kaufman and French, (2025). Still, as has been pointed out by numerous people (see Philippe Eberhard, 1978), it may still be possible to hold a purely epistemological version of special relativity since quantum correlations cannot be used for purposes of sending a signal. I do not try to establish this in the book though. While on the subject of relativity it should be mentioned that science has never been able to measure the velocity of an individual photon since it is always the two-way (and thus average) and not one-way velocity that is measured. It should also be emphasized that in the measuring process light is always interfered with, such as with the mirrors of the Michelson-Morley interferometer, and that this always creates a new source.

The foregoing points concerning the role of special reference frames will be elaborated in my discussions of parallel subspaces in Chapter One, my discussion of electromagnetism in Chapter Two and my discussion of light in Chapter Three. Chapter Four is devoted to a speculative, but realist, discussion of the internal structure of the atom loosely based upon the Bohr model of the atom.

Finally, in Chapter Five I present a sketchy and speculative account of gravity in terms of its being a residual effect of electromagnetism.

Chapter One

The Concept of a Field

I take the existence of fields literally, as opposed to the manner in which positivist positions take them, just as hypothetical entities postulated to help calculate observables such as the accelerations of particles. Also, I distinguish between a space and what "fills" it. Thus, I reject purely geometric characterizations of fields, such as those of Einstein and Hermann Minkowski with their conception of space-time. That is, I hold that space serves as a "receptacle" whereby positions in it are occupied by some sort of entity - or entities. In other words, to use Willard Quine's (1953) language, I make an ontological commitment to the existence of something "filling" spaces and possessing an independent existence apart from those spaces per se; i. e., occupying the locations of the space.

A few points of comparison can be made between my position and the closely-related position of Descartes. For one point, like Descartes (1644/1983, Part 2, par. 11), I hold that even a physical vacuum may be filled by some sort of a substance, whose character I go on to specify in this chapter. Also, like Descartes (1644/1983, Part 2, par. 33), I postulate that what fills physical space is a series of vortices (which I explicate in terms of the motions of "thin shells") capable of circular motion, although I hold that other types of circulatory motion may also be possible here as well. Unlike Descartes though, I do not take the absence of a vacuum to be a matter of conceptual necessity, and thus am also willing to posit the existence of a vacuum (which I term an "empty vacuum") not filled by any substances. In fact I appeal to such a concept in my account of thin shell formation in Section 1.4. Also, Descartes did not claim that ideal liquids fill space.

I begin my discussion by identifying a series of closely-related basic issues concerning the nature of physical fields. One issue involves a distinction that is often drawn between a force field and the energy density of a field. Other issues which are discussed include those of when to sum the effects of a field and the rates at which the intensities of fields decrease as a function of their distance from a source particle. I then cover the question of what fills the space of a field where my answer involves postulating ideal liquids to play this role. This is followed by an account of the creation of superpositions of states by postulating a system of thin shells located in parallel subspaces. I close the chapter by developing a concept of wave-particle unity for resolving the alleged duality of wave and particle properties of matter.

1.1 Force Fields vs. Energy Density Fields

I begin this section by critically discussing the traditional concepts of force fields and energy density fields. After analyzing the traditional concepts of both, I show how to reconstruct the traditional concepts so as to unite them. I hold that there is just one field here that possesses both vector (associated with forces) and scalar (associated with energy density) properties. Thus, I suggest means to integrate the two concepts. I begin by briefly elaborating respectively on the concepts of a central force field and of an energy density field. This is followed by a discussion of at what rate the magnitude of the fields diminish as a function of their distance from a source particle. I then discuss the issue of when to sum the effects of a force.

Central force fields are vector fields in the sense that each point of the space comprising the field has a vector associated with it. The forces of which these fields are comprised are also sometimes called "centripetal forces" from the Latin for "seeking the center." In the space of the fields each point of the space has a vector associated with it which is in the direction of the sources of the fields. Good examples of such central force fields include the electric field and the gravitational

field where each vector gives the force exerted by the field on respectively a unit charge or a unit mass for that particular location.

Energy density fields are scalar fields since energy is a scalar. In the case of electromagnetism they are given by the process of squaring the $\mathbf{E}$ and $\mathbf{B}$ fields with the dot product, which creates the scalar energy density of the fields $\frac{1}{4}\left(\varepsilon\mathbf{E} \cdot \mathbf{E} + \frac{1}{\mu}\mathbf{B} \cdot \mathbf{B}\right)$ where ε is the electric permittivity and μ is the magnetic permeability. In the case of electromagnetic radiation the direction of energy flow per unit area is given by area $\mathbf{S} = \frac{1}{\mu_o}\mathbf{E} \times \mathbf{B}$ where $\mathbf{S}$ is a vector in the direction of propagation perpendicular to the electric and magnetic fields, and μ_o is the magnetic permeability constant of free space. The sense of energy being used here is that of potential energy, a concept which can be further explicated in terms of the potential to raise electron energy levels during the absorption process.

I now turn to the issue of how to reconstruct the concepts of central force fields and energy density fields so as to have a unified concept. It should be emphasized right off the bat that there are obvious tensions between the two concepts since one – the central force field – is a vector field, while the other – the energy density field – is a scalar field. Also there are issues concerning the rate at which the intensities of the fields diminish as a function of distance from a source particle – e. g. whether this is at an inverse linear rate with respect to this distance or an inverse square rate. It turns out that the two subjects are connected and thus I treat them together.

Regarding the issue of whether a field is a vector or a scalar, one key issue is whether the effects are isotropic (the same in all directions), or instead, as with dipole models are anisotropic (varying as a function of direction). Presumably, as I wish to re-emphasize, there is just one field to cover both force and energy density. At least it is simpler to conjecture this. Thus, in order to be consistent, the concept of a field needs to be reconstructed so as to both cover just a single set of directions and to possess the vector property of a force along with the scalar property of

energy. There is a conflict here though with James Maxwell's theory inasmuch as the $\mathbf{E}$ and $\mathbf{B}$ fields are construed as force fields in Maxwell's equations. As I previously noted, the scalar energy density of electromagnetic fields is proportional to $\mathbf{E} \cdot \mathbf{E} + \mathbf{B} \cdot \mathbf{B}$. Also, for the dipole model of radiation, given for example in John David Jackson's (1962/1978, Sec. 9.2) *Classical Electrodynamics*, the magnetic field $\mathbf{B}$ decreases at an $1/r$ rate by $\frac{\mathbf{B}}{\mu_0} = (\mathbf{n} \times \mathbf{p}) \frac{e^{ir}}{r}$, where $\mathbf{p}$ is the electric dipole moment and $\mathbf{n}$ is a unit vector. It can be pointed out that when $\mathbf{B}$ is squared, the resulting energy density decreases at a $1/r^2$ rate. In the reconstructed concept of a field, the field both possesses the vector properties of a central force field and the scalar properties of energy density determined by the magnitude of the force field.

I now turn to the issue of when to sum the effects of a force field. Traditionally this is done initially before summing the effects of the forces themselves, whereby there is a single common field where charge effects are added for each charged particle so as to create a single resultant force $\mathbf{F}_r$ whose strength is given by the overall strength of the field ($\mathbf{E}_r$ for the electric field) at a given location; i. e., $\mathbf{F}_r \propto \mathbf{E}_r$ where

$$\mathbf{E}_r = \sum_{i=1}^n \mathbf{E}_i \quad (1-1)$$

However, in my account this is done after summing the forces whereby the resultant force $\mathbf{F}_r$ is given by the vector summation (superposition) of n distinct force fields $\mathbf{F}_i$ at a given location; i. e.,

$$\mathbf{F}_r = \sum_{i=1}^n \mathbf{F}_i \quad \text{where } \mathbf{F}_i \propto \mathbf{E}_i \quad (1-2)$$

and where each of the n charged particles (approximately 10^{80} for each charged particle in the universe) possesses its own distinct field in a separate parallel subspace.

A point of terminology should be made now. In the traditional literature a "field" is usually used to refer to a total field; e. g., an electric field or a gravitational field. This can cause confusion, in such contexts as those of where the effects of

the components fields cancel out. There is zero net force here even though the component forces, whose effects cancel out, still exist. Usage will usually make it be clear from context, but where it is not, I will specify either the “component field” or the “total field.” It can be noted that this also ties in well with the claim made in quantum electrodynamics (QED) that each electron possesses its own electromagnetic field, sometimes called a “photon cloud” see Franz Mandl and Graham Shaw (1993, pp. 102, 117). A notational point should also be made. In an attempt to minimize confusion I will use **E** and **B** to refer to the respective electric and magnetic fields as traditionally conceived, but in instances where my usage diverges from this, and this divergence may not be clear from context, I will so specify.

From the foregoing discussion it can be seen that I make a sharp distinction between a charged particle and its field. In fact, I distinguish between the two topologically, holding that at least bound charged particles (I do not deal with free charged particles in the book) possess four spatial dimensions while their fields are spatially three-dimensional. Also, while postulating such a large number of subspaces may appear to offend against principles of parsimony, such as Ockham’s razor, I believe that it is necessary in order to adequately account for interference effects. I might remark that this postulation of parallel subspaces was anticipated by David Deutsch (1997) with his variant of Everett’s many worlds interpretation of quantum mechanics. Deutsch however does not discuss the issue of how there can be an interaction among these different subspaces, which is required to account for interference effects. Also, as Timothy Maudlin (2002, p. 5) among others points out, it is not at all clear how to generate quantitative probabilities from the multiple worlds. As I will develop in detail in Chapters Two and Three, my solution to both of these issues is to claim that the whole series of parallel subspaces is “cut” by a four-dimensional bound particle, which thus can be influenced by each of them.

Thus, under my system all of the charged particles in the universe are causally interconnected – in fact in two senses. Both of these senses make the

assumption, which I took note of in the Introduction, that all causal interaction is local; i. e., that there is no action at a distance. One sense is the claim that, since each particle has a set of fields associated with it, each of these fields will eventually reach (and causally interact with) each of the other charged particles. The second sense is the claim that a superposition of the fields associated with each charged particle overlaps (and hence results in a causal interaction with the relevant charged particle) at the location of each charged particle.

I will now address the issue of what a field consists of; i. e., what “fills” a three-dimensional subspace by analyzing the subject of an “ideal liquid.” In the following two sections I develop the concept of parallel subspaces in detail and introduce the concept of “thin shells.” Finally I make some remarks on the subject of wave particle unity.

1.2 Ideal Liquids

To begin my discussion of what fills a three-dimensional subspace, I wish to point out that obvious difficulties are involved if it were to be postulated that the ultimate constituents involved anything like the modern conceptions of solids, liquids, or gases. This is because each of these is postulated as having components – atoms – which themselves are not solid, liquid, or gaseous. There are also problems with postulating anything like Michael Faraday’s lines of force as ultimate constituents both since it is not clear how anything can physically exist if it only possesses one spatial direction and due to issues concerning how, if the lines literally exist they will be constantly getting tangled up with each other.

In spite of the foregoing points, it is possible to cite older PreSocratic concepts such as the concept of an “ideal solid,” as being possible candidates for being the ultimate constituents of space. The concept of an “ideal solid” was perhaps anticipated by Parmenides with his concept of a spherical “plenum” filling space. However, the concept was only appreciably developed by the ancient Greek atomists Leucippus and Democritus. It is described in considerable detail by the Roman follower of Epicurus, Lucretius (c. 60B.C.E./1951), in *De Rerum Natura*.

An ideal solid can be defined as being an impenetrable substance moving in an empty vacuum, with each of its parts also being impenetrable and non-separable from other parts. Thus, classical atoms were thought of as being indivisible. Ideal solids are also conceived of as being perfectly rigid; i. e. their shape is conceived of as remaining constant when put under an indefinitely great pressure.

In contrast to the concept of an ideal solid which has received quite a bit of discussion there has been relatively little discussion of the ideal liquid state in traditional philosophy, although the concept is used in modern materials science (see G. K. Batchelor, 1967, sec. 1.8). For example, John Locke (1690/1959, Bk. 2, Ch. 8) in his *An Essay Concerning Human Understanding* includes solidity but not liquidity in his list of the primary qualities. However, for a variety of reasons, which will become apparent shortly, I prefer the concept of an “ideal liquid” over an “ideal solid” as being the ultimate constituent of space. Thus, I now turn to an elaboration both of what that concept has in common with and of how it differs from the concept of an ideal solid.

I postulate ideal liquids as being like ideal solids (and unlike ideal gases) in possessing constant volumes and hence being incompressible and also of constant density (although not in a sense of being comprised of a more primitive internal structure). It can be noted that a conservation principle of total matter follows from this property, although it is consonant with that matter being rearranged in various manners. I also conceive of ideal liquids as being like ideal solids in being impenetrable from each interior part unless these parts are “pushed aside.”

I also conceive of ideal liquids as being unlike ideal solids in a number of respects. In particular, unlike ideal solids, I conceive of ideal liquids as changing shape under pressure and also as ceasing to cohere together when pulled from different directions. In the latter case I conceive of them as dividing into parts. Conversely, I also conceive of them as being capable of recombining either from the same parts or parts from other ideal liquids. This last topic is connected with the property of “adhesion” whereby I hold that ideal liquids are capable of attaching

to other ideal liquids when they come into contact with them. A good example of this is the liquid drop, whereby numerically different drops may split apart and then recombine in various ways so as to form at least qualitatively identical drops. Looking ahead, this will turn out to be an important property in my explanations of both attraction and repulsion. When not subject to contrary forces, like ideal solids, I conceive of ideal liquids as cohering together, even when accelerated.

The topic of the viscosity (resistance to flow under an applied force) of ideal liquids should also be addressed. As I conceive of them ideal liquids possess no resistance to shear (forces coplanar with their cross sections), and in the language of materials science their shear modulus (measure of rigidity) is zero. In other words, the coefficient of viscosity of ideal liquids is zero; i. e., the liquids are conceived of as being non-frictional. I also conceive of ideal liquids as retaining their shape when not subject to external forces. It can be pointed out that the fact that physical liquids like water lack this property of possessing a fixed shape is not a counterexample to my claim that in the absence of forces ideal liquids possess a fixed shape in spite of possessing zero viscosity. This is because in fact there is a force present in the water example, namely gravity. As I will explain subsequently in my discussion of thin shells in Section 1.4, I do not conceive of forces such as gravity working within these shells themselves. Also, looking ahead, I might note that I appeal to both the property of zero viscosity and of the retention of shape in the absence of countervailing forces with my explanation of the Renninger effect in Chapter Three.

The foregoing can also be expressed quantitatively in terms of the equations of fluid mechanics. For example, the conservation of total mass for an ideal liquid can be expressed by the simple continuity equation

$$\frac{\partial \rho}{\partial t} = -\nabla \cdot (\rho \mathbf{u}) \quad (1-1)$$

where ρ is the fluid density and $\mathbf{u}$ is the fluid velocity. The incompressibility of an ideal liquid can be expressed by the incompressible Navier-Stokes equation which,

ignoring the gravitational field and given my assumption of constant fluid density, is given by

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} - \nu \nabla^2 \mathbf{u} = -\nabla \frac{p}{\rho} \quad (1-2)$$

where ν is the fluid viscosity and p is the fluid pressure. When there is zero vorticity (in my system this corresponds to the absence of a magnetic field), then equation 1-2 implies to a version of the Euler equation

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = -\frac{1}{\rho} \nabla p \quad (1-3)$$

As I conceive of them, Ideal liquids come in two types – positive and negative. Purely as a matter of stipulation I hold that these two types of ideal liquids correspond respectively to positive and negative electrical charges. This constitutes the connection (or so-called "bridge principle") between the basic entities postulated by the model – ideal liquids – and observables inasmuch as the postulation of physical charges has observable consequences. I conceive of opposite ideal liquids as “flowing into” each other when they are spatially indefinitely close to each other. The "force" responsible for such a "flow" is clearly not the Coulomb force inasmuch as that force would tend to infinity along an interface, and the force causing the "flow" here is a contact force. This subject obviously is in need of further elaboration, but I believe that its intuitive sense should be clear.

While the ideal liquids are not observable per se, their postulated identity with positive and negative electric charges constitutes a bridge principle linking these ideal liquids with the empirical content of electromagnetism. In particular, I show that such properties of the ideal liquids as their speed, volume, and topology have corresponding counterparts in electromagnetism. I also show that the postulated oscillation of the ideal liquids accounts for electromagnetic forces. Since the concept of a parallel space is a key component to my subsequent discussion of how fields comprised of positive and negative liquids can oscillate between different sets of dimensions, I now turn to making a few remarks on that topic.

1.3 Parallel Subspaces

I hold that the ideal liquids introduced in Section 1.2 exist in a series of three-dimensional parallel subspaces where a separate subspace is associated with each charged particle in the universe. One possible model here holds that these subspaces are spaced at discrete intervals in an infinite four-dimensional overall Euclidean space (I might note that I will be subsequently postulating a five-dimensional overall space) in which they are embedded. Since there are approximately 10^{80} charged particles this results in approximately 10^{80} subspaces – admittedly a large, although still finite number. A merit of this model (assuming that there is no “action at a distance”) is that holding that the fixed subspaces are located at fixed intervals in the fourth dimension may serve to “insulate” the ideal liquids in the subspaces from causal connections in these different parallel subspaces, except in locations where the subspaces are mutually intersected by a particle possessing at least one higher spatial dimension. A demerit of this account is that the interval spacing in the account appears to be rather *ad hoc* if it cannot be independently motivated. Thus, I will also investigate the merits of parallel subspace models where there is no spatial gap between the subspaces, even though, as I will show, these have problems of their own, including the one of opening up the possibility of causal connections between sequential subspaces.

There are well-known paradoxes dating at least back from the time of Zeno of Elea (such as how to generate magnitude out of something without magnitude) about points in a space being adjacent to each other. The key question is how the points (or subspaces) can be both disjoint but also adjoining. I believe that the answer to these paradoxes involves using topologic concepts instead of metric ones. I now turn to a discussion of the series of topologic concepts of continuity, connectedness and adjacency in this regard. It can be noted that these are necessary concepts to utilize in analyzing the boundaries between subspaces and particles they

causally interact with along with the current topic. The following analyses of possible conditions for the topological conditions of separating spaces (which possess one fewer dimension than the dimensionality of an overall space and are sometimes called “partitions”) and adjacency for parallel subspaces are due to Ernie Kent. I might note that Kent’s analysis of adjacency can also be applied to the intersection between a four-dimensional particle and a three-dimensional subspace cutting it. Both analyses assume a finite fourth-dimensional extent to the subspaces. It should also be emphasized that the two analyses are incompatible with each other inasmuch as the subspaces do not contain limit points (closures) at their boundaries in the first analysis while at least one subspace does in the latter analysis. Using the concept of “connectedness,” it follows that the two sequential subspaces are not “connected” in the first analysis but are “connected” in the second analysis. Since both analyses utilize the ordered geometry concept of “betweenness” I begin with a discussion of that concept. I might note that some of the same points can also be made utilizing the concept of a “sequence.”

I begin my discussion of “betweenness” by dealing with the special case of the relationship between points (zero-dimensional spaces) and a straight line (a one-dimensional space) but this can be generalized to the relationships between n dimensional subspaces and overall spaces of dimension $n+1$. The concept of “betweenness” can be fleshed out now as a principle on a straight line whereby for any three points A , B , and C on the straight line, B lies in between A and C when it is neither the case that A lies in between C and B nor does C lie in between A and B . As I noted, there are well-known issues concerning continuity to simply define adjacent points as being two points where there is no point in between them. However, more can be done by using the concept of a limit point as I will now show. I now turn to Kent’s analysis of the conditions for separating spaces in between the subspaces..

Let each subspace be an open set of the four-dimensional space with a finite thickness in the fourth dimension, and let them alternate along the fourth dimension

with closed three-dimensional sets having a thickness of a single point in the 4th dimension (i. e., a three-dimensional “surface”). Projecting onto the real line lying in the fourth dimension, one would have $(a...b)[b](b...c)[c](c...d)[d](d...e)...$ where the letters are points of the single-point thick closed sets (consisting of infinite three-dimensional surfaces in a four-dimensional space. Thus, each point of a closed three-dimensional “sheet” would be a limit point of the open four-volume subspace on either side and would be a boundary between the two. The boundaries serve as “separators” to partition the subspaces from each other and thus to “insulate” them from causal connections. If this “separation” is not sufficient from allowing the subspaces to be sufficiently “squashed together” so as to allow causal connections among them, it might even be possible to postulate “insulating subspaces” in between each of the other postulated ones so as to avoid these causal connections (this also applies to the subsequent analysis of “adjacency”), although admittedly this is pretty *ad hoc* in both cases. In any event, the resulting space is connected since there is no union of open sets equal to the entire space. It also is continuous since it is possible to approach any point in a closed set as the limit of a function from the open set on either side, thus meeting the definition of continuity at a point. I now turn to Kent’s analysis of the adjacency conditions for the subspaces.

Consider n subspaces P where n is a large number, approximately 10^{80} (which I pointed out earlier is roughly the number of charged particles in the universe). The intersection between any pair of subspaces is null (the empty set). Now, pick a particular subspace M . M will contain as limit points the points of two of the subspaces P_l and P_n . Each of the P_n of which M contains limit points of M , do not contain limit points of each other. Thus, subspace M is adjacent to both subspace P_l and P_n in the sense that there is nothing “in between” it and either of these subspaces. It can also be noted that while M will thus be adjacent to both P_l and P_n , P_l will not be adjacent to P_n . When points (or subspaces) are adjacent in the just-elucidated sense, I will speak of these points (or subspaces) as being

“indefinitely close” to each other. It is possible now to iterate the foregoing procedures over the entire collection of P with each subspace in turn being M .

It might also be possible to just question Richard Dedekind’s (1901/1963, p. 13) postulate about the existence of both least upper and greatest lower bounds for all numbers, at least in this context of dealing with spatial points. In the numeric context the idea then would be to settle for dealing with just the rational numbers instead of all of the real numbers and analogously to settle for just the concept of density (whereby, at least roughly, there is a point in between any two given points) in dealing with the issue of the continuity of sets. This completes my discussion of the continuity and adjacency conditions for the subspaces. Thus, I will now enlarge on previous brief remarks on the topic of nature of dimensionality which is also a key concept for my subsequent discussion of parallel subspaces.

The dimension of a space can be given by a recursive definition due to Menger (1943). Menger builds on an analysis originally due to Henri Poincaré (1912/1963) whereby, "dimension" is defined recursively in terms of the minimum number of dimensions required of a space in order for that space to "cut" or give boundaries to the space whose dimensionality is being tested. The bounding space will then possess one fewer dimension than that of the space being bounded. For example, in the hypothetical case of a one-dimensional space, such as a circle, the space can bound a two-dimensional space. Similarly in the hypothetical case of a two-dimensional space, such as a spherical surface, the space can bound a three-dimensional space. Menger adds the requirement that the space whose dimensionality is being tested must be capable of being given boundaries in each of its infinitely small neighborhoods, by a space of one fewer dimensions. He then defines the dimensionality of a space as being one greater than that of a bounding space for each of the infinitely small neighborhoods of the original space. This addition is to avoid such counterexamples to Poincaré's analysis as of two cones meeting at a point, which could be bounded at the point of intersection by a space of zero dimensions, a point, and of a solid ball embedded in a solid torus which

could be bounded by a space of one dimension, a circle. However, since the subject of physical subspaces would appear to have little in common with examples such as these, I will just use Poincaré's analysis.

I wish to now introduce a fourth spatial dimension to the system being postulated, with the resulting four-dimensional system containing a series of three-dimensional subspaces in which the positive and negative liquids interact. I first show that it is possible for two three-dimensional subspaces to be located as to be spatially indefinitely close to each other in this four-dimensional "over-all" space, in the sense that for each point in one three-dimensional subspace, there exists a point in the other three-dimensional subspace which is indefinitely close to it. While there are obvious issues concerning continuity with this claim, I will not address them, but instead just present an intuitive inductive argument to make the claim plausible. To make this inductive argument, it can first be noted that two lines can lie indefinitely close to each other in a plane, in the sense that for each point on one line, there will exist a point on the other line which is indefinitely close to it. Similarly, it can be pointed out that it is possible for two planes to lie "flat" against each other in a three-dimensional space, where for each point in one plane, there will correspond another point in the other plane which is indefinitely close to it. Thus, by induction it follows that in the case of two three-dimensional spaces lying "flat" next to each other in a four-dimensional space, for a given point in one three-dimensional subspace, there will be another point in the other three-dimensional subspace which is indefinitely close to it.

It can be noted that the inductive argument for the adjacency of points in adjoining subspaces can be repeated an indefinite number of times. However, since in my model each subspace is associated with a separate charged particle, I only repeat the argument 10^{80} times - the approximate number of charged particles in the known universe. While this is a large number, it is still finite, and thus is insufficient for attaining the existence of a four-dimensional embedding space. One must be

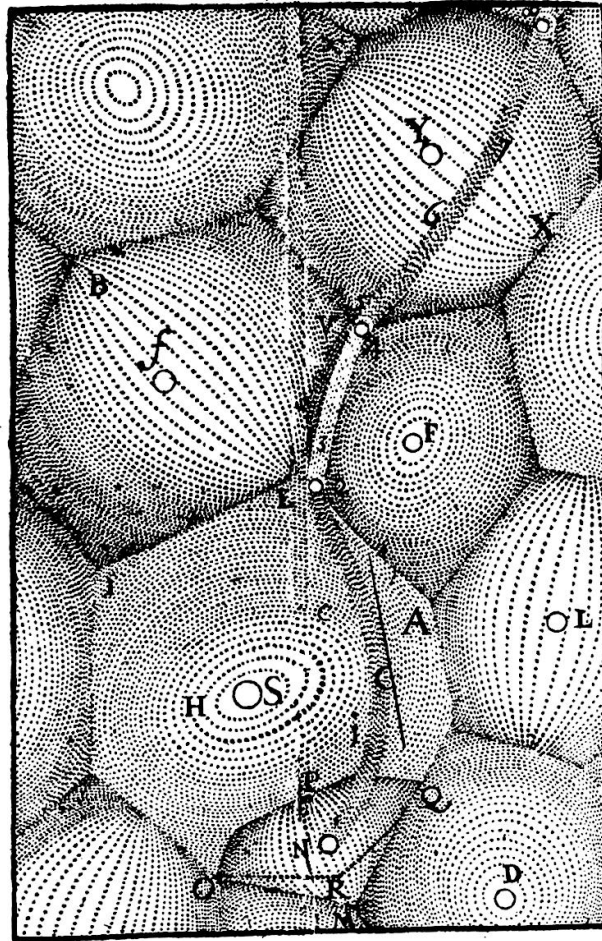


Figure 1.1 Descartes's system of vortices in the *Principles of Philosophy*

careful just appealing to physical intuitions on these matters though since evidently in many ways the physical world is much stranger than might be thought *a priori*.

This completes my treatment of the topic of the nature of parallel subspaces. Thus, I now turn to my application of that discussion in an explanation of how opposite ideal liquids can oscillate in a set of dimensions orthogonal to a series of parallel nested "thin shells."

1.4 Thin Shells

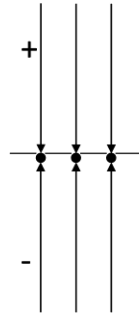
As I remarked in the Introduction, to some extent what I say concerning the subject of thin shells has been anticipated by Descartes (1644/1983, Parts 2 and 3). Descartes's system is illustrated in Figure 1.1. In particular, I agree with Descartes

that the shells are capable of circular motions (Descartes 1644/1983 Part 2, par. 33) and in fact I appeal to such rotations in my accounts of magnetism in Chapter Two and of light in Chapter Three. As previously noted, unlike Descartes I hold that other forms of closed circulatory motions are also possible. However, I do not utilize other forms of circulatory motion in the book.

Also, unlike Descartes, I hold that each thin shell is a complete sphere and is located in its own separate subspace. Otherwise, as Isaac Newton (1687/1934, Bk. III, General Scholium) argues with respect to Descartes's theory of vortices, there is a problem in accounting for the highly eccentric orbits of comets; there being an implicit assumption that such orbits would be blocked by the perfect circular character of the vortices. Regardless of the merits of the issue of whether such an objection refutes Descartes's theory, I wish to emphasize that it does not affect my own account of inverse square rates since I do so in terms of cumulative effects of series of thin shells. I should also point out that if a principle (which I hold in my subsequent treatment) that each thin shell possesses an equal volume of ideal liquid is applied to Descartes's system thin shells either enclosing or partially enclosing large numbers of particles would become enormous. It is not at all clear how this can be made to fit in with the central force inverse square character of the laws of electromagnetism and gravity.

In my account by "thin shells" I refer to an indefinite number of nested equi-volume three-dimensional shells surrounding a charged particle. The shells are "thin" in the sense that the ratio of their respective widths to their respective circumferences becomes indefinitely small as their respective radii become indefinitely great. In my discussion of these shells I make use of the same analysis of "dimension" which I gave in Section 1.3. In particular I make use of the "in the large" analysis of dimension which Poincaré gave whereby the dimensionality of a space is one greater than that of a "cutting space." I also make use of this concept of a "cutting space" in my discussion of the interface along which ideal liquids flow into each other. Looking ahead, I might note that the resulting oscillations of ideal

Parallel lines intersecting points on an orthogonal line



Parallel planes intersecting lines on an orthogonal plane

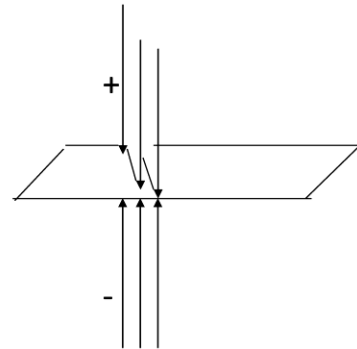


Figure 1.2 Flow of opposite ideal liquids into orthogonal sets of dimensions.

liquids will be relevant for my treatments of the forces associated with electric and magnetic fields in Chapter Two, and with the gravitational field in Chapter Four.

The hypothesis which I wish to make now is that when two opposite three-dimensional ideal liquids “flow” into each other, they are not both destroyed in the process. Instead, I wish to invoke a principle that is analogous to that of the conservation of matter, and hypothesize that the opposite three-dimensional ideal liquids “pull” each other into another indefinitely close three-dimensional subspace which is orthogonal to a series of parallel three-dimensional subspaces contained in a four-dimensional over-all space. It should be emphasized that the existence of this orthogonal set of subspaces is just a matter of postulation in order to make the model work and does not have an independent motivation. I also postulate that the “flowing together” of the ideal liquids in each subspace will only occur along the “interface” between the two opposite three-dimensional ideal liquids. That the dimensionality of the interface is two can be shown by Poincaré’s in the large criterion that the dimensionality of a bounding space be one fewer than the dimensionality of the space being bounded. Since I hold that the space into which three-dimensional liquids are being pulled is orthogonal to the first set of parallel subspaces I claim that there is no causal interaction with the four-dimensional

particles. The three-dimensional subspaces which the opposite ideal liquids “pull” each other into are defined to be exactly analogous to the subspaces from which they came, with the one exception that they are located in dimensions orthogonal to them. I postulate that these orthogonal subspaces do not intersect each other and thus possess no points in common. This is certainly possible since, as illustrated for the one- and two- dimensional cases in Figure 1.2, a series of non-intersecting parallel lines will intersect each of a series of points on an orthogonal line embedded in an overall three-dimensional space, and similarly parallel planes will intersect each of a series of parallel lines on an orthogonal plane embedded in an overall four-dimensional space. By induction, in the case of three dimensions each of a series of three-dimensional hyperplanes will intersect a series of three-dimensional regions embedded in an orthogonal three-dimensional hyperplane embedded in an overall five-dimensional space. Presumably the orthogonal subspaces associated with different particles are also parallel to each other, but nothing crucial in my analysis depends on this point. Also, locations where these lines or planes intersect orthogonal lines or planes constitute interfaces with these orthogonal lines or planes thus allowing for the possibility of causal interactions at those locations among whatever entities (e. g., ideal liquids) may be occupying those subspaces

Once the three-dimensional ideal liquids have achieved their maximum full extension in a new subspace, a restorative force constituted by the opposite ideal liquids still meeting on their same interface will make them “flow” into each other once again. Note that this is a contact force and thus does not pull from the extremities themselves as in the case of action at a distance. In any event, these forces will then pull each other back into the original subspace, where the original process will resume, and so on. Thus, pairs of positive and negative ideal liquids can be seen as continuously flowing into each other as they oscillate between sets of dimensions, presumably in something like a sinusoidal manner, back and forth from one three-dimensional subspace to the other. Also, for reasons that will

become apparent later, I postulate that the average speed at which the positive and negative ideal liquids flow into each other is the speed of light in a vacuum c .

My next series of remarks concern the geometry and topology of four-dimensional particles, which I identify with bound electrons and nucleons. With respect to the issue of dimensionality, an extension can first be noted with the three-dimensional case where, for example, a three-dimensional physical volume of air can be enclosed by the surface of a balloon, which at least approximates a two-dimensional surface. By extension from the preceding case, a three-dimensional hypersphere can be seen to surround a region of a four-dimensional space. The dimensionality of the interface between two four-dimensional particles will also be three-dimensional. I also wish to claim that four-dimensional particles are comprised of four-dimensional ideal liquids, and thus I now briefly elaborate on that concept.

By stipulation I define four-dimensional ideal liquids as being strictly analogous to the three-dimensional ideal liquids except that they occupy four spatial dimensions again using Poincaré's "in the large" analysis of "dimension." As with the case of three-dimensional ideal liquids I also claim that there are both positive and negative four-dimensional ideal liquids. These are also held to flow into each other and switch sets of dimensions precisely analogously to the three-dimensional case. It might also be pointed out that inasmuch as the three-dimensional ideal liquids switch sets of dimensions with respect to a four-dimensional overall space, the four-dimensional liquids will switch sets of dimensions with respect to presumably the same five-dimensional overall space which was previously alluded to.

Consider now the case of a four-dimensional ideal liquid sphere coming into contact with a three-dimensional subspace of tightly-packed positive and negative three-dimensional ideal liquids, all of which are indefinitely smaller than the just-postulated four-dimensional ideal liquid sphere, and which continuously flow into each other as they oscillate between sets of dimensions in the manner previously

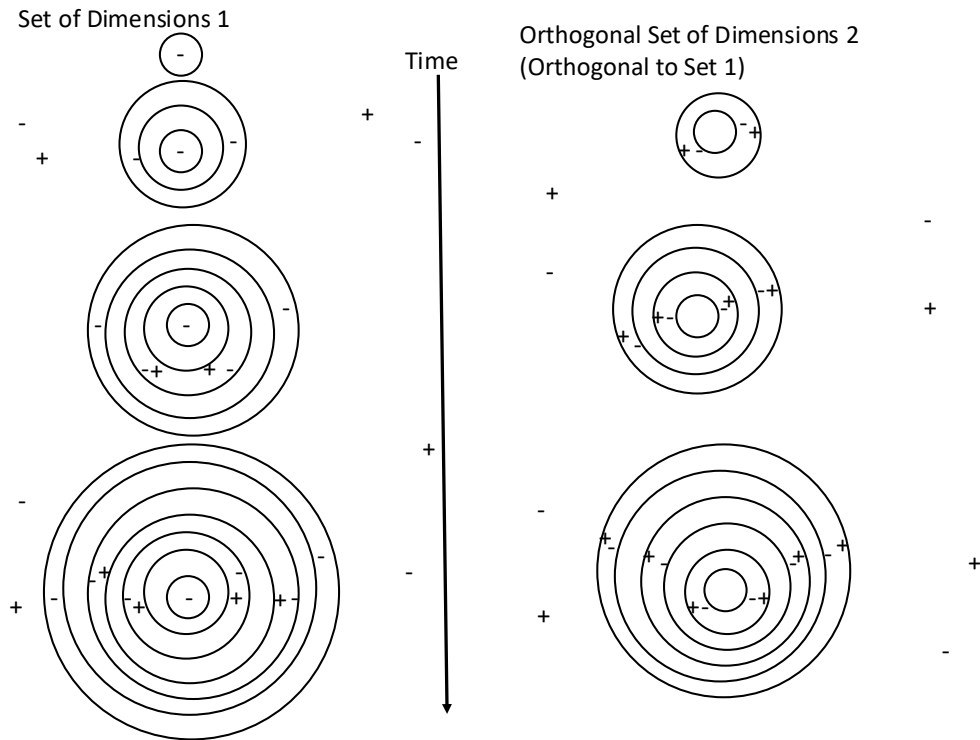


Figure 1.3 Evolution of thin shells. The +s and –s outside of the thin shell system represent indefinitely small portions of ideal liquids which are drawn into the outermost shells.

explained. Along its great circle equator, each four-dimensional sphere will be indefinitely spatially close to an extremely large number of these three-dimensional positive and negative ideal liquids in the three-dimensional subspace. Thus, those three-dimensional ideal liquids which are opposite in charge to the four-dimensional sphere, and which are indefinitely spatially close to it, can be seen to flow into it, resulting in a “thin shell” around the four-dimensional sphere. This shell will itself be only three-dimensional, since it is entirely comprised of three-dimensional liquids.

Once one of these three-dimensional thin shells around the four-dimensional sphere has been formed, there will be a superabundance (the excess

amount of ideal liquid left over after the original ideal liquids flow into each other) of the opposite type of ideal liquid from that which originally flowed into the four-dimensional ideal liquid to form the thin shell, surrounding the new thin shell. This is because a large number of the opposite ideal liquids which were originally there flowed into the four-dimensional ideal liquid in order to form the first thin shell.

This superabundance of this type of three-dimensional ideal liquid will in turn cause ideal liquids of the opposite type from it in the surrounding space to flow into it, precisely analogously to the manner in which the four-dimensional sphere caused the production of the first thin shell. A new superabundance of three-dimensional ideal liquids will then result, producing a new thin shell, and so on. Thus, an indefinitely large series of concentric thin shells will surround the four-dimensional sphere, and the ideal liquids in these thin shells will be constantly switching sets of dimensions, as their oppositely charged halves flow into each other; that is, they will oscillate between sets of dimensions. This process for generating shells is illustrated for the two-dimensional case in Figure 1.3. It can be noted in the diagram that in each thin shell the interface for the positive and negative ideal liquids flowing into each other constitutes a circle. Thus, in this case the orthogonal space being flowed into would be a cylinder. By induction for the three-dimensional thin shell case, which I identify with the electromagnetic field, the orthogonal space would thus be a hypercylinder located orthogonally to the whole series of parallel thin shells.

It can pointed out that, as illustrated in Figure 1.3, thin shells with opposite ideal liquids flowing into each other alternate with shells possessing an empty vacuum where in the case of each empty vacuum region in the orthogonal set of dimensions the corresponding shell is filled with ideal liquids. This alternation with empty shell spaces is only true for the initial process of shell formation though. Since adjacent shells will possess different widths, over time they will come to be out of phase with respect to which set of dimensions they are located in and thus will no longer always be synchronized in terms of always being located in

orthogonal sets of dimensions. Looking ahead I might note that this topic will be relevant to my treatment of adjacent electron spacing in Chapter 5.

Consider the situation where approximately 10^{80} (which I remarked earlier corresponds to the number of charged particles in the known universe) of these four-dimensional particles with their three-dimensional subspaces intertwined with each other. Each four-dimensional particle will then “cut” the whole set of three-dimensional subspaces and thus will be able to causally interact with them. I believe that this both accounts for the way in which superpositions of forces occur and for how interference effects occur since there will be contributions from each subspaces centered at different particles (which I term “originating particles” for the particular subspace in question). I might mention that Deutsch (1997) in his account does not talk about this “cutting” process, but I do not see how the contributions of the various interfering subspace fields can be “summed up” so as to create an interference effect without something like it.

Notice now that each charged particle has a unique set of nested thin shells in which it is centered. As I will explicate in Chapter Three, in the process of photon emission an emitted photon is propagated through a unique set of shells centered on the emitting particle. In contrast, I will explicate how in the process of “photon” absorption there is a superposition of contributions from the different thin shells impinging on a charged particle at a given location. Thus, under my model the causal asymmetry between the emission and absorption processes is mirrored with a corresponding structural asymmetry in terms of the thin shell structures involved.

In cases of either particle antiparticle creation or annihilation (e. g., in cases of either the creation or annihilation of electrons positron pairs) I hold that whole sets of thin shells associated with these particles are either created or destroyed in this process with this creation or destruction being spread at the speed of light.

I now investigate a key property of the set of thin shells produced in the process illustrated in Figure 1.3 – that the thickness of the shells varies as a function of the inverse square of their distances from their sources. As is well known inverse

square laws are ubiquitous in classical physics; e. g., Coulomb's laws of electricity and magnetism, Newton's law of gravity, the intensity of light from a point source, and the formula for the number of electrons in a complete row of the periodic table. To motivate the inverse square nature of the foregoing laws it can first be noted that inasmuch as the volume of a sphere is $\frac{4}{3}\pi r^3$ the area of the surface of a sphere is given by its derivative $4\pi r^2$ in the same way that the circumference of a circle $2\pi r$ is the derivative of the area of a circle πr^2 .

To make my next point it can be pointed out that the volume of each of the shells surrounding a four-dimensional sphere will be a constant. This is because the superabundance of the three-dimensional ideal liquids from which each subsequent shell is produced must equal the amount of opposite three-dimensional ideal liquids which were taken from that space in order to produce the preceding shell. Assuming that the three-dimensional space in which these rings are produced possesses a Euclidean metric structure, it follows then the cross sections of these shells will vary inversely proportionately to the square of their distances from the center of the originating four-dimensional sphere. This is because, as I elaborate on in Chapter Two, the volume of each shell corresponds to the surface area of a sphere in the middle of the shell, multiplied by the cross section of the shell. I now turn to a discussion of the role that the concepts of "waves" and "particles" play with respect to fields.

1.5 Wave Particle Unity

In this section I introduce a concept of "wave particle unity" whereby waves and particles are held to be numerically identified with each other. This involves reworking both concepts so as to be less unlike each other than with the traditional concepts. The approach is obviously very different from the approach of the traditional Copenhagen interpretation of quantum mechanics whereby, with Bohr's doctrine of "wave-particle duality" together with his "principle of complementarity," it is held that it is only possible to demonstrate one property or the other in the same experiment but not both. However, in contrast with Bohr's

approach, I make some positive suggestions for making the concepts more similar. This involves a transformation of both concepts in a position which I, following Hüseyin Yilmaz (2010), call "wave-particle unity." The referent of the transformed concept has also, often facetiously, been termed a "wavicle." Incidentally, Yilmaz has had an experiment run (Yutaka Mizobuchi and Yoshiyuki Ohtaké, 1992) where wave and particle properties are demonstrated in a single experiment. I first sketch ideal extremes of both traditional wave and particle concepts, highlighting their incompatible natures as traditionally conceived. I then show how the traditional concepts can be modified so as to make them at least less incompatible with each other.

Traditionally, waves and particles are thought of as being very different. Particles are thought of as being well-confined spatially and to travel between spatial locations in sharp trajectories. An extreme concept here is Roger Boscovich's (1758/1922, par. 88) concept of a point particle, which is held to completely lack spatial extension. This concept was also developed by Immanuel Kant (1786/2004, Ch. 2) in the context of being the center of a repulsive dynamic force field whose strength decreases at an inverse cube rate from this center. In view of the prominence of the concept of a point particle in modern quantum mechanics, a brief aside is warranted into that usage.

The modern quantum theory concept of a point particle was developed by Paul Dirac (1930/1981) with the somewhat mathematically-artificial concept of the "delta function." In order to conserve probability, the delta function is defined as integrating to a value of 1 with an infinitely high spike at the point it is centered at but being 0 at all other points. Thus, if a delta function is centered at a given point, the probability is 1 that an electron is located at that point. Along with invoking the delta function for the conception of an electron as being a point particle its invocation also entails that the resulting charge and mass density of an electron is finite. The motivation evidently was both to find a relativistically invariant concept (see Albert Messiah, 1958/1999, p. 948) and also to deal with apparent paradoxes

associated with claims that charged particles have a spatial extension. In particular, it was thought that if they were spatially extended their parts would then repel each other and hence the particle would be unstable. However, I postulate special reference frames in my model and thus the issue of relativistic invariance does not apply to it. Also, I deal with the paradox of electrical repulsion both in Chapter 2 and in Chapter 5 by denying that electric fields exist within particles.

In contrast to particles, waves are traditionally thought of as spreading out over all of the physically possible directions for them to travel. In a wave an undulation passes through a medium, but the medium itself does not travel. In contrast, a particle physically travels through space. Waves can pass through each other without disturbing each other, while particles collide when they meet, affecting the subsequent trajectories of both. Particles are thought to be indivisible and indestructible, while waves break up into wavelets when they interact with barriers.

In contrast to the foregoing traditional account of the concepts of waves and particles as being contrasting in character, with my concept of wave-particle unity, “wave-particles” are construed as being spread out over all physically possible paths. These wave-particles are conceived as possessing the wave-like property of being able to “glide past” each other in separate parallel subspaces thus avoiding the issue concerning collisions. I also view them as possessing the wave-like property of breaking up into “partial particles” when elastic scattering occurs. The situation is analogous to that of the ship of Theseus where, at least under Plutarch's account, a ship is rebuilt one plank at a time and the question is raised as to in what sense it is the same ship. Similarly, here the parts do not remain the same even though (at least during the processes of emission and absorption) the overall structure remains the same. Notice that the distinction between numerical and qualitative senses of “identity” starts to break down for these sorts of cases. Also, unlike Yilmaz's theory where a wave and a particle only go one way at a beamsplitter, under my account both the wave and the particle go each way.

The foregoing account ties in well with what I said concerning identity conditions for ideal liquids in Section 1.1. For example, as I noted in that discussion, in the case of ideal liquid drops these drops are also held to be capable of separating into distinct parts and then reuniting with parts from other drops so as to form a new at least qualitatively-identical drop. Another good way to illustrate this point involves invoking Heraclitus's famous aphorism that one cannot step into the same river twice inasmuch as the waters flowing in it keep changing.

In spite of possessing these wave-like properties I also view the wave-particle as possessing a number of particle-like properties. In particular, I hold that both absorption and emission processes involving these entities are discrete events occurring at sharp locations in space and time. However, at times I hold that these events will involve non-local causal processes as outlined, for example, in my discussion of entanglement in Chapter Three. It should be emphasized that under this account at least typically there is no numerically-identical particle travelling along a fixed trajectory between an emitter and an absorber. Instead, there is at most a qualitative identity between what is emitted and absorbed with respect to such properties as wavelength, frequency, energy and virtual mass. Also, I hold that a traveling wave-particle which travels through a medium is a coherent concept. This is discussed in my treatment of thin shell formation in Section 1.3 and is developed in detail with my treatment of the propagation of light in Chapter Three.

This completes my discussion of the particular conceptual model of a field which I will be utilizing in the remainder of this book. I now attempt to demonstrate the efficacy of that model by means of applying it to the specific case of the electromagnetic field. Looking ahead I might remark that the conception of the electromagnetic field in my model is key to both my treatment of so-called "nuclear forces" in Chapter Four (which I explicate in terms of the claim that there is no electromagnetic repulsion among nucleons) and in my analysis of the gravitational field - which I analyze in Chapter Five in terms of the claim that it is a remnant of oscillations in the thin shells associated with the electromagnetic field. Also, I do

not appeal to the existence of a separate quantum mechanical force field apart from the electromagnetic field, i. e., a "quantum potential" as is postulated for example by Bohm and Hiley (1993, p. 29), in my treatment of quantum phenomena. Thus, I end up holding that only electromagnetic fields as instantiated by properties of thin shells in accordance with my model, exist at a fundamental level.

Chapter Two

Electromagnetism

In this chapter I develop models of the electric and magnetic fields based on the system of "thin shells" which I sketched in Chapter One – along the directions connecting a thin shell with its originating particle for electric fields and with rotational effects perpendicular to these directions for magnetic fields. It can be recalled from my discussion in that chapter that I construe fields as having both vector (e. g., the directional properties of a force) and scalar (e. g., the property of energy density) properties. In particular, I hold that the central force field properties are due to the motion of ideal liquids oscillating in the direction of the originating particle. In turn, I explicate the energy density properties both in terms of the vibrational frequencies of longitudinal oscillations of ideal liquids within the thin shells in this chapter and in terms of transverse rotational oscillations of these shells with my explication of light in Chapter Three.

I hold that the electric field is the most primitive of all the force fields. Its dominance is not always apparent since, as I pointed out in Section 1.1, I use the term “electric field” to refer to the total electric field (in the sense of including all of its component parts), even though the effects of the fields of individual thin shells systems typically cancel out. As I noted at the end of Chapter One, I hold that gravity is a residual effect of not completely balancing electric fields and that so-called nuclear forces are due to the non-existence of electric fields within nuclei *per se*. Also, while magnetic fields are traditionally defined in term of relative motions among the electric fields of different charged particles, I do so in terms of vortices associated with electron spin which I take to be ontologically basic. I begin by explicating the nature of the electric field in terms of my model and then turn to my explication of the magnetic field.

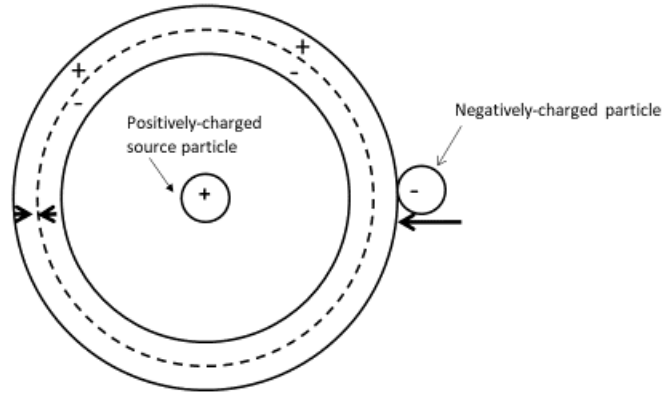


Figure 2.1 Structure of shells for electrical attraction. The arrows indicate the direction of ideal liquid flow.

2.1 The Electric Field

I begin my discussion of the electric field in my model by dealing with the electrostatic case where there is no relative motion between two charged particles. As I pointed out in Chapter One, since the area of a sphere is proportional to r^2 , if the volumes of the thin shells remain a constant their width will approach being proportional to $1/r^2$ as r approaches ∞ . In particular, if n is the ordinal position of a thin shell and if the volume of each thin shell is normalized to 1, then the total volume V_n enclosed by the n^{th} shell will be n . The radius of the outer shell R_n will then be equal to $\sqrt[3]{\frac{3}{4\pi} V_n}$ and the radius of the inner shell R_{n-1} equal to $\sqrt[3]{\frac{3}{4\pi} (V_n - 1)}$. The difference δR_n (i. e., the width of a thin shell) in radii $R_n - R_{n-1}$ is thus equal to

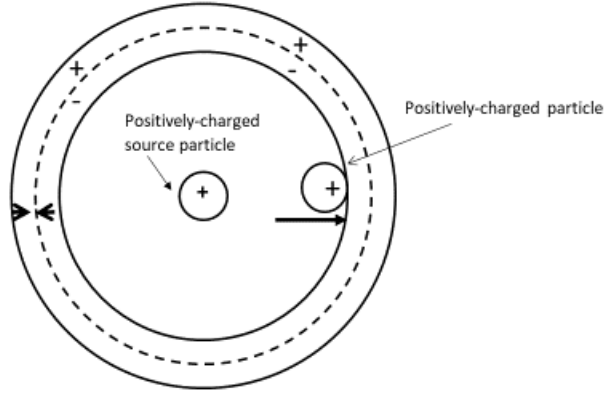


Figure 2.2 Structure of shells for electrical repulsion. The arrows indicate the direction of ideal liquid flow.

$\sqrt[3]{\frac{3}{4\pi}V_n} - \sqrt[3]{\frac{3}{4\pi}(V_n - 1)}$. By the first term of a Taylor expansion this difference is approximately $\frac{dR_n}{dV_n}\delta V_n = \frac{dR_n}{dV_n}$ since by hypothesis δV_n is 1. Now

$$\frac{dR_n}{dV_n} = \frac{d}{dV_n} \sqrt[3]{\frac{3}{4\pi}V_n} = \frac{1}{\sqrt[3]{36}} \frac{1}{\pi^{1/3}V_n^{2/3}} = \frac{1}{4} \frac{1}{\pi} \left(\sqrt[3]{\frac{3}{4\pi}V_n} \right)^{-2} = \frac{1}{4} \frac{1}{\pi} \frac{1}{R_n^2} \quad (2-1)$$

So, approximately

$$\delta R_n = \frac{dR_n}{dV_n} = \frac{1}{4\pi R_n^2} \quad (2-2)$$

Since the electric force field is a central force field and since, as I will explicate next and is illustrated in Figures 2.1 and 2.2, in my model electric forces are constituted by ideal liquids flowing into each other in a direction perpendicular to the circumference of the corresponding thin shell, Coulomb's Law

$$\mathbf{F} = \frac{1}{4\pi\epsilon_0} \frac{Q_1 Q_2}{r^2} \mathbf{r} \quad (2-3)$$

follows. Q_1 and Q_2 are the charges of two respective charged particles, r is the distance between them, $\mathbf{r}$ is a unit vector in the direction between the two charges and ε_o is the electric permittivity constant of free space.

The respective forces of electrical attraction and of electrical repulsion are respectively illustrated in Figures 2.1 and 2.2. Notice that for electrical attraction the interaction between a four-dimensional charged particle and a thin shell constituting the electric field of an oppositely-charged particle is with the outside of the thin shell. For electrical repulsion the interaction with an identically-charged particle is with the inside of the shell. It should also be emphasized that I stipulated in Chapter One that when the ideal liquids constituting a thin shell oscillate, it is into an orthogonal set of dimensions and not parallel ones. Thus, there is no causal interaction with the liquids in this orthogonal set. The forces are exerted by means of the ideal liquids from the thin shell and the particle flowing into each other as shown in the figures. The total force is then given by the vector sum of each impinging shell; i. e., the net force $\mathbf{F}$ is given by

$$\mathbf{F} = \sum_n \mathbf{F}_n = \sum_n \frac{e^2}{r_n^2} \mathbf{r}_n \quad (2-4)$$

where e is the unit charge of an electron and r_n is the distance between the n^{th} charge and the particle. Notice also that the flowing of the ideal liquids constitutes the force in the sense that it produces the resulting movement of the particles by means of being attached to them. In other words, it causes the acceleration in accordance with Newton's Second Law of $\mathbf{F} = m\mathbf{a}$. It should be emphasized again that the forces involved here are contact forces within the context of the model *per se*, and thus need not at least also correspond to any forces in classical electrodynamics such as the Coulomb force or the Lorentz force, although there are bridge principles between the two as previously discussed. The terminological point which I made in Section 1.1 should also be re-emphasized, namely that I am using the term “field”

here to refer to the component fields of individual shell systems and not to the total resultant field unless I so specify.

I will leave it as an open question whether the shell on the opposite side of an impinging four-dimensional charged particle is the immediately adjacent shell as the one on the front side. If this shell is the immediately adjacent one issues arise concerning that the shell will no longer be perfectly spherical – these issues becoming more acute both in cases of extremely dense matter being impinged on and when the radii of the shells have become extremely thin due to extreme distances from their originating particles. If the shell on the opposite is not the immediately adjacent shell to the one on the front side then “contiguity” issues arise concerning the manner in which the shells either have a gap in them or continue to exist at the location of the four-dimensional charged particle in question.

I now turn to my discussion of the electrokinetic case where there is relative motion between two charged particles. The one point which I wish to make for this case is that an obvious complication for my model is that the relative motion between thin shells and the charged particles they interact must, at least for the most part, be “frictionless” in order to be in accordance with the known laws of electromagnetism. In particular, the shells and particles must “glide by” each other “seamlessly” in such situations as where the shells are on the sides of the particles, in the sense that both the internal characters and subsequent motions of the shells are not causally affected by these particles. Unfortunately, I do not know how to account for this property in a non-*ad hoc* manner and so simply stipulate it.

I now turn to my discussion of the potential energy properties of the electric field. It can be recalled that according to standard electromagnetic theory, the electric scalar potential ϕ is invoked here whereby a scalar magnitude (the potential energy of the electric field) is associated with each point of space such that the gradient vector function of these magnitudes is the electric field; i. e., $\mathbf{E} = -\nabla\phi$. Since the field is the spatial derivative of the potential, electric potentials decrease at a $1/r$ rate from a source particle while the energy density of the electric field

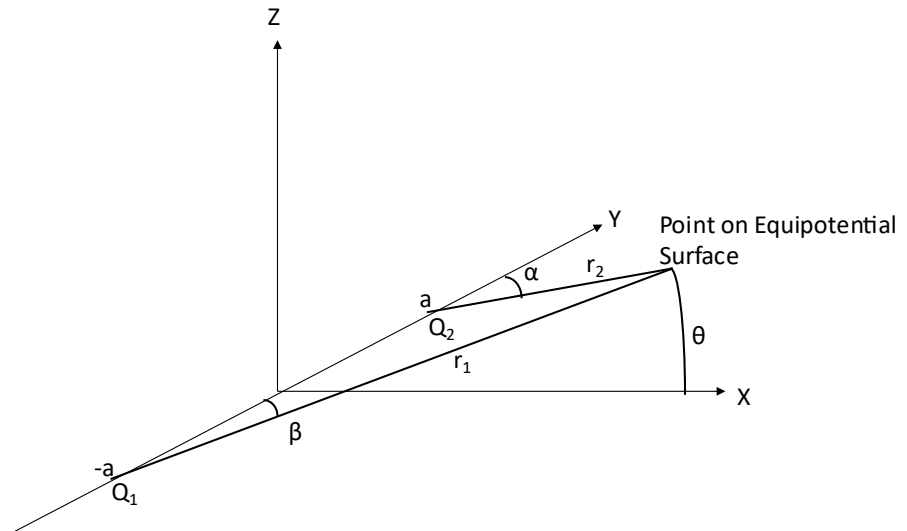


Figure 2.3 Bipolar Coordinate system for an equipotential surface

decreases at a $1/r^2$ rate. However, since in my account I deny the independent existence of the scalar potential apart from the electric field, I do not hold that anything literally exists decreasing at a $1/r$ rate here. Still, an alternative account needs to be given of physical properties of the scalar potential, which is classically defined as the amount of work required in order to move a unit charge across a region of potential difference. As my alternative account I appeal to properties of the ideal liquids in the thin shells. It can be noted that due to the fact that the thin shells associated with a region of potential difference will have different radii from their source particles they also will possess varying widths. One possible suggestion is that the potential energy associated with an electrical potential difference involves the difference in field strengths (as just explicated in terms of thin shell widths), between the two thin shells being moved between.

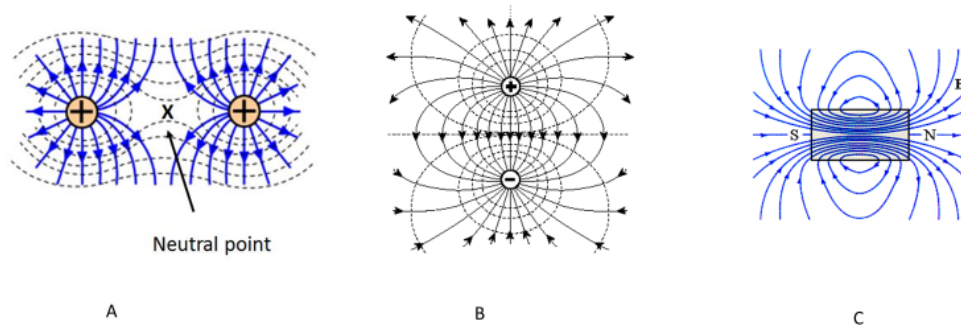


Figure 2.4 Lines of force and equipotential surfaces for A (two like charges), B two unlike charges (an electric dipole) and C a magnetic dipole

Since they also play a key role in my discussion of magnetism, I now turn to a discussion of equipotential surfaces, beginning with the single charged particle case. In standard physics equipotential surfaces are defined as being two-dimensional surfaces with the same electric potential at every point. In order to avoid undue confusions with standard terminology and since it fits in better with relationships of lines of force as they are traditionally understood, in what follows I will use the concept of “equipotential surfaces,” with the implicit understanding that in my model electric potentials are actually a fiction and that I am only making an ontological commitment to electric energy field densities. Due to symmetry considerations, it can be seen that the equipotential surfaces in the single charged particle case will be spherical, and thus correspond to the thin shells of the model. As noted, since the field is the spatial derivative of the potential, electric potentials decrease at a $1/r$ rate from a source particle while the energy of the electric field decreases at a $1/r^2$ rate. However, for the single particle case a spherical shape for the surfaces is retained in both cases with only the relative strengths of the potential

field varying as a function of their radii from a source particle. This is clearly consistent with the thin shell model inasmuch as under that model individual thin shells exist for each of the various cases of both equal potentials and equal field strengths. This is not the true for multiparticle cases though inasmuch as the equipotential surfaces in these cases are typically not spherical.

I now move on to the multiparticle case; first for a two-charge system and then for an indefinite number of charges. I use the phrase “two-charge system” rather than “dipole” since conventionally “dipole” refers only to systems with opposite charges. Equipotential surfaces for two spatially-separated particles with the same unit charges approximate spheroids as their mutual radii increase. This can be most easily shown by using a version of bipolar coordinates given by Rudolf Luneburg (1947, p. 10). The bipolar coordinates α , β , and θ are centered at the two particles Q_1 and Q_2 where α and β give the respective angles from the locations a and $-a$ of the two particles to a point on an equipotential surface, whose polar elevation is given by θ , is illustrated in Figure 2.3. When normalized, the transformation equations to Cartesian coordinates are

$$\frac{x}{a} = \frac{2 \cos \theta}{\cot \alpha + \cot \beta}, \frac{y}{a} = \frac{\cot \alpha - \cot \beta}{\cot \alpha + \cot \beta} \text{ and } \frac{z}{a} = \frac{2 \sin \theta}{\cot \alpha + \cot \beta}. \quad (2-5)$$

The equation for an equipotential surface of a two-charge system is given by

$$\frac{q}{4\pi\epsilon_0 r_1} + \frac{q}{4\pi\epsilon_0 r_2} = c \quad (2-6)$$

where r_1 is centered at one charge and r_2 is centered at the other charge, q is the unit charge of an electron ϵ_0 is the permittivity constant and c is a constant. As r_1 and r_2 increase, as previously pointed out, the resulting shape approximates a spheroid where the two radii are centered at its two foci. It should be emphasized that there is no thin shell which corresponds to this surface though. Still, it is “as if” such a surface existed inasmuch as the equipotential surface in this case corresponds to the resultant (summed effects) of each of the shells comprising the dipole. This case is

illustrated in diagram A of Figure 2.4. Notice that the lines of force are perpendicular to the equipotential surfaces for each of the charged particles. Also notice that the only actual force field in my model corresponds to the central straight-line one. I hold that the other curved lines of force which are depicted do not exist, and are just defined as being perpendicular to the equipotential surfaces. I now extend my analysis to the indefinite multiparticle case.

For the indefinite multiparticle case, the resultant equipotential surface is constituted by the summation

$$\sum_{i=1}^n \frac{q}{4\pi\epsilon_0 r_n} = c \quad (2-7)$$

over the equipotential surfaces for individual charges, where n is the number of charged particles. It can be noted that the indefinite multipole case is reminiscent of the case of configuration space using $3N$ coordinates where N is the number of particles. Notice though that unlike the usages in both classical mechanics and quantum mechanics, I hold that there is nothing unphysical about the spaces in which these coordinate systems are centered.

It should be emphasized that under this model electric forces do not occur within individual charged particles themselves since they do not possess thin shells as part of an interior structure. This will be a key point with respect to the issue of the stability of a model of extended charged particles which I give in Chapter Five on Atomic Physics; i. e. since otherwise it might seem that the parts of a charged particle would cause it to burst apart due to the electrical repulsion among these parts. The account just given of the electric field obviously is in need of further elaboration but I will not do so here. Instead, I will move on to my discussion of the magnetic field. Unfortunately, this discussion is even more sketchy than that of the electric field.

2.2 The Magnetic Field

It is sometimes pointed out that in several key respects the laws of electricity and of magnetism are duals of each other. “Duality” in this context refers to the fact that systematic substitutions can be made between $\mathbf{E}$ and $\mathbf{B}$ with different

fundamental laws of electrodynamics and with Maxwell's equations in particular. Consider Ampere's law $\nabla \times \mathbf{B} = \mu_0 \epsilon_0 \frac{\partial \mathbf{E}}{\partial t}$ and Faraday's law $\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$. When $\mathbf{E}$ and $\mathbf{B}$ are interchanged in the formulas the respective laws are also reversed. Something analogous also occurs with Gauss's law $\nabla \cdot \mathbf{E} = \frac{\rho}{\epsilon_0}$ and Gauss's law for magnetism $\nabla \cdot \mathbf{B} = 0$. The difference between the laws $\frac{\rho}{\epsilon_0}$ where ρ is the charge density and ϵ_0 is the electric permittivity of free space.

There are also a few additional parallels between the electric and magnetic cases. For example, there are two types of charges (positive and negative) and with magnetism there are two types of poles (north and south). Also, there is a partial parallel in that the fact that the electric field is a central force field is paralleled by the fact that the forces either towards or away from a magnet are maximized at the poles of the magnet. Care must be taken in pushing the foregoing analogues between electricity and magnetism too far. For example, while isolated electric charges can exist, this is not so with the magnetic case since there are no magnetic monopoles. Also, magnets are at least typically electrically neutral while they still contain north and south poles.

As with the electric field, I claim both that the inverse square rate of the intensity of the magnetic field strength is due to the corresponding width of thin shells, and also that there is a superposition of effects from different shell systems. Unlike the case of electricity, there are two laws which can be used to express the inverse square law for the nature of magnetism – the Biot Savart law and the Coulomb law for magnetism. The Biot Savart law is more commonly cited than the Coulomb law in this context and is given by:

$$d\mathbf{B} = \frac{\mu_0 i}{4\pi} \frac{d\mathbf{l} \times \mathbf{r}}{r^2} \quad (2-8)$$

where $\mathbf{B}$ is the magnetic induction vector, μ_0 is the permeability constant of free space, i is a current, $\mathbf{r}$ is a unit vector in the direction of a point away from the current and $d\mathbf{l}$ is a tangent vector to the current at a given location. It can be observed that the resulting field is a vector field. Coulomb's law of magnetism is:

$$F = K \frac{m_1 m_2}{\mu r^2} \quad 2-9$$

Here m_1 and m_2 are the respective pole strengths, K is Coulomb's constant and μ is the magnetic permeability.

Using the Biot Savart law in the form of equation 2-8 has the disadvantage in the context of the thin shell model that the magnetic field $\mathbf{B}$ is zero in directions both parallel and antiparallel to the direction of a current (since the cross product $d\mathbf{l} \times \mathbf{r}$ in the numerator of the law goes to zero then). While the Coulomb law does not have this problem it is a scalar law and I make use of the vector character of the Biot Savart law in my later analysis of magnetic moment vectors. The Biot Savart law dates back to 1820 and can be derived from Maxwell's equations. It thus need not be based on the special relativity construal of magnetism.

In his famous 1905 paper "On the Electrodynamics of Moving Bodies" Albert Einstein claims that the fact that reciprocal motion between a magnet and a conductor produces a current regardless of which is considered to be at rest suggests that the laws of physics are invariant over all inertial reference frames. I agree with Einstein that electromagnetic induction is due to relative motion but disagree with his conclusion that this shows that the laws of physics are invariant among all inertial reference frames. In fact, as far as I can see, the reciprocity as far as the cause of the current in the conductor is concerned just shows that the laws of one frame which is stationary with respect to its source particle are equivalent with the laws of another frame with respect to its source particle. Also, unlike Einstein I do not base my account of magnetism in terms of a relativistically-induced Lorentz spatial contraction of a moving electric field as sketched by Feynman (1964, Vol. 2, Ch. 13).

Interestingly, it may be possible to test between the relativistic and other accounts of magnetism such as one which Jean de Climent (2014) has put forward. De Climent postulates that the magnetic field is due exclusively to alignments of the magnetic moments related to the "spin" of individual electrons and not to their external motions such as their speed with respect to a wire, which the relativistic

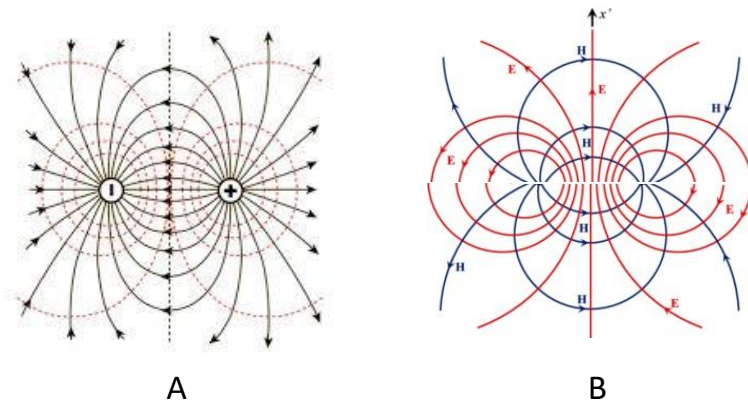


Figure 2.5 Diagram A shows field lines and equipotential surfaces for an electric dipole while diagram B shows magnetic field lines H and electric field lines E from two wires perpendicular to the page. See Richard Becker, 1964, p. 257.

Lorentz contraction account is based on. De Climont also claims that his interpretation is testable by having an electron beam in a Perrin tube be deflected by 90° with an electric field and seeing whether the magnetic moment of the electrons is changed. Since the magnetic moments of the electrons will no longer be parallel to the direction of their translation their magnetic field should disappear under the Lorentz contraction account. It can also be noted that Einstein's attempt to derive the Lorentz transformations from his two postulates of relativity in his 1905 paper has been seriously questioned – see Alasdair Beal, 2024. Despite claims to the contrary, Einstein clearly knew about Lorentz's work in advance since he, like Lorentz, uses β to refer to $\frac{1}{\sqrt{1-v^2/c^2}}$ and not to the modern usage of $\frac{v}{c}$.

I now wish to make a few points concerning prospects of basing a magnetic field ontology on a literal construal of the existence Faraday's $\mathbf{B}$ lines of force. While it may be possible to attempt to give a literal construal of $\mathbf{B}$ lines of force I

do not do so for a variety of reasons. As I pointed out in Section 1.2 it is not clear that there is an ontological format in which one-dimensional lines of force can exist and there are also issues concerning the lines getting tangled up with each other. Still, even though I do not make an ontological commitment to these lines, as I will explicate shortly, I do postulate a system of vortices within the thin shells which play at least some of the same roles. Admittedly, historically, the lines were not thought of as rotating, but still one is a functional analog of the other. Thus, I attempt an approach which I believe is consistent with the thin shell model, although admittedly this model is incomplete. In Chapter Three I reject the dipole model of individual-atom radiation. Still, for the single charged particle case (e. g., an electron) I hold that the particle possesses a magnetic moment which I will analyze in terms of vorticity in my treatment of spin in Chapter Four.

I move on now to my discussion of cases like bar magnets systems which are not in obvious relative motion. I begin with the two-particle case, which I term a “magnetic dipole” with the provision that my usage differs to some extent from the standard physics usage. As with the electrical field case, I utilize bipolar coordinates, in this case centered at each pole. I then move on to the indefinite-number case by generalizing the use of bipolar coordinates centered at the various particles in a manner analogous to the way I dealt with the indefinite-number case for the electric field.

To begin my discussion of a magnetic dipole it can be observed that, using the lines of force formalism, magnetic field lines of force are always perpendicular to electric field lines of force. This is illustrated in Figure 2.5 A for the case of electrical currents flowing in opposite directions (and thus with charges moving with respect to each other) in two wires perpendicular to the page. An isomorphism can now be observed with Figure 2.5 B for an electric dipole where the equipotential surfaces are perpendicular to the electric field lines of force. It might be noted that this configuration can be used to demonstrate Ampere’s force law

whereby wires with currents in the same direction repel each other and in opposite directions attract each other.

Inasmuch as I do not know of any clear-cut cases where there is a divergence between the respective shapes of electric dipole equipotential surfaces and magnetic fields, I will assume that the two are identical and thus identify the equipotential surfaces with magnetic fields. It is just the case that the electrical forces are normal to the surfaces while the magnetic forces are perpendicular to the surfaces. Also, in the pure electrical field case the equipotential surfaces will be stationary in rotation with respect to the primitive rate of rotation given by Ernst Mach's principle as the rotation rate of the fixed stars or alternatively by Newton's water bucket experiment where the water surface is flat. In contrast, in the magnetic field case the thin shells comprising the equipotential surfaces will be rotating with respect to the primitive rate – one direction for the north pole of the magnetic dipole and the opposite direction for the south pole. As with the case of the relationship between **E** lines of force and electric equipotential surfaces, I postulate the existence of magnetic equipotential surfaces rather than the corresponding **B** lines of force.

I base my account of the magnetic field in terms of rotations of the ideal liquids comprising thin shells when there is relative motion among the thin shells from different systems. To some extent I base this fluid dynamic account of magnetism on earlier models given by Maxwell and Hermann Helmholtz. Maxwell (1861) gave a mechanistic model of the magnetic field with a system of gears complete with idle wheels to allow the gears to spin in the same direction. However, Maxwell also cites with approval a hydrodynamic model of the magnetic field postulated by Helmholtz (1858). It is perhaps also worth mentioning that William Thomson (Lord Kelvin) used some of Helmholtz's ideas regarding vortices in ideal liquids with his model of vortex atoms (Thomson, 1867). It is possible to attempt models of the magnetic field in a shell system both in terms of postulating a constant angular momentum for the whole shells and for a constant rectilinear velocity as a

function of the polar angle within the shells. Since the former case would appear to have problems accounting for the opposite poles of a magnet, I just deal with the latter case. I now turn to the details of my fluid dynamic model for that case for a two-particle system.

In my version of fluid dynamic model of magnetism for a two-charged-particle system, the density of the magnetic field $\mathbf{B}$ at a particular location corresponds to the curl of the fluid velocity field $\mathbf{\Omega} = \nabla \times \mathbf{v}$ where $\mathbf{v}$ is the relative velocity field of the liquid flow and $\mathbf{\Omega}$ is the vorticity vector of an ideal liquid at that location. Vorticity is a fluid dynamic concept corresponding to the circulation per unit area in the limiting case of an indefinitely small loop. Being a curl the vorticity itself is a macroscopic property. However, being the area of an infinitesimal loop, the regional subject matter of vorticity is infinitesimal and thus a property in the small. Parenthetically it might also be mentioned that an alternative name for the curl is "rotation" (abbreviated "rot") which emphasizes the rotary nature of what it is used to characterize. I might also mention that Feynman (1964, Vol. 2, p. 40-5) points out analogies between the rules of vorticity and the laws of magnetostatic (in the traditional physics sense of the magnetic fields of non-varying currents) fields in free space.

It can be noted that the curl of the magnetic vector potential $\mathbf{A}$ is also the magnetic field density since $\mathbf{B} = \nabla \times \mathbf{A}$. Thus, it is tempting to identify the magnetic vector potential with the velocity field. It can also be noted that the magnetic field is invariant under gauge transformations which would correspond to the choice of different reference frames for the velocity field. Thus, a reference frame for the velocity field can be chosen so as to work with the rest of the system. In particular, the so-called transverse gauge $\nabla \cdot \mathbf{A} = 0$ (also called the Coulomb or radiation gauge) can be chosen so as to account for the propagation of light. It can also be noted that there is evidence for the existence of the vector potential with the Aharanov-Bohm effect where electrons in a long solenoid show a phase shift in the

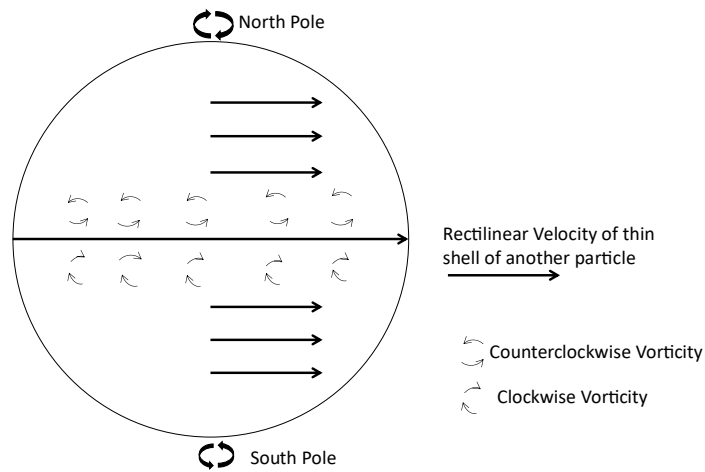


Figure 2.6 The magnetic field of a two-charged-particle system in the rest frame of one of the particles in relative motion with the other.

absence of either an electric or a magnetic field. An adequate theory would have to account for this phenomenon and hopefully someone can do so in terms of the velocity field.

Inasmuch as magnetic phenomena involve relative motion among electric charges, it follows under the thin shell model they also involve relative motion between charged particles and adjacent thin shells. My claim then is that 4-dimensional charged particles “drag” against thin shells in such a manner as to create a constant rectilinear motion. This is illustrated in Figure 2.6. Notice that the direction of the axis of the magnetic pole is determined by the direction of this rectilinear motion. It can also be observed that the constant rectilinear motion corresponds to “streamlines.” It can also be observed that these streamlines result in small “eddies” which rotate counterclockwise above the “equator” and clockwise below the equator. This situation can be described mathematically by vorticity vectors $\mathbf{\Omega}$ inasmuch as the resulting angular velocities $\mathbf{\omega}$ will vary as a function of the secant of the polar angle – see Batchelor, 1967, sec. 7.7 The rotational reference frame for zero rotation here corresponds to the primitive rate of rotation. That this

is the primitive rotational rate for the magnetic field can be readily ascertained by rotating ring magnets on a pole and noticing that a ring magnet on the pole above a rotating one does not also rotate.

The magnitude of the total magnetic field $\mathbf{B}$ of a thin shell will be proportional to the definite integral of the vorticity vector $\mathbf{\Omega}$ over the thin shell as given by

$$\mathbf{B} \propto \int_{-\pi/2}^{\pi/2} \mathbf{\Omega} \cos(\Phi) d\Phi \quad (2-13)$$

where Φ is the polar angle of the thin shell. The differential $d\mathbf{B}$ in turn is an "in the small" property of the thin shell corresponding to the vorticity at a given location. Of course, there are singularities at the poles with such an account, but other than to take note of them, I will not attempt to deal with them in this book.

Looking ahead, I might note that there are isomorphisms between the diagram in Figure 2.6 in and the diagram for polarization in Figure 3.4 in Section 3.1 and the diagram for spin in Figure 4.4 in Section 4.2. As I will argue in these sections the fact that the direction of the axes of rotation in these diagrams is determined by the direction of rectilinear velocity is a key point in the explanation of why the introduction of a third element in between two others (oriented oppositely to each other so as to cancel a given outcome) and oriented at a different angle from both of them in an experimental setup brings back the outcome. Also, in my discussion of spin I take note of a possible connection with the magnetic field, whereby spin up and spin down states are correlated with the north and south magnetic poles and macroscopic magnetic effects are caused by the manner in which alignments of spin occur among electrons from different atoms, where these spins were previously randomly aligned. I should emphasize that this sketch of an account is seriously incomplete in that it does not specify a "linking mechanism" between these spins and various stationary magnetic "phenomena" such as the fact that opposite poles of two different magnets attract each other and that the same poles repel each other. A step in the right direction in an explication of these

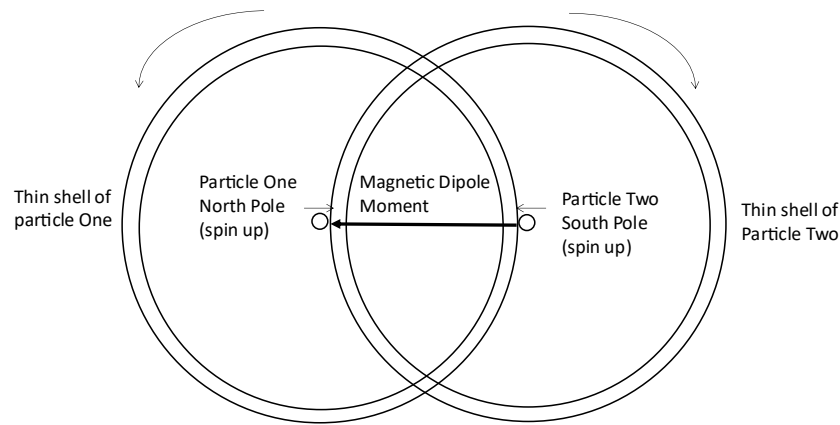


Figure 2.7 “Magnetic Dipole”

phenomena might be to note the existence of a centrifugal force which is perpendicular to the axes of the spins and I will develop this move to some extent in my subsequent discussion. It should be admitted upfront though that a subtle reworking of some of the relevant concepts being used in the model may be necessary in order to adequately handle these cases.

As far as I can tell there is no good empirical evidence for the existence of isolated magnetic monopoles. The lack of isolated magnetic monopoles fits in with my account inasmuch as both clockwise and counterclockwise rotating eddies are created in pairs as I just sketched. Also, since I account for magnetism in terms of the relative motion among charged particles, under my account the simplest magnetic case will consist of two charged particles in relative motion with respect to each other (i. e. are current elements) where one particle is a north magnetic pole and the other a south magnetic pole. This “magnetic dipole” case is analogous to the electric dipole case. Even though I do not hold that isolated magnetic monopoles physically exist I am still justified in using the Biot Savart law by interpreting it as referring to pseudo- magnetic monopoles. The law gives predictions for moving

charges which are equivalent to magnetic monopoles from the thin shell point of view.

Since, under the shell system charged particles carry their fields with them it follows that for the two-particle cases the reciprocal fields from each particle impinge on their respective opposing particles resulting in mutual “pushes” and “pulls” on those particles. Thus, the resulting magnetic forces can be explicated in terms of the actions of fields from the two respective poles acting on their reciprocal particles. Notice that these forces are perpendicular to the line segment connecting the particles (the direction of the electrical force between them) and thus are parallel to the directions of the shells per se. This is illustrated in Figure 2.7 under my construal of a two-particle “magnetic dipole”; e. g., an electron pair with opposite spins.

The analog of equations 2-6 and 2-7 where the scalar strength of an electric field using bipolar and a generalized form of bipolar coordinates were given, can now be used for the case of the magnetic field where the respective coordinates are centered at each magnetic pole. As with the Biot Savart law, I use the vector version of these equations, in this case using the magnetic dipole moment μ in the numerator. The analog for equation 2-6 then is:

$$\frac{\mu_1}{4\pi\mu_0 r_2} + \frac{\mu_2}{4\pi\mu_0 r_2} = c \quad (2-10)$$

where μ_0 is the magnetic permeability of free space and μ_1 and μ_2 are the respective magnetic dipole moments. I connect these magnetic moments with the spin up and spin down states of the electron pair in Section 4.1 where I give a non-relativistic account of spin. In the multi-dipole case for permanent magnets and electromagnets in the presence of electric currents the magnetic dipoles are aligned. The analog of equation 2-7 for the generalized case then is:

$$\sum_{i=1}^n \frac{\mu_n}{4\pi\mu_0 r_n} = c \quad (2-11)$$

where μ_n is the n th particle's magnetic dipole moment. I now turn to my discussion of the indefinite number of particle case for magnetic forces.

In my model magnetic forces are accounted for in terms of centrifugal forces associated with vortices on other particles. It can be noted that according to classical electromagnetic theory, the magnetic field exerts a tangential force on a moving charge in accordance with the Lorentz force

$$\mathbf{F} = Q(\mathbf{E} + \mathbf{v} \times \mathbf{B}) \quad (2-12)$$

where Q is the charge of a particle moving with velocity $\mathbf{v}$. It can be noted that the $\mathbf{v} \times \mathbf{B}$ portion of the Lorentz force law was already given by J. J. Thomson in 1881 who postulated it in virtue of experiments diverting cathode rays (free electrons) with magnetic fields. Thus, like the Biot Savart law, the law predates relativity.

I now turn to a discussion of kinetic systems where, as with applications of the Lorentz force, there is clear motion among the different systems involved. I begin by discussing vector (directional) properties of the magnetic field. Since both magnetic fields and magnetic forces are perpendicular to the line segments connecting particles with their thin shells (which is the direction of electrical fields and forces), I identify magnetic forces with rotations of these shells with their directions being perpendicular to the direction of the corresponding magnetic fields. I now briefly sketch a few previous attempts to base magnetism on rotations and then show how my model can utilize rotations to account for the Lorentz force.

An obvious complication for any multiparticle interpretation of magnetism in terms of thin shells is that, even if the $\mathbf{B}$ lines of force are interpreted as merely a calculation device, these lines are typically not straight but instead are typically more convoluted than even $\mathbf{E}$ lines of force. For example, in the case of a bar magnet the lines of the magnetic field are conceived of as circling back on each of the poles of the magnet. This raises issues concerning how well it fits in with the rest of the thin shell model since under that model each thin shell is construed as being located at a fixed distance from its respective source particle. It might appear that any rotations associated with magnetic fields would have to be located within

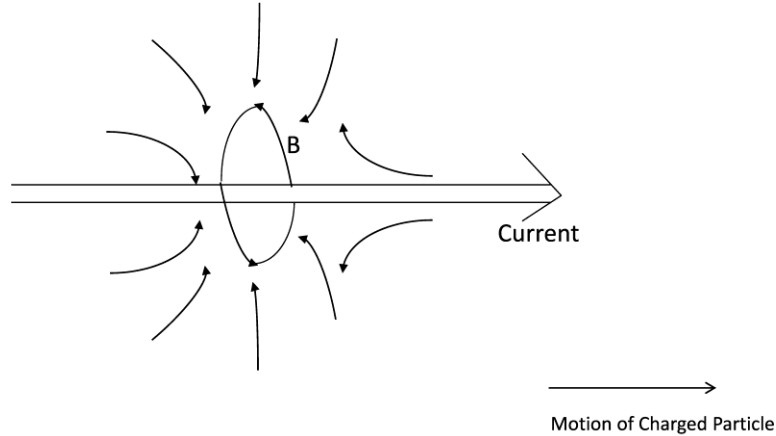


Figure 2.8 Deflection of a charged particle in a magnetic field **B** (using the right-hand rule). Notice that the **B** force is perpendicular to both the **B** field and to the rectilinear velocity vector **v** of the charged particle

the confines of an individual thin shell and not from among different thin shells. A possible strategy for dealing with this point is to claim that instead of being construed as a single system the postulated circling magnetic field lines are to be explicated as being a result of many individual thin shell system.

It can be noted that the tangential character of the magnetic force is also exhibited by the magnetic dipole moment μ (exhibited by bar magnets and current loops) which results in a torque $\tau = \mu \times \mathbf{B}$. Thus, unlike the electric field for a single charge, the magnetic field is not a central force field. Also, magnetic effects always involve relations between charges moving with respect to each other. According to my model they also involve moving respective fields. Thus, the moving fields create vorticity effects on the particles. Note that the net effect of the Lorentz force law remains the same if the force fields **E** and **B** are replaced by the corresponding reconstructed field concepts possessing both scalar (for energy density) and vector (for the force direction) features. In particular, the operation of the cross product $\mathbf{v} \times \mathbf{B}$ needs to be reconstructed here so as to also result in a microscopic rotational effect.

The foregoing account has also been developed by Charles Lucas (2013) in terms of the claim that Lenz's law (that an induced current and its accompanying magnetic flux will appear in such a direction as to oppose the change that produced it) is not reducible to Faraday's law where it is responsible for the negative sign of the law. Instead of construing Lenz's law this way Lucas claims that since currents inherently involve moving charges Lenz's law should be considered as a separate law from Faraday's law which just applies to static cases. He also holds that the law can be explained in terms of electrical feedback effects on finite-sized charged particles, instances of which would be the source particles of the fields involved.

It should be emphasized that it is only the effects of the magnetic force on a moving charged particle which are observed and not the $\mathbf{B}$ field itself which is perpendicular to the force. This is illustrated in Figure 2.8, which shows the result of using the right-hand rule for the cross product $\mathbf{v} \times \mathbf{B}$ in the Lorentz force. As Maxwell (1873/1954, pt. 3, p. 470) pointed out, this is like a centrifugal (from the Latin for "away from the center") tangential force where a rotating body - e. g., an eddy in an ideal liquid - exerts a torque on an adjacent body. A possible connection could then be made between repulsive magnetic forces and centrifugal forces from rotations in the same direction (either clockwise or counterclockwise) and between attractive magnetic forces with centrifugal forces and rotations in opposite directions. This also ties in with my discussion of "spin up" and "spin down" in Section 4.2. Obviously, these topics need further development though.

Of course, the issue of the $\mathbf{B}$ force being perpendicular to the velocity vector $\mathbf{V}$ of the charged particle along with being perpendicular to the $\mathbf{B}$ field itself needs to be addressed in order to make the model work, and Maxwell does not address this issue. A key difference between my model and Maxwell's is that under my model though I hold that the electrons comprising the current carry their fields with them. It can also be noted that the direction of this current is the same as that on which the special relativistic explication of magnetism holds that a length contraction occurs. Thus, the rationale for at least the direction of the resultant $\mathbf{B}$

force can be the same as the one which Feynman (1964, Vol. 2, Ch. 13) gives, perhaps with the added claim that the moving $\mathbf{B}$ field causes a “push” against the charged particle. While the model gets the direction of the magnetic force right, it should be pointed out that, as with the electric field, the magnetic field also possesses scalar energy density properties (for example from the way in which contributions from the fields from multiple sources combine together) and these need to be accounted for. Hopefully, someone else can do more as far as a quantitative analysis here goes.

Chapter Three

Light

In this chapter I present a physically-realist account of both wave and particle properties of light in terms of properties of electromagnetic fields oscillating at a single frequency and enclosed in finite wave packets. The doctrine of wave particle unity developed in Chapter One is invoked so as to associate photons with wave packets. These are held to divide up during elastic scattering processes although the emission and absorption processes are still held to be quantized. Much of my discussion is based on my treatment of the subject in French, 2008.

I begin my discussion by developing the electromagnetic account of light. First, I attack the dipole model of radiation at least at the level of emission from individual atoms. Instead, I give an account in terms of transverse rotational waves of the thin shells which propagate at the speed of light in all directions, with the amplitudes corresponding to energy densities. I then move on to show how that account can be used to explicate the phenomenon of light entanglement which has often been held to defy a physically-realist account. In particular, I give an account of polarization entanglement in terms of two-photon absorption and realistically-construed advanced waves. I end the chapter with a discussion of delayed-choice experiments.

3.1 An Electromagnetic Account of Light

From the work of Maxwell (1873/1954 pt. 3, ch. 20), I take some version of the electromagnetic origins of light to be well established. As Jackson (sec. 6.3) points out it is possible to derive a wave equation for light just in terms of the magnetic vector potential $\mathbf{A}$ using the transverse gauge and the electric scalar potential Φ . In this section though I will just work with Maxwell's account using the electric field $\mathbf{E}$ and magnetic field $\mathbf{B}$ per se. According to this version of the theory the propagation of light is accounted for by solving the respective wave

equations for the electric field $\nabla^2 \mathbf{E} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{E}}{\partial t^2}$ and for the magnetic field $\nabla^2 \mathbf{B} = \mu_0 \varepsilon_0 \frac{\partial^2 \mathbf{B}}{\partial t^2}$ (derived from Maxwell's equations) by a plane wave solution $\Psi = e^{i(\mathbf{k} \cdot \mathbf{r} - \omega t)}$. Here μ_0 and ε_0 are respectively the magnetic permeability and electric permittivity of free space, $\mathbf{k}$ is the wave vector, $\mathbf{r}$ is a position vector, ω is the angular frequency, and Ψ is a probability amplitude. This results in the speed of light in vacuum being given by $\frac{1}{\sqrt{\mu_0 \varepsilon_0}}$. In contrast with this account, at least for single particle sources, I hold that light consists of a spherical wave propagating in all directions and not a plane wave propagating in just one direction. It can be pointed out though that in the case of multiple sources for the waves, a superposition of such spherical waves can approximate a plane wave.

Regarding a closely related topic to the previous one, I disagree with the claim that light (at least in the case of emission from individual atoms) is the result of a dipole, as outlined for example in Jackson (1962/1988, sec. 9.2). In defense of this latter claim it can be noted that at least for the far field (the Fraunhofer case) the intensity of light decreases at an inverse square rate with respect to its source. This is consistent with the traditional concept of the energy density of a field as opposed to a force field, and it can be recalled that in Chapter One I argued that it is possible to reconcile the two conceptions. Also, I see no need to postulate dipoles in cases where there is no independent evidence for their existence. Admittedly, in the case of such multi-atom emission systems as radio antennas dipoles are involved, but I hold that the effects here are from superpositions of effects from individual atoms. Also, even for the near field (the Fresnel case) as far as I can see the additional terms for light intensity can be accounted for in terms of additional scattering processes (I expand on this topic shortly) and thus do not apply for light being emitted from a single atom. A few remarks on what sense of the concept of a wave packet which I am using are also in order. I begin that discussion by first discussing two concepts of wave packets which have been historically important

but which are difficult to construe realistically. Both of these accounts hold that wave packets consist of a range of frequencies.

The first multi-frequency concept of a wave packet which I discuss is an abstract sense used by wave mechanics. This concept is of an enclosed range of wavelets where there is a superposition of wavelets possessing different frequencies, typically an infinite number for free particles. An example of this is the continuum of frequencies associated with the phenomenon of *bremsstrahlung* when a free particle collides with a target. This results in a finite-sized wave packet due to constructive and destructive interference among the different frequencies, with the length of the wave packet being determined by the region of constructive interference. Fourier transforms are then appealed to for purposes of switching between the time domain (the uncertainty of time of emission) and the frequency domain (the range of frequencies involved).

At least according to the Copenhagen interpretation of quantum mechanics, the "waves" of wave mechanics are not interpreted realistically but instead as probability amplitudes. In fact it is not at all clear to me that a coherent realist account can be given of a packet with a range of frequencies if these are all construed realistically as occurring literally during the same period of time. In particular, two obvious difficulties for realistic interpretations are the following. One point is that these waves are construed as being located in a $3N$ -dimensional configuration space (where N is the number of particles) instead of being contained in 3-dimensional physical space. A second point is that the waves would seem to be constantly hindering and blocking each other in the sense that the occurrence of one would be in the way of, and thus blocking, the occurrence of another one. This is because distinct individual modes existing separately in a wave packet are not the same thing as a superposition of the waves whereby the amplitudes are added.

The second account of multi-frequency wave packets which I wish to discuss are purportedly-realist, multi-frequency theories of radiation that build on Max Planck's *ad hoc* hypothesis for avoiding the "ultraviolet catastrophe" with

black body radiation whereby only certain finite oscillations of electromagnetic radiation are allowed. I take it that the basic insight behind Planck's hypothesis is that the emission of radiation occurs in discrete multiples of energy $h\nu$ where h is Planck's constant and ν is the frequency of the emitted radiation. This in turn implies that the emitted radiation also possesses discrete frequencies.

An example of an "intermediary" physically realist theory which builds on that of Planck is that of Bohr, Hans Kramers and John Slater's (1924) BKS theory. This theory holds that a classical electromagnetic field contains all frequencies that an atom can either emit or absorb a photon between only certain allowed frequencies and it was held that any given atom "will communicate continually with other atoms" at these frequencies. It is noteworthy that the same blocking difficulties as those just mentioned for literal construals of wave mechanics occur for these accounts, at least if they are construed literally as occurring distinctly and simultaneously. Also, subsequent experiments (see William Cross and Norman Ramsey, 1950 for an experiment involving Compton scattering) have shown, contrary to the predictions of the BKS theory, that conservation of momentum and energy apply to individual scattering processes and not just to statistical averages of them. It is noteworthy that this last criticism does not apply to my physically realist construal where I hold that light within individual wave packets is monochromatic. I now turn to the details of that account.

In contrast to both the foregoing abstract wave packets postulated by wave mechanics and to the multiple-frequency electromagnetic accounts, with my realist concept of a wave packet, I hold that the wavelets within the wave packet are all of the same frequency (i. e., are monochromatic), with the length of the wave packet being determined by a fixed time of emission ΔT . These emission times can vary over a wide range of values but for typical processes of electron transfers within individual atoms are of an approximate magnitude of 10^{-8} s. Under this construal the light of any given wave packet will be monochromatic. Thus, the frequencies of oscillation in each individual wave packet are held to be exact, although we do

not know with certainty what they are since there is a range of possible frequencies for identical originating processes. For example, the spectra of light both from molecules and from pulsed lasers is broadband over a range of frequencies. In the case of pulsed laser Chandrasekhar Roychoudhuri (2010) has argued that the mode locking involved with pulsed lasers only physical modes are actually involved and light waves themselves do not interact. In the molecular spectrum case, with the advent of high-resolution broadband spectroscopy (see A. Muraviev, D. Konnov and K. Vodopyanov, 2020), it has been demonstrated that the apparent broad-spectrum emission bands associated with molecules are in fact made up of very fine lines corresponding to a wide range of electron transitions. Also, when individual photons are detected in both the pulsed laser and the molecular cases they possess just one wavelength.

The similarity theorem for Fourier transforms between position and frequency (whereby a stretch in the space domain results in a contraction in the frequency domain and vice versa, albeit with a change in amplitudes) hides the claim that different senses of probabilities are involved in my account of wave packets. In fact, Roychouhuri and Femto Continuum (2005) have argued that the range of physical light frequencies is typically not continuous, as is implied by the transform. It should also be emphasized again that experimental results which purport to refute BKS theory do not apply to my account. This is because the wave packets associated with photons being emitted and absorbed are at the same frequency; i. e., are monochromatic. Therefore, both energy and momentum are conserved at an individual event level according to the account.

As I pointed out in the Introduction almost all measurements of the speed of light have been two-way and not one-way. It is true that Ole Rømer in 1676 measured the one-way speed by timing eclipses of the Jupiter's moon Io, but his measurement of 227,000 km/sec was off by 25 percent and thus lacks the precision needed to refute an emission theory. Also, Einstein (1905), by stipulation, defined simultaneity in terms of the two-way measurement and devised light-clock thought

experiments and his relativistic interpretation of the Lorentz transformations to account for the null results of measurements apart from c . It can be pointed out, in counterpoint to Einstein, that in the case of measuring the two-way velocity of light between a source and a reflecting mirror moving with relative velocity magnitude v , then, assuming that the magnitude of the velocity reflected from the mirror is $c-v$, it follows that

$$\left[\frac{(c+v)+(c-v)}{2} \right] = c \quad (3-1)$$

Since v cancels out a two-way measurement is independent of the velocity of the source.

Also, I conceive of wave packets as existing in the format of wave particle unity which I developed in Chapter One. It can be recalled from my discussion of wave particle unity there that I introduced the concept of “wave-particles” there which are construed as being spread out over all physically possible paths, gliding past each other in separate parallel subspaces, and breaking up into “partial particles” when elastic scattering occurs. It should be emphasized that this concept of a wave packet differs from the traditional concept in two respects. First, it differs in that I hold that the particle, along with the wave can be divided at a beamsplitter, whereas under the Copenhagen interpretation of quantum mechanics the probability wave is split but not the particle. Secondly, it can also be recalled from Chapter One that unlike the traditional conception of a field where there is only one field which each charged particle contributes to, I hold that each charged particle carries its own electromagnetic field in a separate subspace.

From the foregoing considerations it follows that I hold a version of the emission theory of light, where it is held that the velocity of light depends on the velocity of its source. This obviously goes against Einstein’s (1905, 1920, p. 25) Second Principle of Relativity which Einstein states in terms of the claim that the speed of light is independent of the velocity of the source. It should be noted though that the Second Principle is sometimes (see Max Born, 1920/1964, p. 232) also stated in terms of the claim that the speed of light is both independent of the velocity

of the source and of the velocity of the receiver. For my version of an emission theory, this does not make a difference since I hold that it is the relative velocity between the source and the receiver which is the relevant factor.

I should also emphasize that unlike Ritz's (1908) original emission theory though, my variant is not a ballistics theory where light is conceived of as being like a bullet that retains the velocity of a gun firing it. Also, it can be pointed out that the evidence against emission theories is not that overwhelming and that in particular the Michelson-Morley experiment does not refute them since the interferometer of that experiment constitutes a new source (see the discussions of John Fox, 1965, and of Beckmann (1987, sec. 1.3). Fox points out that Willem de Sitter's (1913) argument from double stars is invalid since a gaseous corona around both stars would constitute a new source. It is true that Kenneth Brecher (1977) considers the case of the emission of X rays from binary star systems since their extinction lengths would be considerably longer. It can be noted that under Ritz's ballistic account a source is like a gun in the sense that once a photon is emitted its subsequent motion would be independent of subsequent motion of the source. In contrast, under my account since the source carries its field with it, the subsequent motion of a photon would be continuously tied to that of the source.

Counterpoints, such as those just given, can be made to at least some of the purported refutations of emission theories involving rotating mirrors such as that of Albert Michelson (1913). Michelson's experiment was conducted in air but a subsequent version in high vacuum was performed by Beckmann and Mandics (1965) also with a negative result. The claim to negative results of both of these experiments assume that light reflected from subsequent mirrors retains the added velocity from the original source. Both Michelson and Beckman (1987, p. 40) claim that their experiments also refute versions of emission theories that hold a mirror serves as a new source and where the reflected light would be c with respect to the new source. Two points can be made in response. One point is that, as with my response to Brecher's argument from X ray emissions from binary star systems,

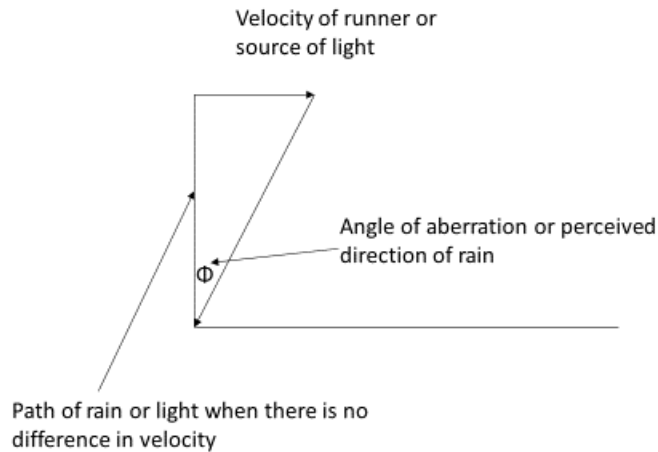


Figure 3.1 The Aberration of light under an emission theory

under my account the field moves along with the source and thus the emitted light from the moving mirror would continue to be correlated with the motion of the mirror. A second issue involves the point that Michelson cites the work of Richard Tolman (1912) as refuting emission theories holding that light is reflected from a mirror at velocity c . Interestingly, Tolman in a footnote states that he does not address theories in which the velocity of light is held to change during its path and that “it might be very difficult to test theories in which the velocity of light is assumed to change on passing through narrow slits or near large masses in motion, or to suffer permanent change in velocity on passing through a lens.” Since under my account narrow slits or lenses would constitute new sources, I do not see that Tolman refutes my version of an emission theory.

It can also be pointed out that the aberration of light is accounted for very naturally under an emission theory. The situation is analogous to the case of raindrops falling perpendicularly to the ground, where a runner will view them as

falling at a slant as is illustrated in Figure 3.1. The diurnal and orbital motions of the earth with respect to the source of the light would correspond to the motion of the runner and the angle of aberration is analogous to the perceived direction of the rain.

Emission theories also give a very intuitive explanation of the Sagnac effect where a phase shift is observed between two light beams traveling in opposite directions in a rotating interferometer. In spite of the fact that the distances among the mirrors of the interferometer remain the same when the interferometer rotates, the key point is that the total distance covered by the light changes with respect to the primitive rate of rotation – i. e., the rotational reference frame with respect to which a Foucault pendulum does not rotate and where, with respect to Newton's water bucket experiment, the surface of the water is flat. Since light only travels in a straight line with respect to this primitive rate of rotation, it must compensate for the rotation of the interferometer in order to reach the same position of a subsequent mirror as in the non-rotating case. This increases the pathlength of the light for one beam and decreases the pathlength for the other beam resulting in the phase shift between the two beams.

Another topic that comes up with respect to emission theories is the Doppler shift. As Tolman (1910) pointed out at the time of the initial disputes between the relative merits of relativity and the Ritz theory, there is only a second-order difference in the predicted Doppler effect for frequency between the Ritz theory (where the change in frequency is $\frac{c+v}{\lambda}$ where v is the magnitude of the velocity of the source) and the relativistic formula $n_0 + \frac{v}{c} + \frac{v^2}{c^2} + \dots$ where n_0 is the observed frequency of an observed spectral line. Obviously, there would be no Doppler effect for wavelength under the Ritz theory. Quirino Mojarana (1918) has discussed experiments purporting to show a difference between Doppler shifts due to wavelength properties of light and those due to frequency properties. He claims then that when light from a moving source impacts the optical components of these experiments there is a change in wavelength which cannot be accounted for under

an emission theory. As far as I can see this argument at most only applies to the original Ritz ballistic version of the theory and not also to the version where the velocity of light is c with respect to each new source since the optical components constitute new sources. It is also questionable whether the second-order term in the relativistic Doppler effect has been experimentally verified since the predicted effect is extremely small.

Beckmann (sec. 1.6) points out that there is no good direct empirical evidence to support either Lorentz's or Einstein's versions of length contraction. He also touches on alternatives to special relativity explanations to both purported mass increases in particle accelerators and purported evidence for time dilation as a function of relative velocity. I discuss his treatment of purported mass increases in my discussion of gravity in Section 5.1. With respect to purported evidence for time dilation, such as evidence from rates of fast-moving neutral pion decay (both from celestial and terrestrial sources), Beckmann questions whether this is merely evidence for physical processes slowing down rather than time itself (Beckmann, sec. 1.9.3). A possible move is to suggest that the asymmetry between the two rates is due to the fact that there are many more charged particles (and hence thin shells) in the earth-centered reference frame than in the reference frame where a neutral pion is at rest. Perhaps some principle akin to Mach's principle can then be invoked to at least partially motivate a difference in the decay rates. The principle would have to be in the context of postulating a primitive reference frame for time instead of the primitive rotational reference frame of the fixed stars which Mach postulates. I will not develop this move in more detail in this book though.

Another feature of the fields associated with wave-particles involves cases when they are not absorbed by a potential absorber - the Renninger effect. This effect was first discussed by Mauritius Renninger (1960) and the relevant measurements were sometimes known as "interaction-free" or "non-demolition" measurements. As Daniel Greenberger (1983) points out, the non-detection of a particle here keeps the phase of the wave coherent, and instead, only changes its

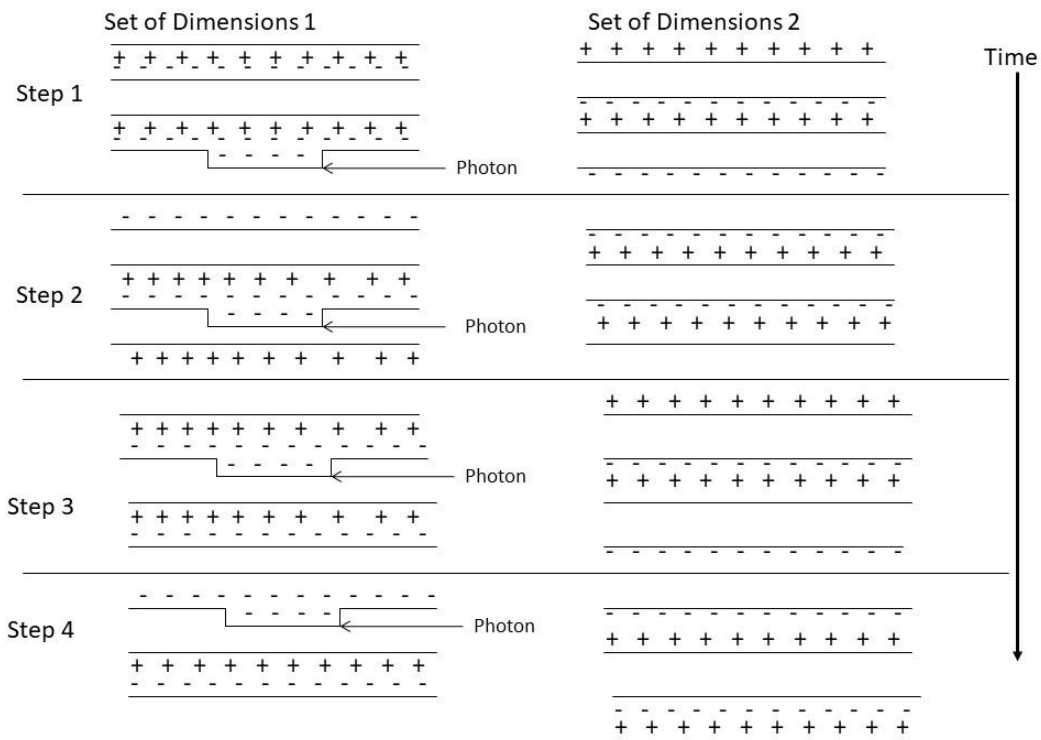


Figure 3.2 Propagation of "photons" in thin shells (by being "conducted" over the shells). The position of the shells is kept constant and the "photon" is propagated from the bottom shell to the top one in the same set of dimensions.

amplitude. I account for this in terms of the claim that these fields are "pushed aside," by a potential absorber in these cases (creating a "shadow" in the process) and thus increasing the energy density in other directions. In other words, the ideal liquids constituting particles of light (i. e., "photons") are rearranged in such a manner that they are concentrated in directions where there is a finite probability that they will be absorbed.

I now discuss the subject of particle (photon) propagation within a thin shell system and will then turn to a more extended treatment of the wave properties

of light. This treatment has some features in common with the model which Descartes(1644/1983, Part 2, par. 64, Part 4, par. 28) propounds in his *Principles of Philosophy*. Analogously to Descartes's claim that light is propagated through a series of vortices, I hold that it is propagated through a series of thin shells. However, Descartes held that light consists of longitudinal pressure waves while I hold that it is a transverse wave. Also, in accordance with what is known today, under my model light propagates at a finite speed, contrary to the claim that the speed of light is infinite which Descartes made in his Letter to Isaac Beeckman (Descartes, 1897, vol. 1, pp. 307-312).

I hold that a photon is transferred from thin shell to thin shell in the manner illustrated in Figure 3.2. As can be seen, once a photon has come into contact with the inner side of a shell, it becomes temporally attached to that side, and is then "passed on" to the next shell as the two sides of the original shell flow into each other and pass on into the other set of dimensions. This transferring process will take place at the speed of light, inasmuch as that is the rate at which the two sides of a thin shell flow into each other. Also, even though photons are portrayed as possessing charges here, it should be emphasized that they do not carry fields (i. e., their own thin shells) with them and thus do not function as being charged in any stronger sense.

I hold that light consists of spherical waves oscillating perpendicular to the line of propagation. The perpendicular **E** waves and **B** waves will be in adjacent subspaces oscillating on orthogonal axes, with the frequency of oscillation corresponding to ν , the frequency of the corresponding wavelength of light. At this point a choice needs to be made between two possible models of polarization with rotations of thin shells – one based on a constant angular velocity for the whole sphere and one based on a constant rectilinear velocity. Both of these variants in turn have variations where the rotations are constant over time and variations where the rotations oscillate. Variants where the rotations oscillate have the distinct advantage in that they clearly give a better account for interference effects since

these involve the phases of repeating oscillations. I have several considerations for choosing to opt for the constant rectilinear velocity variant over the constant angular velocity variant. For one point there is the consideration of the constancy of the velocity of light, in the context of the transverse velocity. In my model this corresponds to the maximum transverse velocities of the electric and magnetic fields in a wave packet. Also, I will argue it can better account for the opposite rotations associated with circular polarization. Notice that this transverse oscillation of the thin shells does not conflict with the account of the electric force given in Section 2.1 since the effect of the oscillations averages out so as to coincide with the direction of the electric force given there.

I now elaborate on the interpretation where the rotational waves possess constant rectilinear velocities. The rotational waves are defined as having a radial propagation of c and a transverse propagation in the form of an oscillating rotational wave with a maximum speed of c . They will also be defined as being rotationally perpendicular to each other and so as to be out of phase by π . I hold that the resulting transverse velocity holds along lines of latitude in all directions of the sphere. Notice that this constant rectilinear velocity results in the angular velocity increasing monotonically inversely proportional to the cosine of the polar angle away from the equator, as was also given in my discussion of the magnetic field in Section 2.2. This results in streamlines possessing equal velocities at different latitudes. It can also be noticed that this results in singularities at the two poles.

I now turn to my discussion of the polarization of light under my model. I hold that angles of polarization are determined by the axes of rotation of the $\mathbf{E}$ waves in accordance with the experimental results of Otto Wiener, 1890. I then correlate these axes of rotation with positions on a sphere, which is in certain respects analogous to the one postulated by Henri Poincaré (1889, Vol. 2, Ch. 12) for mapping different polarization states onto a sphere. In my version though,

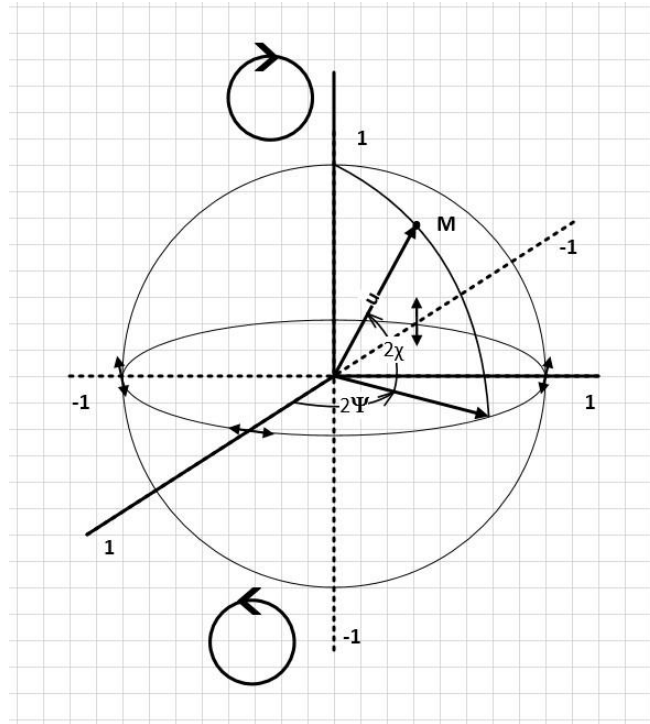


Figure 3.3 Poincaré-like sphere for polarized light. The poles represent spherically polarized light; points on the equator represent linearly polarized light and elliptically polarized light is represented by polar angles in between these two.

illustrated in Figure 3.3, oscillations along the equator of the sphere correspond to horizontally polarization, oscillations on an axis perpendicular to these correspond to vertical polarization, ones along the poles to circular polarization, and intermediary positions to elliptical polarization. Of course, there are singularities at the poles, but as I stated in my discussion of magnetism in Chapter Two, I will not deal with this issue in the book.

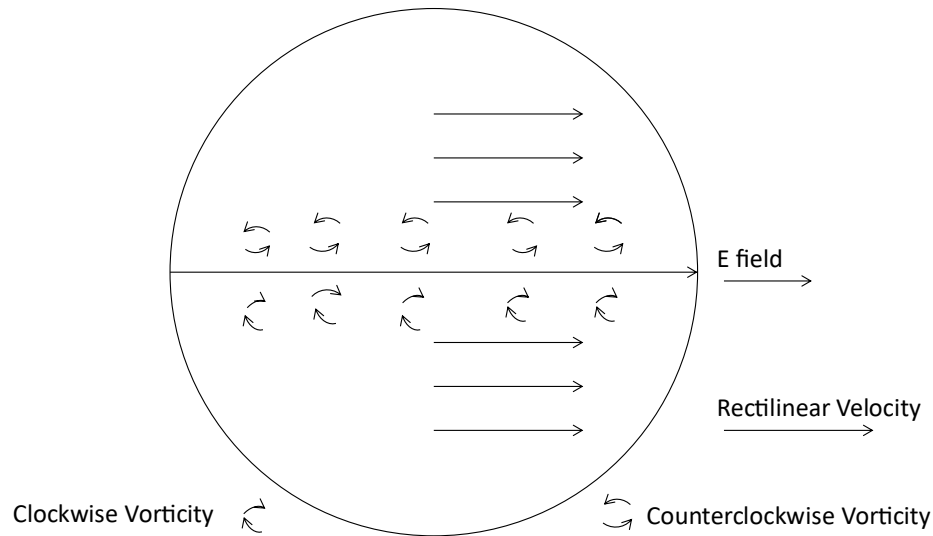


Figure 3.4 The relations between the HV (in the large) and the helicity (in the small) bases for polarization.

Figure 3.4 illustrates the connection under my model between characterizations of polarization states for the HV (horizontal – vertical) basis and for the helicity basis. An isomorphism can be noted between the diagram and Figure 2.6 for magnetism. This isomorphism suggests that there is at least a linkage (if not an identity) between the two “mechanisms.” Using the HV basis, polarization states are characterized in terms of a function of the vectors representing horizontal and vertical polarizations. In contrast, with the helicity basis, polarization states are characterized as a function of the vectors for right-handed (clockwise) circular polarization and left-handed (counterclockwise) circular polarization. In my model, the helicity basis involves the direction (clockwise or counterclockwise) of rotations in the small. In contrast the HV basis involves rotations in the large along

an axis for the entire sphere. Due to the resulting symmetry of increasing angular velocities on each side of the equator, it can be seen that the constant rectilinear velocity over a spherical surface results in microscopic eddies possessing opposite directions of rotations on opposite sides of the equator. This is illustrated in Figure 3.4, and results in opposite rotations for the two circular polarizations at the two poles. An experimental confirmation of such a realistic construal of circular polarization is that a twist is given to matter when it interacts with circularly polarized light (see, for example, Galstyan and Dranoyan, 1997). Also, the well-known experimental fact that placing a third polarizer at 45° in between two orthogonal polarizers results in some light coming through, fits in well with my account since it holds that new fields (with their own polarizations) are created by each polarizer instead of holding that the polarizers act like filtering devices.

The equations for the effective angular velocities Ω_1 and Ω_2 for two oscillating fields at a given angle of polarization can now be defined as

$$\Omega_1 = (c/r)a\cos(\omega t) \quad (3-2)$$

and

$$\Omega_2 = (c/r)b\sin(\omega t) \quad (3-3)$$

where r is the radius from the source particle at time t , ω is the angular frequency of the light wave, and $\mathbf{a}$ and $\mathbf{b}$ are orthogonal radial unit vectors centered at the source particle and aligned with the axes of rotation respectively of the $\mathbf{E}$ and $\mathbf{B}$ fields. The respective field strengths here would involve the inclusion of proportionality factors in equations 3-1 and 3-2. I now connect with a construal of probability amplitudes for absorbing light with the path integral approach of Richard Feynman.

3.2 Feynman path integrals

In this section I link the foregoing account of light in terms of oscillating electric and magnetic fields with a realist construal of Feynman's path integral approach of quantum mechanics – see Feynman and Hibbs, 1965. In particular, I construe the paths which Feynman sums over realistically in terms of actual

physical trajectories of a particle. I begin by discussing the linkage between rotational waves and probability amplitudes, whose squares constitute the probabilities for events.

I treat probability amplitudes as scalars (Erwin Schrödinger's treatment) rather than as vectors (as with Dirac's bra $\langle\psi|$ and ket $|\psi\rangle$ state vectors) but presumably a treatment with state vectors could be done as well. I deal first with the case where there is just a single path linking the particles emitting and subsequently absorbing the photon, and then deal with the multiple path case. In the single-path case I wish to introduce two probability amplitudes Φ_1 and Φ_2 . These correspond to the real (cosine) and "imaginary" (sine) terms of the Euler identity expansion for the probability amplitude (using Feynman's path integral approach) of a physically possible path between two points

$$\Phi = e^{iS/\hbar} = \cos(S/\hbar) + i \sin(S/\hbar) \quad (3-4)$$

where S is the action between the points, $\hbar$ is Planck's constant divided by 2π , and Φ is the probability amplitude for the path.

The path integral approach is based on concepts from the calculus of variations. In particular, it makes use of the concept of a functional (class of functions), whereby each function represents the trajectory of a possible path. When the first derivative of the functional goes to zero, the corresponding path is called an "extremal," and this corresponds to the shortest path. As Feynman points out, developing a remark by Dirac (1930/1981, p. 129), it is only paths close to extremals (i. e., paths close to the classical paths) which contribute significantly to the probability amplitudes. As Feynman (1964, Vol. II, p. 19-9) states "all the paths that give wildly different phases don't add up to anything." This is because contributions cancel out from paths which are significantly different from the classical ones all of which are, at least roughly, in phase (Feynman and Hibbs, 1965, p. 29; Feynman, 1985, pp. 53, 54).

I wish to construe probability amplitudes realistically in terms of properties of the foregoing rotational waves Ω_1 and Ω_2 . I use spherical polar coordinates r, θ

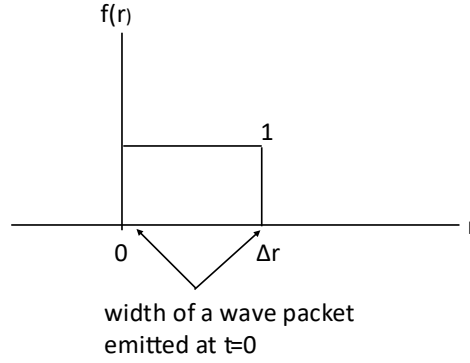


Figure 3.5 Step Function for a Wave Packet with width Δr

and Φ which are respectively the radius, polar and azimuthal angles for a thin shell. A step function $f(r) = u(r - \Delta r)$ on r can then be defined where the energy density u per subtended solid angle θ , Φ of the wave packet with respect to its originating particle is normalized to 1. Figure 3.5 illustrates such a step function for a wave packet pulse emitted at $t=0$. This can be generalized for a wave packet between radii r_1 and r_2 by a second step function $u(r) = f(r - r_1)$ where $r_2 = r_1 + \Delta r$. The probability amplitudes, Φ_1 and Φ_2 can now be defined as follows:

$$\Phi_1 = \frac{f(r-ct)}{\sqrt{4\pi\Delta r r}} \cos(\omega t) \quad (3-5)$$

$$\Phi_2 = \frac{f(r-ct)}{\sqrt{4\pi\Delta r r}} \sin(\omega t) \quad (3-6)$$

where r is the distance from the source particle for a wave packet of width Δr emitted between times t_1 and t_2 traveling at c over a finite distance to a potential absorber. The width of the wave packet then is $\Delta r = c(t_1 - t_2)$ and thus the difference between the inner and outer radii $r_1 - r_2$ of the wave packet remains a constant as the packet propagates at the velocity c . I interpret the amplitudes physically as

corresponding to the magnitudes of the $\mathbf{E}$ and $\mathbf{B}$ force fields. The probability P for a photon being absorbed is traditionally given by multiplying a probability amplitude by its complex conjugate. In my notation this corresponds to summing the squares of Φ_1 and Φ_2 . Thus, the probability density for absorption, as normalized, is given by:

$$P = \Psi\Psi^* = \Phi_1^2 + \Phi_2^2 = \frac{f^2(r-ct)}{4\pi\Delta r r^2} = \frac{f(r-ct)}{4\pi\Delta r r^2} [cm^{-3}] \quad (3-7)$$

It can be pointed out that inasmuch as f is normalized to 1, $f^2=f$ and $\sin^2 + \cos^2 = 1$.

While the magnitudes of the probability amplitudes associated with physically-possible paths of light rays vary at an inverse ratio with respect to their distances from their source particles, the probability for absorption is inversely proportional to the square of that distance. This corresponds to the energy density (intensity) of the electromagnetic field. This is also linked with the fact that the time-dependent Schrödinger equation is a first order differential equation with respect to time, unlike the standard wave equation which is second order with respect to time. As Messiah (1958/1999, Ch. 2) emphasizes this condition is required so that when the ψ function is multiplied by its complex conjugate ψ^* it results in a probability. It is thus an artifact of using complex numbers in quantum mechanics, and I question the necessity of this in Appendix A. I will now deal with the multiple path case.

The multiple path case involves interference effects from among the contributions from the different physically possible paths. I wish to explain interference effects in terms of the claim that there is a superposition of rotational effects from among the previously-mentioned rotational waves when they meet a potential absorber. Since potential absorbers will impact each of the subspaces of the different rotational waves, there will thus be a superposition of their various effects. The probability for absorption then is given by the absolute square of the sum (the "kernel" as defined by Feynman (1965, p. 26)) of the probability amplitudes associated with individual physically-possible paths. Phase factors of these probability amplitudes account for constructive or destructive interference

among the different paths. Notice that all of the fields contributing to the interference effects have a “presence” at the location where the absorption effects either occur or do not occur. It is just the case that, analogously to the cases of the effects of electric fields cancelling out discussed in Chapter Two, this presence is only actually “detected” when the effects do not cancel out or complete destructive interference does not occur.

Since I am using sine and cosine notation, kernels in my interpretation of Feynman’s account will be comprised of two parts K_1 and K_2 , corresponding to the summations of the respective probability amplitudes Φ_1 and Φ_2 . K_1 and K_2 will thus correspond to the resultant rotational effects, taking all of the rotational waves together respectively of the waves for the **E** fields and the **B** fields. The sum of these two kernels will then determine the probability of absorption of the individual wave packets; e. g., the probability P for light to travel between two points a and b would be given by adding the squares of the two kernels:

$$P = \Psi\Psi^* = K_1^2 + K_2^2 = (\sum \Phi_1)^2 + (\sum \Phi_2)^2 \quad (3-8)$$

The summations are over all physically possible paths from a to b , and Φ_1 and Φ_2 are the probability amplitudes, as previously characterized, associated with wave packets for each physically-possible path from a to b when these have been suitably normalized. The overall probability for absorption thus corresponds to the intensity at a given location of a superposition of the electromagnetic fields from the various source particles. Feynman (1965, Ch. 4) has shown that the resulting differential equation is the Schrödinger equation, although I will not summarize his derivation in this chapter. It need not be the numerically same photon as that which is emitted from one source which is absorbed, but that rather a discrete amount of energy is drawn from a "pool" to which many different sources may contribute (see Harry Paul, 1986.) To elaborate on this idea a bit more, I hold that each physically possible source of light of the same wavelength which impinges on a potential absorber contributes to the pool. Within a closed system the total photon number from the various emitting particles will remain integral (although the total may

change in integral amounts in the just mentioned emission and absorption processes) even though the photon numbers for the parts, after elastic scattering, are not in integral amounts. A nice illustration of this involves the principle behind how a laser works, whereby coherent (in phase) light from multiple sources is absorbed in discrete amounts. I now turn to my discussion of elastic scattering along with some more remarks about the Renninger effect.

In the case where light is not absorbed, I hold that two processes occur. First, elastic scattering will occur, where I hold that only a partial collapse of the wave packet takes place, with photons being literally divided into distinct portions in the process. These distinct partial photons will subsequently be propagated in different subspaces of spherical rotational waves, each possessing the same frequency as the original rotational wave and centered at the location of scattering. The second process involves light which is not scattered being "pushed aside" (the Renninger effect discussed earlier in this section) creating a shadow in the given direction and thus increasing the magnitude of the presence (and hence also the chances of absorption) from other locations. The change in the magnitude of the probabilities for absorption and scattering from other directions in this case are given by the ratios of the solid angles of shadow regions and the complete solid angle for a sphere - 4π sr (square radians). I will not derive the relative ratios (i. e., the cross sections) for scattering and absorption.

It can next be noted that a beamsplitter involves both the transmission and the reflection of light, and thus splits a light beam in two. According to standard quantum mechanics, a beamsplitter splits a probability wave but not a particle. Since I am identifying waves with particles though, the particle must also be split at the beamsplitter. It should be emphasized that I do not hold that probability waves possess an independent existence apart from the oscillating electromagnetic fields constituting the wave packets. It can also be noted that beamsplitters are key optical components in classical interference experiments, where they are needed to separate beams before they are recombined, with a phase differential, at a detector.

They are also key components with polarization entanglement experiments, which, after a brief digression on the absorption process, I will discuss in my next section.

I hold that in the absorption process a discrete amount of energy ($E=h\nu$) is drawn from the distinct subspaces of each wave-particle (e. g., a partial photon) impacting on the absorbing particle. The relative ratios drawn from each subspace correspond to the partial photon densities at the location. The energy is drawn from along past trajectories until a node, involving a four-dimensional particle, is reached. The node plays the role of providing a four-dimensional "link" between three-dimensional subspaces. The energy is then drawn from along other possible trajectories in subspaces centered on the node to other locations where the partial photon already has a "presence." The sense of "presence" here is the same as the sense in which the absorbed photon had a presence at its detector; i. e., it had a potential to be absorbed there. I hold that this backwards and forward (in a spatial not temporal sense) wave process must take place in the present; i. e., during the absorption time. Thus, both waves must travel faster than c , which requires a special reference frame. The concept of a backwards wave here is analogous to the concept of an advanced wave developed by David Klyshko (1988), only under my conception of these waves, the waves do not go backwards in time and instead act instantly in the present.

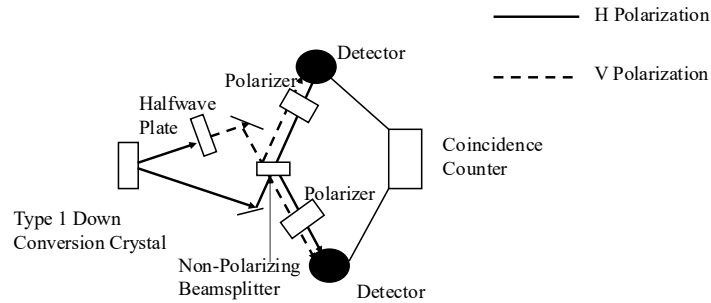
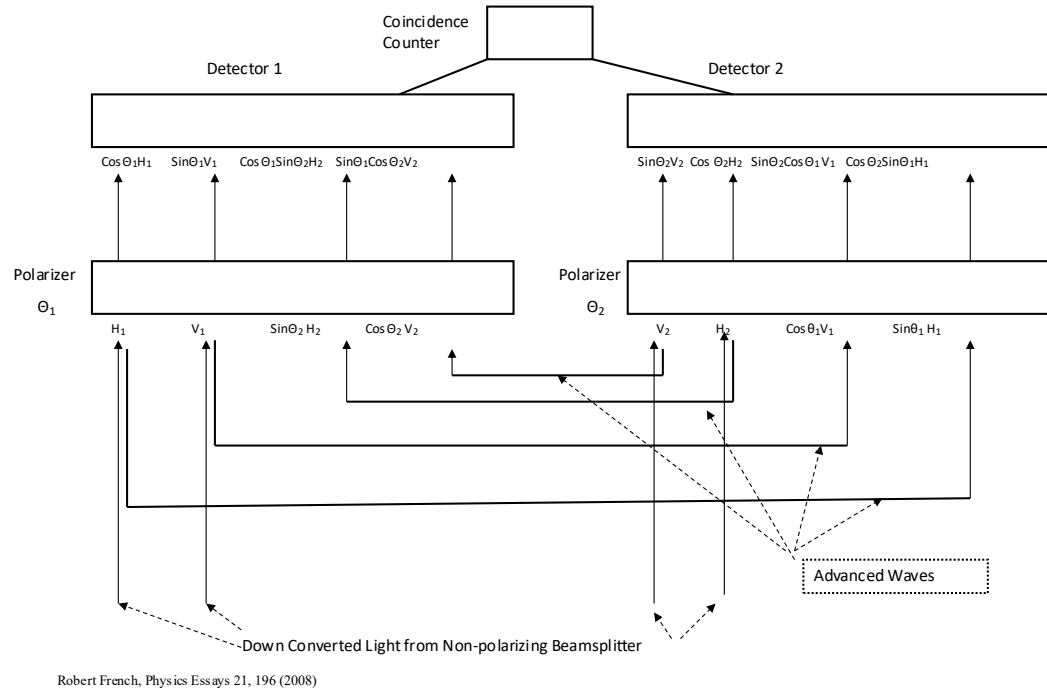


Figure 3.6 Shih and Alley (HOM) interferometer

3.3 Entanglement

I hold that the backward and forward wave process just sketched utilizing advanced waves acting in the present is the key for explaining the correlations at a distance which occur in polarization entanglement experiments. Thus, I now turn to a discussion of that topic. Polarization entanglement experiments utilize an interferometer illustrated in Figure 3.6 originally developed by Carroll Alley and Yanhua Shih, 1987. In the interferometer laser light is “down-converted” in a non-linear crystal so as to produce a pair of correlated photons (sometimes called the “biphoton”) one of which is made to be orthogonally polarized with respect to the other. After beams of the two orthogonally polarized photons are mixed in a beamsplitter beams and, using the language of standard quantum mechanics, “probability amplitudes” for both polarizations are made to overlap at each of two detectors with the results being sent to a coincidence counter.



Robert French, Physics Essays 21, 196 (2008)

Figure 3.7 Advanced waves for force fields (for probability amplitudes interpreted as rotational waves)

It can be noted that the corresponding force field strengths at each detector are changed respectively at a $\cos\Theta$ and $\sin\Theta$ rate by a polarizer placed at an angle Θ to the original bases angles. Since the intensity (energy) of a field is given by the square of the force field strength, energy density fields (from which photons, possessing a discrete packet of energy, are drawn) emerging from a polarizer are cut in intensity I in accordance with Malus' law $I \propto \cos^2\Theta$. It should be emphasized that, under my account, the original fields are not being cut in strength or filtered by the polarizer. Instead, as I mentioned in Section 3.1, since in effect I am defending an emission theory of light, I hold that new fields (driven by the old ones) are being created by the polarizer. A new basis of polarization is then given by the relevant Jones operator (see Robert Clark Jones, 1941) of the polarizer.

It is now possible to identify the field components from which the energy is drawn from in both singles counting and coincidence counting in polarization entanglement experiments. First there are energy density fields associated with the original signal and idler fields as they converge together at each individual detector. These energy density fields are given respectively by $\sin^2\theta_1 + \cos^2\theta_1$ and $\sin^2\theta_2 + \cos^2\theta_2$ terms. Since $\sin^2\theta + \cos^2\theta = 1$ the total singles counts from combining the two energy density fields will remain constant as a function of polarization angle.

I now turn to the utilization of my model to explain the non-local coincidence counting results of polarization entanglement experiments. This explication contains some subtlety and invokes two steps. The first step is illustrated in Figure 3.7. This involves the force fields that I explicated as rotational waves and which have a simultaneous presence at each of energy density fields will remain constant as a function of polarization angle. the two detectors. I then invoke advanced waves analogous to those proposed by Klyshko (1988) except as I previously noted I do not hold that they go backwards in time. For the purposes of this discussion I leave it as an open question as to whether these waves are generated by the polarizers or by the detectors themselves. For the $\Psi^\pm$ Bell states (Kwiat et al., 1995) I hold that advanced waves from each of the rotational waves at one detector are cut by the polarizers at the opposite detection system. This results in a $\sin\theta_1\cos\theta_2$ wave being present at one detector and a $\sin\theta_2\cos\theta_1$ wave being present at the other detector. The angle sum and difference identities $\sin(\theta_1 \pm \theta_2) = \sin\theta_1\cos\theta_2 \pm \sin\theta_2\cos\theta_1$ can then be invoked to show how this is equivalent to the sine of the sum or difference of the angles between the two polarizers. Similarly, with the $\Phi^\pm$ Bell states the rotational waves $\sin\theta_1\sin\theta_2$ are present at one detector and $\cos\theta_1\cos\theta_2$ at the other. The identity $\cos(\theta_1 \pm \theta_2) = \cos\theta_1\cos\theta_2 \mp \sin\theta_2\sin\theta_1$ can now be invoked on the rotational

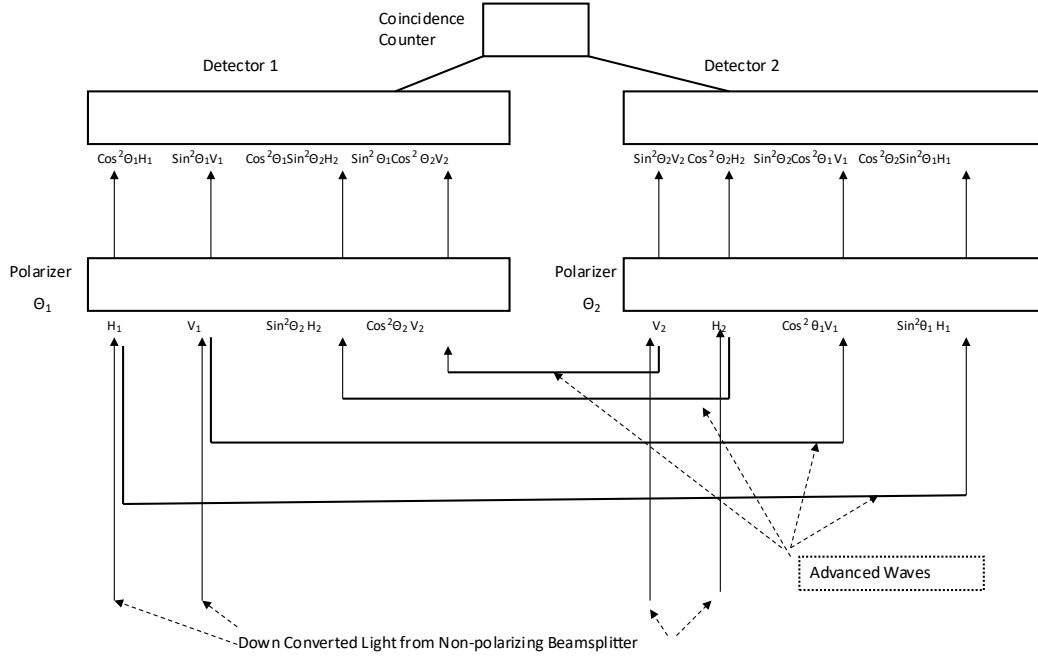


Figure 3.8 Advanced waves for energy fields (for probabilities).

waves at the two detectors. It can be noted that when the terms for any of the Bell states are squared this results in an interference term $\sin \Theta_1 \cos \Theta_1 \sin \Theta_2 \cos \Theta_2$ which cannot be factored (Shih et al., 1994).

The second step involves joint energy density field absorption as is illustrated in Figure 3.8. As shown, the absorption process also involves advanced waves, only in this case for energy density fields which are cut respectively at $\sin^2$ and $\cos^2$ rates in accordance with Malus' law. For the $\Psi^\pm$ Bell states this results in $\sin^2 \Theta_1 \cos^2 \Theta_2 H_2$ and $\cos^2 \Theta_1 \sin^2 \Theta_2 V_2$ energy density fields being present at one detector and $\cos^2 \Theta_2 \sin^2 \Theta_1 H_1$ and $\sin^2 \Theta_2 \cos^2 \Theta_1 V_1$ energy density fields being present at the other detector. In the absorption processes associated with coincidence counting energy is redistributed in a joint process so that the energy density fields associated with both sine terms are absorbed at one detector and the energy density fields associated with both cosine terms are absorbed at the other

detector. It should be pointed out that the energy for two photon absorption here is continuously present at the two detectors while the corresponding wave packets are present at them. It is the probability for the joint absorption process which is modulated by the first step process which includes the cross term $\sin\theta_1\cos\theta_1\sin\theta_2\cos\theta_2$. Similar remarks hold for the case of the $\Phi^\pm$ Bell states except in these cases the advanced waves are cut by polarizers on a new set of basis beams changed in polarization by 45° by a quarter wave plate.

By the preceding considerations, the energy of the “partial photons” present at each detector can be jointly absorbed in either of two alternative ways $\sin^2(\theta_1 \pm \theta_2)$ or $\cos^2(\theta_1 \pm \theta_2)$ corresponding respectively to the $\Psi^\pm$ and $\Phi^\pm$ Bell states. The joint energy associated with these states can be measured by the difference between the polarization vectors of the two polarizers with coincidence counting. It should be emphasized again that the photon absorption process at each separate detector draws energy from the sets of energy density fields jointly present at both detectors. It can also be noted that in the case of each of the Bell states, due to the Young inequality $ab \leq a^p/p + b^q/q$ where $p = q = 2$, the cross term $\sin\theta_1\cos\theta_1\sin\theta_2\cos\theta_2$ is always equal to or less than the squared terms $\sin^2\theta_1\cos^2\theta_2$ and $\sin^2\theta_2\cos^2\theta_1$ or $\cos^2\theta_1\cos^2\theta_2$ and $\sin^2\theta_2\sin^2\theta_1$ and thus negative energy is never involved. The preceding can be interpreted as the process of correlated two-photon absorption from the combined energy density fields present at the two detectors with correlated photon pairs from the fields being jointly absorbed by the process of correlated photon absorption (Hong-Bing Fei et al., 2010). It has also been argued that this process can occur with two absorbers at a macroscopic distance from each other (Ashok Muthukrishnan et al., 2004).

My claim is thus that the energy density fields associated with the polarization entanglement experiments are absorbed in tandem over the spatially extended region encompassing the two detectors. As Maudlin (2002) emphasizes a special reference frame (e. g., that of the source particle) is required here. In particular, he points out that the correlations shown in polarization entanglement

experiments require both superluminal causal connections and superluminal information transfer, even though this information transfer cannot be used to send a signal in any conventional sense.

It can be noted now that since the electromagnetic field properties depend on the polarization angles of both polarizers, they can only be measured by coincidence counts from both detectors. It can also be noted that Alain Aspect's (1982) experiments have shown that a common cause explanation of the correlations does not work. In his experiments the set of polarizers being sent to is changed in flight by fast acousto-optic switches after the photons have left their source. Since there is a space-like separation between the absorption events at the detectors the correlations cannot be explained by any subluminal communication between the detection events. It can next be pointed out that at very low intensities of down-converted light (e. g., at the single photon level), there are anticorrelations, showing photon number squeezing, between detection events at two detectors after a beam has been separated by a beamsplitter (Philippe Grangier et al., 1986). Thus, as in the joint-absorption case just discussed, the energy for the photon being absorbed is drawn from along past and forward trajectories in such a manner as to provide a link with the second detector. Depending on the nature of the experiment involved the key node, as previously defined, for creating the link may be in a beamsplitter or even in a down-conversion crystal. I should emphasize again that the foregoing account parallels the account in terms of advanced waves given by Klyshko (1988) and also the transactional account of John Cramer (1986) only it does not involve backwards in time causation, which I find to be quite implausible.

It might be thought that a fast rotating mechanical chopper could be used to block the foregoing advanced waves. While such choppers are commercially available as fast as 100 kHz there is a problem with creating sharp edges with such systems. As an alternative Shih and Sanjit Karkamar designed an experiment for blocking advanced waves using a pockels cell triggered by a pulsed laser as

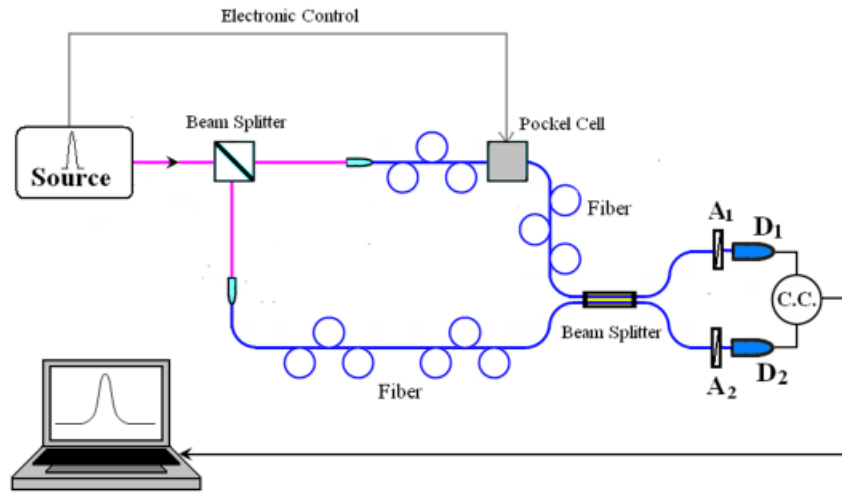


Figure 3.9 Design for testing for advanced waves with a down-conversion source

shown in Figure 3.9. When the pockels cell is temporally coordinated with pulses generating down-converted photon pairs so that it changes polarization once a pair has passed it, this should interfere with the advanced waves required for a joint absorption event, and this should result in different statistics for coincidence counting. Shih reports that the experiment was performed by a former graduate student of his at the University of Maryland Baltimore County in 2022 but with a null result. In particular, the standard correlation function was still observed during the time window when the advanced wave would have been “blocked” by the pockels cell. Since pockels cells work using a phase shift it might be questioned whether they actually succeed in blocking the advanced waves though. Thus, there may be something to be said for trying to use a mechanical shutter system with slower entangled particles, such as possibly atoms to also test the account.

It can also be pointed out that polarization entanglement experiments can be performed over a space-like separation of the absorption events at the two

detectors by feeding the light into optical fibers which are a few km in length (the record now for sending down converted light through optical fibers is over 100 km. Hubel et al., 2007). Given the distances involved with a satellite system based in space, where entanglement was demonstrated by Juan Yin et al. (2017) at some point it may also be possible to use a mechanical optical shutter system (such systems are commercially available at speeds as short as 5 ms) for blocking advanced waves even in the case of light.

Since my explanation is similar to the one just given for entanglement I will now make a few remarks concerning double slit experiments. As with entangled systems I hold that the probability for absorption is given by the local fields at a detector, only in this case with partial photons being emitted from both slits. In the resulting absorption process I hold that the energy of a complete photon is drawn from the linked trajectories. I will now close the chapter with a brief discussion of delayed-choice experiments.

3.4 Delayed Choice and Quantum Eraser Experiments

I believe that the just-given account of entanglement invoking advanced waves may also be the key for understanding what is going on with so-called "delayed-choice" and "quantum eraser" experiments. These experiments were originally discussed by John Wheeler (1978) who proposed a variety of methods (e. g., using a hinge to change a mirror or detector setting and by removing a pin in front of a hole which light might go through) in an experimental setup after a photon has already left its source. He then predicted that the resulting interference or "*welcher weg*" (which way determination) effects would be resolved by the result of the final setup at the time of detection and not the initial setup at the time of emission. These experimental predictions of Wheeler have been confirmed a number of times; e. g., by Yoon-Ho Kim et al. (2000) and by Vincent Jacques et al. (2007).

Note that it is the detector (absorption) process and not the source which determines either the nature of interference effects or their lack under my account.

This agrees with both Wheeler's predictions and the results of the experimental tests. Under my account even if a choice concerning the experimental setup is made after a photon leaves its source, it is still the final experimental setup at the time of detection which determines the outcome. For example, in the Jaques et al. (2007) experiment a fast optical switch is used to determine which of two interferometers (which are respectively in open and closed configurations) is being sent to.

An added point is that, as I see it, it is not the last-nanosecond process of "choosing" (however this is fleshed out, whether in terms of mental or even random processes) how to set up an experimental setup which causally determines whether interference occurs or not. Instead, as I see it, the key factor in determining whether interference patterns will be detected or not is the absorbing processes taking place at the detector(s) together with the processes of emission and absorption of advanced waves discussed in Section 3.2. In particular, I hold that these are causally responsible for whether interference effects or *welcher weg* information is being demonstrated. It can also be pointed out that what actually occurs is either an entangled interference pattern between two detectors (when *welcher weg* information is not present) or the lack of this pattern between two detectors (when *welcher weg* information is present). It should be emphasized that the changes in the experimental setups between these two cases occur before absorption takes place in the detectors.

A set of experiments which are closely-related to delayed choice experiments are the "quantum eraser" experiments originally proposed by Scully and Drühl (1982). These are thus also relevant for this discussion. It is noteworthy that these experiments – e. g., those by Herzog et al. (1995) and by Peng et al. (2014) always involve either adding or taking away path-length compensators or otherwise changing the nature of possible pathways (e. g. by the insertion of a quarter wave plate to change the polarization of a beam) to a set of detectors. Thus, as with the pure delayed-choice experiments, it is the probability amplitudes located

at the detectors (which I flesh out physically in terms of the final configuration of states of partial photons) which causally determine the probabilities of outcomes.

Chapter Four

Atomic Physics

In this chapter I extend my model to the atomic realm. In particular, I develop a highly-speculative variant on Bohr's model of the atom, with significant differences regarding certain details. It should be emphasized that, in at least partial justification of this procedure, as Alfred North Whitehead (1925, p. 106) points out in *Science and the Modern World*, the laws of nature may be quite different in strikingly different environments, and thus for example within atoms the basic laws may be fundamentally different. Also, as Ruggero Santilli (1981) has argued, both special relativity and quantum mechanics may not apply to the nuclear level where there is independent evidence for a spatial structure. Santilli claims that this also goes against the quark model of nucleons which is based on a point particle conception.

There is very strong empirical evidence from scattering and images from electron microscopes, for the existence of entities, atoms, of a width of approximately one angstrom - $10^{-10} m$. However, as far as an internal structure of atoms goes, the evidence for the nature for any particular postulated structure is much more indirect and nebulous. This is because it is not possible to perform internal probes for that structure without disturbing the structure in the process. For example, due to Einstein's formula explaining the photoelectric effect, $E=h\nu$, the energy of a photon with a wavelength short enough to probe the internal atomic structure will be in the x ray range of the electromagnetic spectrum, and thus will be too energetic to serve as such a probe. This lack of direct evidence of the nature of any internal structure leaves room for speculation concerning that structure such as that which I conduct in this chapter.

As should be clear from my discussion in Chapter Three, one point of difference which I have with respect to Bohr's model of the atom is that I reject the dipole model of radiation at least for individual atomic sources. In fact, I hold that

the dipole model is responsible both for the *ad hoc* character of Planck's avoidance of the so-called "ultraviolet catastrophe" associated with the Raleigh Jeans law of blackbody radiation and also the *ad hoc* avoidance of the result of classical electromagnetism that an electron in an orbit radiates due to its acceleration. Also, instead of claiming that electrons are like finitely-sized miniature planets in an orbit (or even worse anything like Boscovich's point particles) I claim that they are spread out over entire orbitals. It might be noted that this move of claiming that electrons are spread out over whole orbitals was to some extent anticipated by Arthur Eddington (1928, Ch. 9) in *The Nature of the Physical World*. I should also emphasize that I only deal with circular and not elliptical orbitals. Thus, I do not account for the fine structure of spectra as being due to changes of energy when adding components from changes in orbital radii. However, I do not give an alternative account for this fine structure in the chapter either.

I should note that I am not the first to claim that electrons possess a finite size, since Hendrik Lorentz in *The Theory of Electrons* (1916) develops a theory claiming this as do David Bergman and Paul Wesley (1990). As I pointed out in Chapter Two, this does not entail that electrons will be unstable due to the electrical repulsion among their parts, since under my account the thin shell structure only is in place outside of charged particles and not within them. This is consonant with Whitehead's point.

As far as I know, others have not previously postulated that electrons are spread out over whole orbitals. In any event, I agree with Bohr both in being a realist about the existence of electrons independently of observation (unlike the modern construal, the so-called "Copenhagen interpretation," associated with the later Bohr) and in the angular momenta of electrons being quantized. I also agree with Bohr in explaining absorption and emission spectra of atoms in terms of literal switches between electron "orbitals" when these are suitably construed. I now elaborate on some considerations for my differences with the Bohr model.

With respect to the issue of distinguishing between an orbital and a bound electron occupying it, one consideration is that of Ernest Rutherford's scattering experiments with alpha particles - helium nuclei. Since these particles easily transverse an "electron cloud" before being scattered by a much smaller nucleus it follows that either the cloud is easily penetrated by at least certain particles, or is comprised of much smaller components. However, since I hold that electrons are spread out over entire orbitals only the former conjecture is relevant. I will not speculate in this chapter on how this could be the case though. Also, instead of appealing to a so-called "strong nuclear force" to hold nucleons together, I instead appeal to the presence of the initial ground state electron orbital to prevent the nucleons from separating. Since electric forces are not present within the orbital itself, it follows that there is no electrical repulsive force for a nuclear force to overcome. It might also be noted in this regard that, as Feynman (1964, Vol. 2, p. 2) points out, an atomic bomb actually works from the electrical repulsion among protons in a uranium nucleus when it is tapped lightly by a slow neutron and not from the influence of a separate nuclear force per se. Admittedly, this process would not also account for the nuclear fusion which occurs in the hydrogen bomb since that works by combining nucleons and not separating them.

I agree with Bohr (and also Schrödinger) that the radii of successive orbitals are a quadratic function of the principal quantum number n . However, I both disagree with the rationale which Bohr gives and also add a subtlety concerning unoccupied "shells" existing at linear intervals. As is well known Bohr (1913/1967, p. 136) gave a theoretical consideration for the quadratic dependence involving the claim that a Coulomb electrical central force field holds electrons in their orbitals which results in the centripetal acceleration mv^2/r . However, as was also known at the time, the resulting orbitals (at least if classically interpreted) are unstable, since accelerated charges radiate energy. This is the primary failure of the point particle orbital model and there does not seem to be a convincing non-*ad hoc* way around the problem. Hence, it is appealing to suppose that electrons are not orbiting point

particles. It is interesting that Bohr (1913/1967, p. 141) also cites one piece of empirical evidence here, that being that in a high vacuum more spectral lines of the Balmer series for hydrogen are apparent. To the best of my knowledge though the subject of the magnitude of the orbital radii of excited states in different degrees of a vacuum has never been systematically investigated from terrestrial sources, and it is not possible in the case of the spectra of stars to independently quantitatively establish the degree of the vacuum. This raises questions concerning how much precision on these matters there actually is empirical evidence to support and thus more testing would appear to be called for with respect to these issues.

After dealing with the foregoing issues concerning the nature of electron orbitals, I discuss the nature of the spin of the electron and use that treatment in a discussion of the nature of covalent chemical bonds under the model. I conclude the chapter with some remarks on the subject of what "reduces" wave packets; i. e. what makes so-called "measurements" have definite values.

4.1 Electron Orbitals

The following is a speculative account of the structure of electron orbitals. It is based on the well-known fact that there are four degrees of freedom in order to adequately account for the known facts from spectral analysis experiments concerning the orbital shells of electrons in an atom. These degrees of freedom are characterized in terms of four integral or half-integral quantum numbers. Only the principal quantum number is needed to account for the Balmer formula for the spectrum of hydrogen. An additional two quantum numbers (the azimuthal quantum number l for orbital angular momentum and the total angular momentum quantum number j) are needed to account for the effects of a magnetic field with the normal Zeeman effect and a fourth, the spin quantum number s , to account for the anomalous Zeeman effect and the Pauli exclusion principle. I now turn to the details of the account.

I identify electron orbitals with the same set of thin shells which I introduced in Chapter One in my discussion of the electromagnetic field. As I conceive of

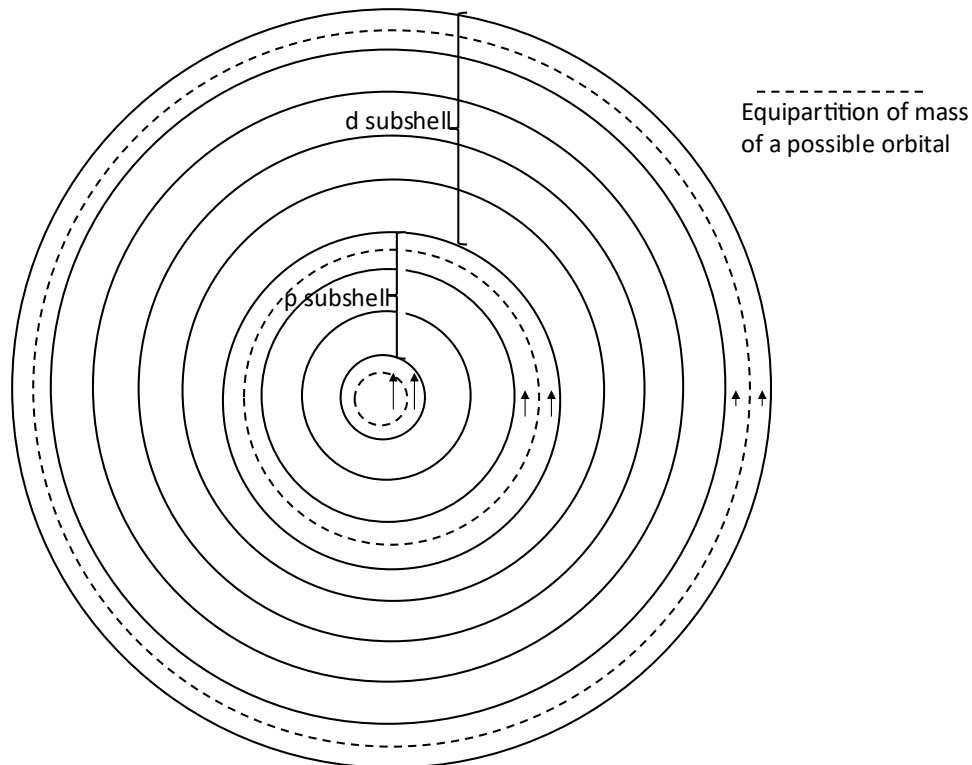


Figure 4.1 Possible electron orbitals with radii proportional to n^2 together with integrally-spaced shells. The arrows represent that the electrons possess equal rectilinear velocities in their orbitals.

them, these shells may or may not be either actually occupied or even potentially occupied by an electron just as quantum states may or may not be occupied. Hence the shells are identified with the principle quantum numbers. The shells will be integrally spaced but at least for the case of hydrogen can only be occupied at quadratically-spaced intervals for higher energy levels; i. e., when $r \propto n^2$ as is illustrated in Figure 4.1. Notice that the first orbital is adjacent to the nucleus. As I argued in Section 1.4, this claim is consistent with the fact that in Figure 1.3 of that section adjacent shells are shown to be in orthogonal sets of dimensions since the ideal liquids come to be out of Figure 4.1 synchronization through repeated oscillations. I also postulate that the integral spacing of electron orbitals is consistent with the claim that subsequent orbitals are not synchronized in sets of

dimensions since I hold that the creation and evolution of the thin shells temporally precedes the creation of the orbitals.

I now present two brief speculations as to what may possibly be responsible respectively for the integral spacing of possible electron orbitals (in a sense which I will be elaborating on shortly for orbitals which can be occupied for elements with atomic numbers greater than one) and the quadratic spacing of orbitals for higher energy levels than for the ground state energy. With respect to the integral spacing of possible orbitals, my hypothesis is that in the process of initial hydrogen atom formation a proton and a spherical “electron” are created. The spherical “electron” is postulated to possess a diameter the same as the width of the ground state orbital and to rotate in different directions so as to osculate against the outer surface of the innermost thin shell. This would result in a new shell with the same width as a previous one. It may be thought to be somewhat tempting to identify individual bound electrons with these rotating spheres, possibly sometimes rotating in opposite directions in an inner shell. However, it must be remembered that the shells need not be occupied by actual electrons, and I will not attempt to develop an alternative concept of potential electrons in this chapter.

I now turn to the issue of motivating the quadratic spacing of orbitals for higher-energy electron orbitals. In particular I show how Figure 4.1 suggests a model for the structure of electron orbitals for atoms with atomic numbers greater than 1. According to the periodic table of elements the total number of electrons in the n^{th} shell of an atom is $2n^2$ where n is the principal quantum number. Also, for the series of electron subshells l in each shell up to the level n^{th} shell is given by $2(2l + 1)$ where l is the azimuthal quantum number. Using standard chemistry notation $l = 0$ is the s level, $l=1$ the p level, $l = 2$ the d level and $l = 3$ the f level. Thus, there are 2 electrons in an s subshell, 6 electrons in a p subshell, 10 electrons in a d subshell, and 14 electrons in an f subshell. Under my model these sublevels correspond to the regions in the diagram between occupied excited levels of hydrogen (including the excited levels themselves), with the claim being that these

regions are successively filled in at integral intervals. Thus, as is illustrated in Figure 4.1, there are respectively 3 and 5 subshells for the first two intervals. The next interval, corresponding to the f level with 7 subshells, is not shown. Since two electrons (with opposite spins) occupy each orbital this structure is in accordance with the arrangement of the periodic table.

Obviously this spacing of orbitals differs from the spacing for the thin shells when they are not occupied by electrons though, inasmuch as I showed in Chapter One that that spacing varies at an inverse square rate with respect to the distance from a source particle. I should also note that as far as I can see it is also possible to have variants of the model where the spread-out electrons occupy whole orbitals, but I will not discuss these variants further in this chapter and obviously an alternative account would then have to be given of the subshell structure. From the foregoing considerations, it follows that the difference in volumes V_r between an electron orbital and the preceding subshells is given by

$$V_r = \frac{4}{3} \pi r_n^3 - \frac{4}{3} \pi (r_n - 1)^3 \quad (4-1)$$

where n is the principal quantum number (the ordinal number of the orbital) and where r_n is the radius of the n^{th} orbital.

Since I hold that electrons are ideal liquids, I hold that the rectilinear velocities of the circular streamline contours for electrons in their orbitals are a constant as a function of changes in the orbitals' radii. It can be pointed out next that the expression $L = \sqrt{l(l+1)}\hbar$ for the relationship between the azimuthal quantum numbers l and the magnitude L of the orbital angular momentum vector $\mathbf{L}$ approximates $(l + \frac{1}{2})\hbar$ as l increases. One possible construal of the expression $L = \sqrt{l(l+1)}\hbar$ then would be to have the radius of the orbital be proportional to $n^2 - \frac{1}{4}$ and the rectilinear velocity of the orbital be proportional to $1/(n - \frac{1}{2})$ since

$$\frac{n^2 - \frac{1}{4}}{n - \frac{1}{2}} = n + \frac{1}{2} \quad (4-2)$$

Obviously other construals which do not come out with such simple ratios are also possible. Furthermore, as I previously pointed out, it is questionable how much precision there is in the empirical evidence for the radii of excited states of atoms, so it is not at all clear that this construal is ruled out.

One possible suggestion for a variant on the foregoing account (with a different radius) is to have it correspond to the radius of a spherical surface for the equipartition of mass distribution in an orbital (which $n - \frac{1}{2}$ approximates as n increases). Since the spherical surface constituting the equipartition of mass distribution in an orbital will also have equal volumes on either side, and since the width of each ring is 1 (due to the integral spacing of their radii), this results in equipartition volumes V_{em} of successive orbitals being given by

$$V_{em} = \frac{4\pi}{3} \left[\frac{r_n^3 - (r_n - 1)^3}{2} \right] \quad (4-3)$$

where r_n is the outer radius of the n^{th} possible orbital. The resulting radius of the n^{th} equipartition surface is then given by

$$r_{em} = \sqrt[3]{r_n^3 - \frac{(3r_n^2 - 3r_n + 1)}{2}} \quad (4-4)$$

The resulting rectilinear velocity $\mathbf{v}$ for this orbital would be

$$\mathbf{v} = \frac{(n + \frac{1}{2})\hbar}{r_{em}m} (\hat{\mathbf{r}} \times \mathbf{a}) \quad (4-5)$$

where m is the mass of the electron and $\hat{\mathbf{r}}$ and $\mathbf{a}$ are unit vectors respectively in the direction of the equator where the rectilinear velocity is a maximum and in the axis of rotation direction.

It can be pointed out that under the foregoing model the three-dimensional volumes of electron orbitals asymptotically approximate being proportional to r^2 with increasing r values. There may well be other possible interpretations here as well which I will not speculate on in this chapter. In any event I now turn to a

discussion of how a superposition of electron states is created during the period in between the absorption and the emission of a photon by an electron.

When a photon is absorbed an electron switches orbitals (i. e., changes shells) from a lower orbital (e. g. the ground state) to a higher orbital (an excited state), and when an electron transitions from a higher orbital to a lower one a photon will be emitted. I claim that when a photon is absorbed, the absorbing electron acts like an ideal liquid in that it is temporally "split in half," with one "half" remaining in the original orbital, and the other half "jumping" to an orbital with an energy higher by a factor $h\nu$, the energy of the photon. In effect this constitutes a superposition of the two states inasmuch as each orbital is occupied, albeit each with only half of the electron. The time for a realignment of the electron halves is determined by the period T (i. e., the reciprocal of the frequency, or $1/\nu$) of the light ray being absorbed. This fits in well with the old Bohr quantum theory, where Bohr (1913/1967, p. 136) also equated the frequency associated with an emitted photon with the difference in frequencies of revolution of electrons in the orbitals being jumped between. I will not speculate on the nature of any mechanisms which might be causally responsible for determining the frequency of the emitted and absorbed light here other than to suggest that they may involve the radial angle of a portion of the spread-out electron occupying one orbital "catching up" and thus realigning with the radial angle of a corresponding portion of the spread-out electron occupying the other orbital.

I now turn to the details of my derivation of the Balmer formula for the spectrum of hydrogen:

$$\frac{1}{\lambda} = R \left(\frac{1}{n_1^2} - \frac{1}{n_2^2} \right) \quad (4-6)$$

where n_1 and n_2 are respectively the principal quantum numbers of the lower and upper quantum states and $n_2 > n_1$. R is the Rydberg constant. The circumference c_1 of the original orbital is directly proportional to the original radius r_1 ; i. e., $c_1 = 2\pi r_1$. Similarly, the circumference of the orbital jumped to will be directly proportional to the radius of that orbital; i. e., $c_2 = 2\pi r_2$. The angular speeds Ω_1 and

Ω_2 of the electron orbitals can now be derived. In order to have the angular momentum be an integral function of n , these angular speeds are proportional to $1/r_n \propto 1/n^2$. Thus the difference in angular speeds (the time for the radial angle of one orbital to realign with the radial angle of the other orbital) is given by $1/n_1^2 - 1/n_2^2$. When this is multiplied by the Rydberg constant it gives the Balmer formula. I now turn to my discussion of the energy of orbitals.

It can be recalled that I agree with Bohr that the angular momenta of bound electrons is quantized. Thus, at each subsequent orbital the angular momentum will be proportional to n and also be equal to $r\mathbf{v}$ where $\mathbf{v}$ is the rectilinear velocity of the electron in its orbital. Since $r \propto n^2$, the angular momentum is also proportional to n , and thus the rectilinear velocity is directly proportional to $1/n$. It can be recalled that I hold that the rectilinear velocities of electrons are constants as a function of changes in radii of their circular contours within a particular electron orbital. Thus, since the kinetic energy is proportional to the square of the rectilinear velocity, it follows that the kinetic energies of the respective orbitals will be proportional to $1/r^2$. When this is multiplied by the Rydberg constant, and using the Planck formula $E=h\nu$ it gives the Balmer formula for the hydrogen spectrum. It can be noted that this makes the kinetic energy of the lower orbital less than that of the upper one. It can be pointed out though that bound states are held to possess negative energies both under the Bohr theory (Dirk Ter Haar, 1967, p. 36) and under wave mechanics (David Bohm, 1951/1989, p. 247). Thus, less kinetic energy is subtracted from the higher orbital and it will have less total potential energy. It should be emphasized that the concept of "negative energy" being used here need not be paradoxical since it just refers to negative potential energy; i. e., positive energy is required in order to dislodge (free) an electron from its orbital.

It is true that it is usually held that the Bohr theory has been superseded by the modern quantum theory of wave mechanics. Thus, obviously in order to be at all complete much more is needed here in accounting for the successes which modern quantum theory has had with such matters as predicting transition

probabilities together with the resulting intensities of spectrum lines and the spectrum of helium. I will not tackle any of these subjects in this chapter though.

4.2 Spin

It is well known that a fourth degree of freedom "spin" is necessary in order to account for such phenomena as the anomalous Zeeman effect (splitting of spectral lines in a magnetic field) and the Pauli exclusion principle. In fact the latter principle is being appealed to in Section 5.1 with the claim that each orbital contains two electrons of opposite spin. It is perhaps noteworthy, in regard to the relationship of spin to defenses of the Bohr model, that Von Neumann (1955/1983, p. 5) asserts that "almost all difficulties of the model disappear" when it is supplemented with Samuel Goudsmit and George Uhlenbeck's (1926) account of spin and the magnetic moment of the electron. In this section I will just deal with the case of spin $1/2$; i. e., for the case of fermions, and only for bound electrons in particular.

It can be recalled that in keeping with my avoidance of the ignorance interpretation, my model differs from traditional treatments of quantum mechanics in that it holds that electrons are not point particles, but that instead they are "spread out" over entire orbitals. It both follows that no distinction is made between the spatial location of an electron and its orbital, and also that the spin angular momentum cannot be sharply separated from the orbital angular momentum.

The foregoing considerations suggest an account of spin in terms of an alternative to the traditional account of spin in terms of a literal precession ("the Larmor precession") of the axis of rotation of the electron in the presence of a magnetic field where the spin is held to be responsible for the precession. Instead I hold that spin involves a property over the entire shell constituting an electron,

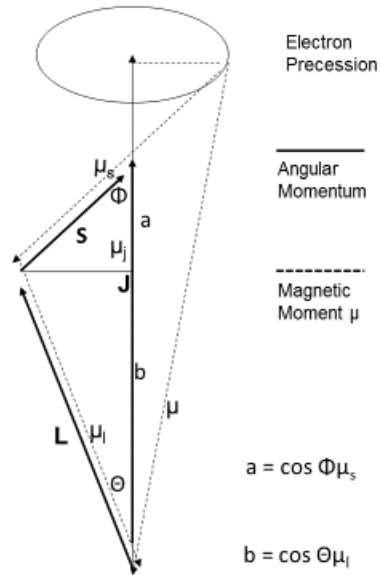


Figure 4.2 Precession of electron in the presence of a magnetic field with magnetic moment μ

namely the vorticity. Figure 4.2, adapted from Gerhard Herzberg (1937/1944, p. 109), illustrates the manner in which the quantizing the orbital angular momentum $\mathbf{L}$ and spin angular momentum $\mathbf{S}$ together with the total angular momentum; $\mathbf{J}$ results in an electron precession.

Orbital angular momentum $\mathbf{L}$ and spin $\mathbf{S}$ effects only occur in the presence of a magnetic field; the normal and anomalous Zeeman effects, although the spin effects cancel out for the normal Zeeman effect for atoms possessing an even-numbered number of electrons. This suggests that spin and orbital angular momentum are not intrinsic properties of electrons but rather, along with the total angular momentum $\mathbf{J}$, are activated as a coupled system by the presence of a magnetic field. Thus, unlike the traditional conception of a fixed-valued quantized “electron spin” interacting with the magnetic field to result in the splitting of spectral lines associated with the Zeeman effects, I instead hold that the magnetic

field creates the splitting. The result is context dependent, and thus has something in common with the contextualist account of spin postulated by John Bell (1987, Ch. 17) in the context of presenting a counterexample to “proofs” (such as those by John Von Neumann (1932/1955, Ch. 4, Sec. 2) and Simon Kochen and Ernst Speiser (1967) against at least local hidden variable theories. Bell’s example is discussed in some detail by Bohm and Hiley (1993, Ch. 10).

The presence of a magnetic field causes the magnitude of the split in spectral lines to be determined by two factors. As illustrated in Figure 4.3, one factor is the component of the magnetic moment of the electron in the field direction. Notice that this magnetic moment is parallel to the axis of rotation given by Equation 4.5 associated with the electron's orbital angular momentum. The second factor is Alfred Landé’s g factor which is defined as $1 + \frac{J(J+1) + S(S+1) - L(L+1)}{2J(J+1)}$ where j , l and s are respectively the total angular momentum, orbital angular momentum, and spin quantum numbers. The g factor in turn is a component of the formula for the Larmor frequency $-\frac{egB}{2m}$ where B is the magnitude of the magnetic field, m is the mass of the electron, and e is its charge

With respect to rectilinear velocities, my hypothesis is that the magnetic field causes a deviation of a spherical electron orbital away from that of a rigid sphere in the sense that there is relative motion among the internal parts. In particular, I claim that the rectilinear velocities (resulting in the orbital angular momentum) are a constant as a function of polar angle. The resulting difference in rectilinear velocity as a function of polar angle here causes a coupling effect which in turn results in a rotation (the vorticity) which I associate with spin. The situation is analogous to the one for the Poincaré sphere model for polarization and is

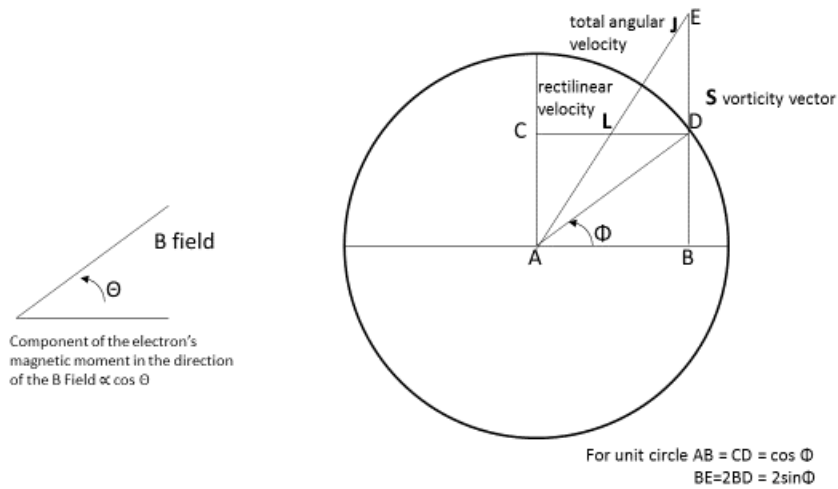


Figure 4.3 Account of spin in terms of vorticity

illustrated in Figure 4.4. Thus, as with the case of polarization, since the resulting angular velocity increases monotonically as a function of polar angle away from the equator, this will result in microscopic circulatory movements, or eddies, with opposite rotations on opposite sides of the equator. This is also akin to the effect of the so-called “Coriolis force” on the earth’s surface whereby water in drains in the northern hemisphere rotates counter-clockwise and in the southern hemisphere rotates clockwise.

Since I also analyzed magnetism in terms of vorticity in Chapter Two, I hold that in many respects (aside from the fact that its magnitude is quantized) spin is the analogue within an atom of an external magnetic field. It can be recalled from my discussion of magnetism in Chapter Two, vorticity Ω is a hydrodynamic concept corresponding to the circulation per unit area for an infinitesimal loop. As I pointed out in that chapter, being a curl $\nabla \times \mathbf{v}$ where $\mathbf{v}$ is a velocity field, the vorticity itself is a macroscopic property. However, as I also pointed out in Chapter

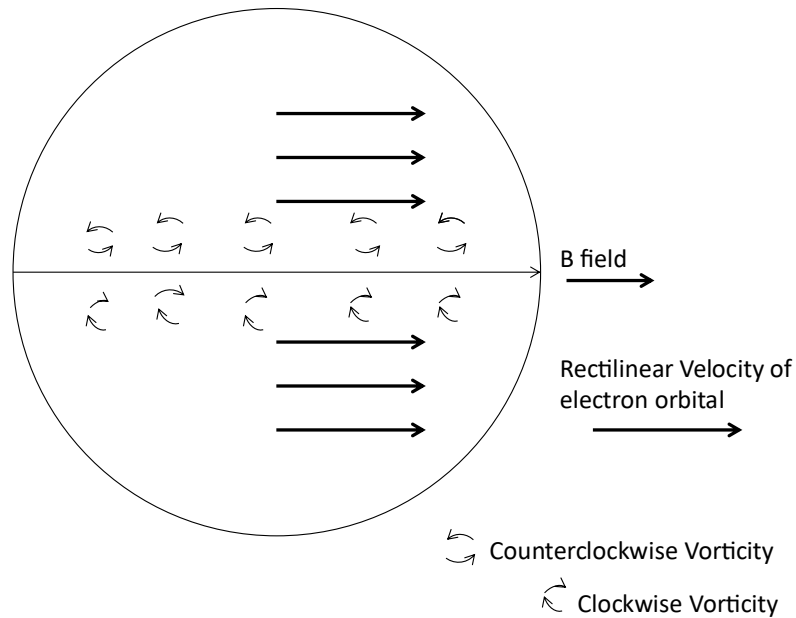


Figure 4.4 Spin Analogue of the Poincaré sphere for polarization

Two the regional subject matter of vorticity, the circulation per unit area in the limit of an indefinitely small loop, is a subject in the small. The velocity field corresponds to the magnetic field, and the couple creating the vorticity (spin) corresponds to the infinitesimal loop. Also, since vorticity is defined as the curl of the velocity it should be emphasized that the direction of the vorticity vector is perpendicular to that of the velocity field.

My next series of remarks involve Landé's g factor which was just defined. It is closely related to the gyromagnetic ratio (the ratio of the electron spin's angular momentum to the electron's magnetic moment). The traditional account of the g factor being two from Dirac's (1930/1981, p. 266) relativistic wave mechanics is not available here since it presupposes special relativity. Thus, an alternative account for this ratio being two is necessary. Feynman (1964, Vol. II, p. 40-5) points out that the vorticity of a fluid is twice the magnitude of the local angular

velocity. This would make the vorticity correspond to the magnetic moment of spin and thus the macroscopic local angular velocity would then be associated with the spin angular momentum vector $\mathbf{S}$. This would then agree with the observed g factor of two.

In view of the preceding considerations a possible suggestion for motivating the quantization of total angular momentum $\mathbf{J}$ can be sketched as follows. As illustrated in Figure 4.3, the orbital velocity $\mathbf{L}$ varies as function of the cosine of the polar angle Φ . I identify the angular velocity of the electron orbital with $\mathbf{J}$ quantized as a complete rotation. Thus, the ratio of $\mathbf{L}$ with respect to $\mathbf{J}$ will vary as an inverse of the cosine (i. e., the secant) as a function of Φ . It should be emphasized again that the spin angular momentum vector $\mathbf{S}$ (which I identify with the vorticity vector) is orthogonal to $\mathbf{L}$ inasmuch as the vorticity is the curl of the local velocity field. Since the g factor is 2 (due under my account to the vorticity being twice the local angular velocity), the total rotation vector $\mathbf{J}$ possesses a precession of the angular location of a complete rotation. It can then be recalled from my account of Section 4.1 that a complete rotation plays a key role in determining the nature of spectra. Thus, a precession in the location of a complete rotation results in a shift of these lines by changing the magnitude of the vorticity vector and thus resulting in a greater split in spectral lines depending upon the direction of the precession.

Isomorphisms can be noted between the diagram in Figure 4.3 and those for polarization in Figure 3.4 and for magnetism in Figure 2.6. As with the postulated linkage between polarization and magnetism discussed in Section 3.1 this suggests a linkage in mechanisms between spin and magnetism. In particular, a possible hypothesis would be, in spite of their possessing different dimensionality (four dimensions for electron orbitals and three dimensions for thin shells), to extend spin from electron orbitals to adjoining thin shells and linking it with magnetism (as was suggested by De Climont, 2014). As explicated in Section 1.4, the linkage would only be to the set of shells centered at the particle in question. The positive and negative poles of the magnetic field would then correspond to spin up and spin

down states – i. e., to clockwise and counterclockwise rotating vortices. As I pointed out in Section 2.2, macroscopic magnetic effects would then be due to the alignments of the spins of individual electrons of different atoms. Such an account, while intriguing, obviously requires more working out such as on the nature of the linkages between the electron orbitals and the thin shells.

I now address the issue as to why it is the case that electrons have opposite spins when paired together in an orbital. In the presence of a magnetic field this corresponds to the vorticities having opposite directions – clockwise and counterclockwise. Spin states can be thus be characterized as being “spin up” or as being “spin down” depending upon whether the rotation associated with the vorticity vector is clockwise or counterclockwise. It can be seen that, depending upon whether these rotations are clockwise or counterclockwise, they will either contribute positively or negatively to the velocity of the orbital, thus accounting for the splitting of spectral lines in a magnetic field with the anomalous Zeeman effect. It should be emphasized again that there are paradoxes associated with any non-contextualist traditional conception of a fixed-valued quantized “electron spin” interacting with the magnetic field to result in the splitting. Instead, I hold that the magnetic field creates a non-fixed-value vorticity (which the “spin” is identified with), which in turn creates the splitting.

For electron orbitals, according to Pauli’s exclusion principle two electrons share the orbital with opposite spins. I discuss two variants of the model here. In the first variant (which I term “spin model 1”) I hold that since the electrons come from different nucleons each of the electrons is located in a distinct parallel subspace – the subspace originally associated with that nucleon. The opposite spins of the paired electrons bound together in an orbital will then be, so to speak, “actualized in tandem” in the presence of a magnetic field. That is, I hold that they come jointly into existence in the process of being “measured” as with spectroscopic measurements of the splitting of spectral lines associated with the normal Zeeman effect which applies when there is an even number of electrons. I

do not construe what is going on epistemically, but instead in terms of energy either being added or subtracted (depending upon the direction of the vorticity) to the electron in the absence of a magnetic field by the presence of a magnetic field. Admittedly, it might seem then that, in the absence of a magnetic field, there would be no way to distinguish between the two spin states. However, once a magnetic field is present, the two directions of rotation are not determined in an *ad hoc* manner inasmuch as these directions are determined in terms of which side of the equator they are on as previously discussed.

In the case of electron orbitals with only one electron (i. e., for elements with an odd atomic number) it might be thought that the electron would be confined to just one hemisphere since that is the only way for it to spin in only one direction. However, this appears to be rather *ad hoc* since it arbitrarily confines the electron to just one hemisphere. Instead, I strongly suspect that it is not so much a fixed spin in one direction but rather a superposition of spinning partial electrons over the whole orbital; i. e., with a partial electron spinning in one hemisphere and a partial electron, with the opposite spin in the other hemisphere. Again, this is analogous to the model of light polarization illustrated in Figure 3.3. This also suggests my second variant of an electron pair (which I term "spin model 2"). In this variant both electrons comprising an electron pair exist in a single subspace. This would have to involve electron from separate nucleons in some manner "collapsing" into a single subspace. I will not develop this model further in this chapter except to allude to it in my treatment of covalent bonds in Section 4.3.

Variants on the experiments by Otto Stern and Walter Gerlach (1922) can also be cited as confirming at least some of the foregoing points. These experiments are akin to those of optical experiments mentioned in Chapter Three using horizontal and vertical polarizers to completely block light, and where light re-emerges when a third polarizer with an intermediary angle (e. g., 45°) is inserted in between the two other polarizers. Thus, I do not hold that the magnets of a Stern Gerlach device do anything like filtering for a given spin. Instead, I hold that in

effect they change the angle of rotation of an electron orbital, and hence also the angle of the equator or the orbital. As I speculated in the case of light, it may possibly also be the case that in effect numerically new electrons are created in this process. I will not speculate any further on these matters in this chapter though.

I now turn to a discussion of the nature of chemical bonds under the model. It turns out that my analysis of spin plays a key role in that discussion.

4.3 The Chemical Bond

I confine my discussion of chemical bonds to the case of covalent bonds. What came to be known as "covalent bonds" were originally hypothesized by Gilbert Lewis (1916) when he claimed that "an electron may form part of the shell of two different atoms and cannot be said to belong to either one exclusively." Linus Pauling (1939/1960, p. 5) characterizes Lewis's position as postulating a "single bond" that "involves two electrons held in common by two atoms." Pauling (1939/1960, p. 61) adds the requirement that the paired electrons be of opposite spins. The question arises though as to where exactly these electrons are located. Are they each attached to each atom in the sense that each one is literally in the orbitals of each atom? Or are they attached separately to each atom? Neither alternative is very attractive and may not even be coherent at least for physically realist interpretations as I now show.

Consider the following dilemma which can be pressed with respect to realist construals of Lewis's definition if electrons are conceived in accordance with traditional quantum mechanics as being indivisible point particles. If the electrons are conceived as being each in the locations of each atom then they would have to be spread out which goes against the point particle conception. If they are not conceived as each being in the locations of both of these atoms, then it is not at all clear how this differs from the case where there is no bond at all. It is true that the concept of a superposition is often invoked here whereby it is claimed that each physically possible state exists concurrently at the same time until a measurement is performed. However, for a physically realist account (and if it is also claimed

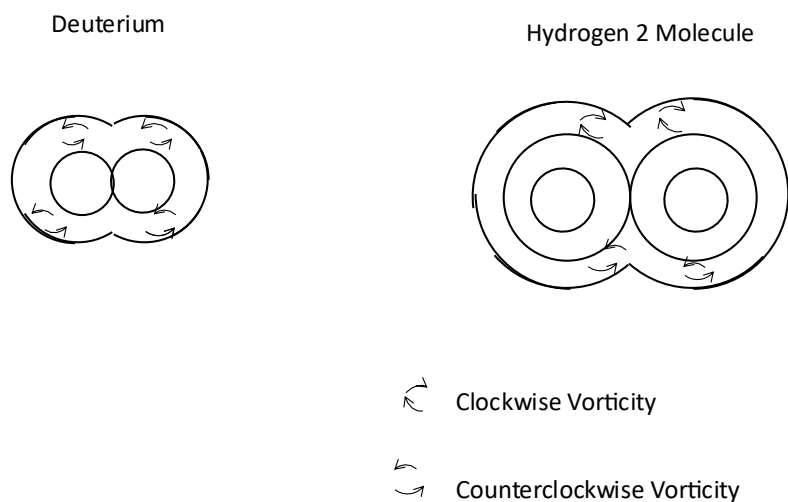


Figure 4.4 Structure of covalent bonds ahead, I explore some of these issues in more detail in Section 5.4 on the reduction of wave packets.

that electrons are indivisible), then the foregoing dilemma still applies.

Looking My solution to the foregoing dilemma is to claim that both electrons are in each atom's orbitals (in separate parallel subspaces for spin model 1 and in the same one for spin model 2) since they are both located in the innermost shell (which I numerically identify with the "shared" electron orbital) surrounding both atoms as is illustrated in Figure 4.4. It should be emphasized again that in my model the first electron orbital is adjacent to the nucleus and that this claim is consistent with my treatment of thin shell evolution in Section 1.4. Also, notice that the difference between deuterium and the hydrogen 2 molecule under the model is that the innermost shell surrounds both nucleons together in the case of deuterium, but each nucleon individually in the case of the hydrogen 2 molecule. Rather than claim that there is an attractive force holding the bonded atoms together, I claim

that the surrounding shell itself prevents the separation of the atoms. That is, I hold that the innermost enclosing shell plays the role of restraining the bonded atoms from detaching from each other. For more complex molecules than diatomic ones, just the individual atoms involved with the bonds will be surrounded by the enclosing shell. It can be pointed out that subsequent shells will more closely approximate perfect circles as their ordinal position away from their central charged particles increases. This allows for the occurrence of the oscillatory rotational waves discussed in Chapter Three.

With respect to bringing in the concept of spin into the account, one key point is that the electrons in separate subspaces in a "shared" orbital will possess opposite vorticities (i. e. when one is rotating clockwise the other will be rotating counterclockwise). In other words, one electron will be spread out over an entire orbital in one subspace with spin in one direction and in another parallel subspace an electron will be spread out over the same orbital (in the sense that it has the same radius from the nucleus) but with its spin in the opposite direction.

I now close the chapter by making a few remarks on the subject of the so-called "collapse" or "reduction" of quantum mechanical wave packets; i. e., what it is about a so-called "measurement" that results in a quantum superposition of all physically-possible values of an "observable" property "collapsing" into a precise value for that property.

4.4 Reduction of Wave Packets

In this section I address the issue of what it is that reduces wave packets; i. e., the issue as to when particles possessing spread out values come to have definite values. A concrete example in my account would be the question as to when a photon which is spread out over a series of thin shells comes to have a sharp precise location. This issue can also be stated using the symbolism of traditional quantum mechanics in terms of when the Born rule kicks in (the so-called "Heisenberg cut") whereby probabilities of eigenvalues are generated by $\langle \Psi | P | \Psi^* \rangle$ where Ψ is the probability amplitude for the bra vector, Ψ^* (the complex conjugate of Ψ) is the

probability amplitude for the ket vector, and P is a projection operator. In other words, I am raising the question as to at what stage a statistical mixture is created which is characterizable by a density matrix where (at least apart from considerations from non-diagonal elements) elements along the diagonal correspond to subjective probabilities concerning the degree of ignorance of the actual state of the system.

According to the standard Copenhagen interpretation of quantum mechanics a "measurement" on a quantum system collapses a wave packet in the sense that it yields a definite quantity. For example, according to Dirac's 1930/1986, p 36) analysis, "the disturbance caused in the act of measurement causes a jump in this state of the dynamical system." Somewhat similarly, Von Neumann (1932/1955, Ch. 6) held that a measurement changes a probability amplitude Ψ into a definite element in a density matrix. Bohr (1958, p. 73) claims that wave packet reductions occur when they involve "phenomena" which are defined as involving irreversible amplifications such as registrations on photographic plates. Also, with his correspondence principle Bohr held that quantum physics reproduces classical physics for macroscopic effects. However it is notoriously difficult to make a sharp delineation concerning where exactly the distinction between the microscopic and the macroscopic occurs, since clearly the distinction admits of degree. Finally, it can be pointed out that some physicists, such as Von Neumann (1932/1955, Ch. 6, Sec. 1) and Eugene Wigner (1962), have even held that a wave packet reduction only occurs when consciousness is involved.

As far as I can see the issue of whether a measurement (with its implication that a measurement operation be performed) is made, or even can be potentially made, is actually irrelevant to the issue of wave packet reduction. Also, I do not believe that consciousness is required for such a reduction. Instead of postulating macroscopic measurements or consciousness as being responsible for reducing wave packets, I hold that wave packets are reduced in the microcosm with the absorption process. For example, it can be noted that particle detectors, such as the

Geiger counter in the Schrödinger cat thought experiment, work off of the photoelectric effect, which inherently involves absorption, and then amplify the resulting signal. Similarly, in the case of human vision it can be pointed out that the photoelectric effect occurs with the absorption of light by the pigment rhodopsin in the rods of the retina of the eye. This also occurs with the absorption of light by the photo emulsions used in photography.

To flesh out some more details of the account, I hold that physical properties, such as position, only become specific when a particle is absorbed, with them being spread out over all physically-possible situations prior to this. An apt metaphor for what is going on here is the boxer Muhammed Ali's dictum "float like a butterfly and sting like a bee." The spreading out of a wave packet (which I construe realistically) corresponds to the "floating of the butterfly" and the absorption process corresponds to the "sting like a bee."

It should be emphasized that such optical processes as elastic scattering, reflection, refraction, diffraction, and even parametric (light interacting with light) processes in nonlinear crystals, such as up or down conversions, do not involve absorption and thus do not reduce wave packets. Instead, in agreement with what Feynman (1985) has argued, I hold that these processes involve breaking down light so that it takes all physically-possible paths between an initial emitter and a final absorber. Hence, I hold that such optical components as mirrors, lenses, or even nonlinear crystals do not reduce wave packets although detectors, including the rods and cones of the retina of the eye, do because they work off of the photoelectric effect and thus involve absorption. It should also be pointed out, that at least under my account in Chapter Four, entanglement phenomena only become exact during absorption processes, although these involve two-photon (or higher photon number) correlated absorption. Thus, the existence of these phenomena does not constitute a counterexample to my analysis.

There may also appear to be counterexamples to my claim that absorption always constitutes a reduction of a wave packet in cases of superpositions of atoms

in both excited and unexcited states, such as those used in quantum computers. It should be emphasized though that no photons per se actually exist in these states. Also, as Sabine Hossenfelder has pointed out, it is entanglement which actually is responsible for the effects and not that particles are actually in more than one state at the same time. Of course this needs to be fleshed out more but I think that such an account may well be at least in the right direction.

It also wish to argue that so-called FAPP (for all practical purposes) attempts to explain reductions of wave packets in terms of decoherence effects from increases in entropy (see Roland Omnès, 1994) do not work. Decoherence involves wave functions becoming orthogonal (thus disallowing interference) due to coupling with environmental wave functions so that the probabilities of not being orthogonal are extremely low, and for all practical purposes non-existent. The reason that this does not work at a fundamental level is that even extremely remote possibilities still exist, and thus allow for very small superpositional and interference effects. Also, the property of entropy admits of degree and thus does not create a sharp-line cutoff between cases where it exists and cases where it does not exist.

I now discuss two distinct but closely-related dilemmas which can be posed with respect to the issue of the ontological status of these possibilities and probabilities. The first dilemma involves the senses of "possible" and "probable" used in the decoherence interpretation. It can be pointed out that both "possible" and "probable" are ambiguous between ignorance construals and those in terms of something like propensities. Under the ignorance senses of "possible" and "probable" the problem is that under these construals it is presupposed that a wave packet reduction has already occurred even though we do not know which way it has occurred. Thus, under these construals the increases in entropy would not actually cause the reductions inasmuch as definite physical properties (which were merely not known about) would already exist. In contrast, under the propensity senses of "possible" and "probable" the problem is that even remote possibilities

are never actually reduced in the sense of becoming completely nonexistent. That is, it is still physically possible for them to occur even given a set of fixed physical initial conditions. Thus, under these construals these remote possibilities are never completely ruled out from occurring since they remain physically possible. It follows that wave packet reductions never actually occur under these construals.

The second dilemma concerns the ontological status of propensities. In particular, a dilemma can be pressed with respect to the issue as to whether or not propensities construed as potentialities actually exist or not prior to the time when they are actualized. First, it can be pointed out that a potential entity must either exist or not exist – there is no "in between" middle ground. It follows then on the one hand that if these propensities do have a prior existence in some other form, as Aristotle for example held in at least some cases (see Aristotle's discussion of different uses of "potentiality" in his *On Interpretation*, Ch. 12 and in Book Theta of his *Metaphysics*), then they must actually exist, albeit in a different format the details of which we may be ignorant. On the other hand, if it is held that they do not have a prior existence even in some other format, as Heisenberg (1962, Ch. 10) evidently held, then they do not exist. If they do not exist they cannot have a causal influence on the outcomes of subsequent measurements.

On a closely related subject to that just discussed it should be mentioned that purely abstract senses of "potentiality" can also be invoked. In particular, a purely abstract sense of "potentiality" of mere logical possibility; i. e., logical consistency can be appealed to. It might be pointed out that this includes an abstract sense of "physical possibility" of being allowed by the ideal laws of physics (not necessarily as we believe that we know them) either with or without a set of fixed physical initial conditions. Clearly mere logical possibility, while a necessary condition, is not also a sufficient condition for an existence claim since it merely refers to the lack of a logical contradiction. Also, if no ontological commitment is being made to the actual existence of the physical initial conditions then obviously

also no ontological commitment is being made to what would occur given their existence.

I would also like to make a few remarks on alleged connections between interference effects and epistemology. In particular, I disagree with epistemic criteria for the occurrence or non-occurrence of interference in terms of whether trajectory paths for particles are indistinguishable (where there is interference or indistinguishable (where interference effects disappear), as advocated by Bohr (1949), and where the "distinguishability" "indistinguishability" vocabulary was introduced by Feynman (1964, Vol. 3, Ch. 3). Instead the key issue for me is whether the wave packets overlap or not; interference only occurring when the wave packets overlap as I indicated in Chapter Three. As I showed in my treatment of entanglement in that chapter, this can be extended to multi-particle interference effects at a distance, as with the case of entangled-photon (which Klyshko terms "biphoton" for the two-photon case) interference. This shows that appeals to indistinguishability, such as those made by Shih (1999), are unnecessary in accounting for this type of interference as well.

It should perhaps be noted that there has been at least one claim (Shahriar Afshar, 2005) to demonstrate both interference of light and "*welcher weg*" which way information in the same experiment by passing coherent light through dual pinholes and then placing thin wires in regions of destructive interference immediately in front of a lens while still showing a constructive interference pattern at the image plane later. No reduction in intensity results from the placement of the wires and it is inferred that no light is present at the locations of the wires. As I see things though, light is present at the wires, but since there is no absorption there due to the destructive interference, the Renninger effect occurs, pushing the field densities to other directions.

In summary, in my model a reduction of a spread-out wave packet, whereby the packet becomes well-situated in a specific location with specific properties, only occurs during the objective process of absorption. Everything that happens is

construed realistically throughout the whole process here. I take this to be a decided advantage over such alternative accounts as the traditional Copenhagen interpretation whereby it is held that measurements collapse spread-out wave packets. I find the Copenhagen interpretation to be extremely problematic for both subjective and physical construals of the measurement process. Subjective construals of measurement (such as in terms of knowledge or observations or potential knowledge or potential observations) would appear to inherently invoke anthropic considerations whose relevance to a physical model is far from obvious. This is particularly clear in view of the failure of the ignorance interpretation of quantum probabilities to explain such phenomena as interference effects. Anthropic considerations also clearly arise with respect to completely physical construals of measurements, at least if these are construed in terms of anything inherent to the process of measuring per se, as opposed to certain portions held in common among all versions of it such as the photo-electric effect. I find it to be incredibly unclear as to why the issues of human knowledge, even potential knowledge, or the physical process of measurement per se, whether by humans or machines, should play such a role in determining the nature of the physical world.

I now close the chapter by briefly discussing the Einstein Podolsky Rosen (1935) paradox. In their well-known paper of 1935, Einstein, Podolsky Rosen (EPR) speculated that it may be possible to measure the results of non-commuting operations, such as measuring momentum and position, if these operations are conducted at disparate locations. In practice, it is easiest to test this with mutually-incompatible polarization states of photons, as I argued in my section on entanglement in Section 3.2 of Chapter Three. In particular, as I argued there, I hold that what actually is going on in the sorts of situations envisaged by EPR is multiphoton absorption from disparate locations. I then hold that due to the presence of advanced waves an interference pattern is also created. Interestingly, this is even possible for joint measurements of momentum and position with an experiment originally conceived by Popper (1934/1959, sec. 77). In this

experiment, after an interaction causing their states to be correlated, two particles are separated and the position of one is measured and the momentum of the other is also experimentally determined. Kim and Shih (1999) were able to verify Popper's prediction with a biphoton pair generated by down-conversion and more recently Peng et al. (2015) have also demonstrated it with chaotic-thermal light. This would appear to be a counterexample to Heisenberg's uncertainty principle for momentum and position at least if the principle is construed in terms of lack of knowledge, as Heisenberg (1927) clearly does in his original paper on the subject where he uses the German "können" for "know how;" i.e., an "ability" sense of "knowledge." Non-anthropic construals of the uncertainty principles may still be defensible though, such as in terms of the fact that a wave and its Fourier transform cannot simultaneously both be made arbitrarily small; see Messiah, (1958/1999, p. 130). I will not discuss the merits of such an interpretation in this chapter.

Also, as pointed out in Section 3.3, any purported counterexamples to my account based on the results of delayed choices or quantum eraser experiments do not apply since the changes in the experimental setups occur before absorption takes place at the detectors. I now to my final chapter which is a highly speculative discussion of gravity in terms of the model.

Chapter Five

Gravity

In this chapter I attempt to model the force of gravity. Inasmuch as this book is not based on the special theory of relativity, my account is not based on the theory of gravity given by the general theory of relativity – since the general theory is based on the special theory. Instead, my account is based off of the older theory of Newton (1687/1934). Also, my account is reductive in the sense that it does not appeal to additional fields besides the electromagnetic one. One methodological point in favor of such a reduction involves invoking a principle of parsimony whereby it is clearly simpler to postulate just one field to account for both electromagnetism and gravity rather than to postulate separately existing fields for each. In particular in my account I attempt to explain the gravitational force in terms of its being a residual effect of electromagnetism. This is not the first attempt to give an electromagnetic account of gravity. Johann Zöllner (1883) attempted one in terms of the claim that the attractive force between opposite charges is very slightly greater than the repulsive one between charges of like sign. Henrik Lorentz (1900) at least at one time endorsed a similar theory. Henri Poincaré (1906) also explored possible linkages between explanations of gravity and of electromagnetism and even Maxwell (1873/1954, vol. 1, p. 42) expresses some sympathy with the idea but he does not elaborate on it. Also, various electric dipole models have been investigated such as ones by Beckmann (1987, sec. 3.5), Andre Assis (1992), and Lucas (2013, Ch. 8).

Zöllner's theory is *ad hoc* in the sense that it does not postulate a reason to account for why the attractive force would be stronger than the repelling one. A possible alternative might seem to be the claim that there is a surplus of either positive or negative charges in macroscopic matter in order to account for the asymmetry. However, this clearly does not work since, inasmuch as the surplus

charge is completely comprised of the same charge, the net effect would be for these charges to repel each other, rather than to attract. In view of the foregoing considerations it would appear that other alternative versions of linking gravity with electromagnetism are at least worth investigating. It is in this spirit that the speculative theory proposed in this chapter is put forward.

It should be warned right at the beginning of my discussion that this account is the least satisfactory of those in the book, and at least portions of it may strike the reader as being extremely *ad hoc*. However, I believe that aspects of the model may be on the right track and hopefully someone will be inspired to make improvements to the *ad hoc* portions. I begin by a discussion of the relationship between mass and charge, since I hold that these are more closely aligned than traditionally thought. I then move on to an analysis of the gravitational force as being a residual effect of electromagnetism. Finally, I briefly evaluate the merits of some of the purported experimental evidence which has been cited as favoring Einstein's theory of general relativity over Newtonian accounts.

5.1 Mass and Charge

In this section I explore the relationship between the concepts of mass and charge since this topic is integral to my subsequent analysis of the gravitational force as being a residual effect of electromagnetism. It can be recalled from my discussion in Chapter One that I identify positive charges with positive ideal liquids and negative charges with negative ideal liquids. Also, notice that I have not postulated a third type of ideal liquid to correspond to mass. This raises the question as to how mass is to be accounted for in my system.

My strategy is to account for mass properties in terms of properties of charges *per se*. Inasmuch as there is a fixed ratio - 10^{-39} - between the magnitudes of gravitational and electrical forces, it might be thought that there would be similar fixed ratio between the magnitudes of charge and mass. However, the subject is clearly not quite this simple as is shown by such facts as that while electrons and protons have equal fixed charges, the proton's mass is approximately 1836.15 times

as great as that of the electron. I have two possible suggestions to make with respect to this subject, neither of which is very satisfactory, but which perhaps could be developed further. Unfortunately, both of these accounts are incomplete as given in the sense that not all aspects of the subject are covered. Also, much of what I say is both quite programmatic and tentative.

The first suggestion involves the ratio between the respective observed masses of the proton and electron. To account for the measured ratio being approximately 1836, my claim is that a neutron consists of a series of 1836 layers (possibly in parallel subspaces) of paired opposite charges. Obvious problems for such an account include the fact that the ratio between the observed masses of the proton and electron is not exactly integral. Also, while in the process of beta decay a free neutron decays into a proton and an electron, the proton does not continue to decay into a particle with a higher multiple of a unit charge. Thus, there is no independent evidence for the existence of multiple opposite charges. It might also be noted that while it is usually claimed that a neutrino is produced in the beta decay process, this has been questioned (as discussed by David de Hilster, 2011) since it is only when the mass of the neutron is relativistically adjusted that there is a need to postulate the neutrino for conservation purposes.

The second suggestion involves attempting to work with some variant on the quark model (although possibly not so as to involve any point particles) in the so-called "standard model" of particle physics. Under the quark model baryons, such as protons and neutrons, are held to consist of three quarks. Up quarks are held to possess a $2/3$ positive charge and down quarks a $1/3$ negative charge. A proton, consisting of two up quarks and one down one will thus possess an integral positive charge, and a neutron, consisting of one up quark and two down ones will possess a neutral charge. While quark theory is based off of relativity, it can be pointed out that the claims about combining non-integral charges could be made independently. Still, it would both seem to be rather *ad hoc* where the $1/3$ factor for charge in the

model comes from and also it should be emphasized that unlike my program the standard model posits mass as being distinct from charge.

I now make a number of points of both comparison and contrast between mass and charge and then move on to discuss the plausibility of accounting for gravity in terms of its being a residual effect of electromagnetism. Under their classical physics conceptions mass and charge have some properties in common since they are both scalar properties of matter and have vector fields associated with them – respectively, the gravitational field for mass and the electrical field for charge. Also, as I develop in more detail in Section 4.2, both of these vector fields are central force fields (a concept which I explicated in Section 1.1) with their intensities being inversely proportional to the squares of their distances from their source particles. There are also some key differences between mass and charge. For example charge is quantized while mass is not. Also, there are two types of charges – positive and negative – while there is only one type of mass. Hence, the effects of mass do not cancel out in the way that opposite electrical ones do, and thus the mass of each massive particle contributes to the overall total mass.

A distinction can be drawn between gravitational and inertial mass. As is well known, Newton in his *Principia* (1687/1934, Book 2, sec. 6) described a series of pendulum experiments where the oscillations of pendulum bobs are timed when the bobs are comprised of different substances, and thus argued that gravitational and inertial mass are directly proportional to each other. In this regard Einstein (1922/1974/ p. 56) cites the torsion balance experiments of Lorand Eötvös which were considerably more precise than Newton's. It should be emphasized though that both Newton's and Estros's experiments were only conducted in reference frames which were mutually stationary between the observer and the instruments. This leaves it as an open question as for what happens when this is not the case. In fact, I hold that it is not the case for non-inertial reference frames as I discuss next.

The first point I want to make concerning inertia (i. e., the resistance to acceleration by a countervailing force) is that the concept can be applied to charge

as well as mass. In particular, I hold that charge inertia involves the resistance of a charged body to an electrical force field. The ratio of the inertial mass to the inertial charge is presumably the same ratio 10^{-39} ratio as the ratio of the gravitational and electrical forces. Beckmann (1987, p. 184) makes this point as well and I believe that it is a matter which merits further experimental investigation.

I believe that the distinction between gravitational mass and inertial mass is key for understanding claims about purported mass increases in particle accelerators. In particular, I hold that this apparent mass increase just pertains to inertial mass and not gravitational mass. Beckmann (1987, sec. 1.7) and Tom Bethell (2009, Ch. 8) discuss a similar alternative to the special relativity explanation of purported mass increases in particle accelerators. In particular, they argue that instead of the quantity of matter increasing in particle accelerators it is a change in the resistance to a force changing a body's momentum. In other words, more energy is required to accelerate a given fixed mass since the magnitude of the effective force decreases as a function of the relative velocity of the fixed mass with respect to that of the origin of the force.

It is noteworthy, as Bertrand Russell (1925/2009, p. 94) pointed out back in 1925, that the increase of mass was well known earlier than Einstein with experiments with accelerated electrons. In fact, Lorentz (1904/1952) derives a similar formula from considerations concerning the transverse and longitudinal electromagnetic masses of the electron. It is also interesting that at least purported derivations of the formula for mass increase from the conservation of momentum in popular expositions of special relativity, such as those by Bohm (1965) and Max (1920/1962), are extremely convoluted at least compared with the derivations of the formulas for time dilation and length contraction from the Pythagorean theorem. This suggests that, instead of having been used for purposes of discovery, the derivations of mass increase were in fact made *post hoc* where the derivers knew in advance what they were looking for.

It should also be pointed out here that typically the accelerated body will consist of a very large number of charged particles, and also that typically the effects from opposite charges will cancel out. This is not the case with inertial mass though, and hence its effects will often dominate. I might also mention that, in the case of light, I do not hold that photons possess gravitational mass since I do not hold that they possess their own shell systems. However, as I noted in section 3.1, there is good empirical evidence that photons possess momentum, and hence interial mass. I now turn to a discussion of the linkage between charge and gravitational mass in terms of properties of the fields associated with each.

5.2 Gravity as a Residual Effect of Electromagnetism

As pointed out in my discussion of points of comparison between mass and charge in Section 4.1, both the electric and gravitational fields are central force fields and their strengths vary at a rate which is inverse to the square of their distances from their source particles. In particular, Coulomb's law

$$\mathbf{F} = \frac{1}{4\pi\epsilon_0} \frac{Q_1 Q_2}{r^2} \mathbf{r} \quad (5-1)$$

discussed in Chapter Two can be compared with Newton's law of universal gravitation

$$\mathbf{F} = G \frac{m_1 m_2}{r^2} \mathbf{r} \quad (5-2)$$

where G is the universal gravitational constant. The fact that these two laws have the same structure at least suggests that the gravitational field may be the result of the same mechanism as that responsible for the electric field, possibly being a remnant of it.

It can be observed that in spite of his dictum "hypotheses non fingo" from the General Scholium to his Principia (1687/1934) Newton apparently believed that the gravitational force was transmitted instantaneously. However, it can also be pointed out though that prior to Einstein's General Theory of Relativity, Paul Gerber (1898) had a field theory of gravitation in which it was held that the speed of the field is the same as that of light. As noted in Section 5.1, the ratio between the magnitude of the electric force and that of the gravitational force is

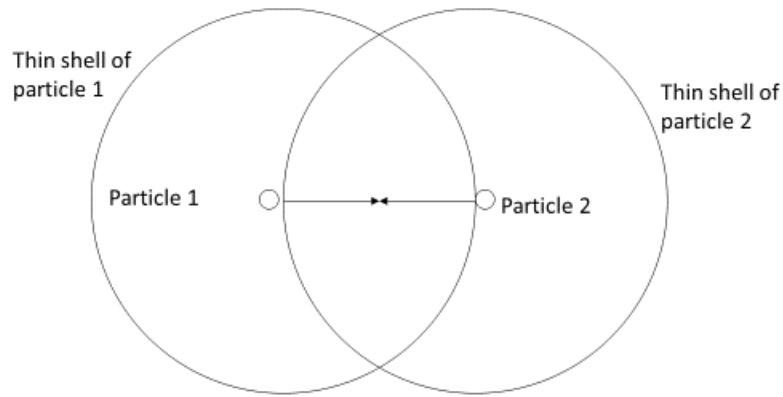


Figure 5.1 Reciprocity between thin shells of two particles

approximately 10^{-39} . Thus, the 10^{-39} ratio for the relative strengths of the electrical and gravitational forces associated with each thin shell should remain the same for each of these shells, even though the absolute magnitudes for both decrease at an inverse square rate. For electrically neutral bodies the electrical forces cancel out which will leave as a remnant the gravitational force. I now discuss the relative merits and demerits of a series of residual models, none of which are entirely satisfactory.

I first consider models which hold that the magnitude of a remnant decreases in one step with the ratio of the residual width to shell width decreasing at a $1/r^2$ rate with respect to each shell where r is the distance from the center of the shell system to the shell in question. It can be noted that this results in the ratio of the residual width to the overall radius to decrease at a $1/r^4$ rate. Also, it can be seen that the same 10^{-39} ratio between the magnitudes of the electrical and gravitational forces will be in effect for the initial thin shell after an originating particle as well

as for each subsequent shell. An attractive feature of models based on such residuals, at least with respect to other models which I will go on to discuss, is their non-composite character. However, they are not at all clear with respect to the issue of what creates the residual ratio in the first place. Also, they are *ad hoc* both with respect to the issue of what could be responsible for such a small residual ratio and for why the resulting force would be an attractive force.

A somewhat better motivated one-step theory involves taking note of two alternative processes for a causal connection between thin shells and charged particles – one from the shells to the particles and the other from the particles to the portion of the shells immediately in front of the particles. The reciprocity between the natures of these two processes is illustrated in Figure 5.1. Notice that with respect to the first process, which was discussed in detail in my treatment of electromagnetism in Chapter Two, the inverse square rate is determined by the width of a thin shell at the location of a particle. With respect to the second process notice that, inasmuch as the width of the thin shell there varies at a rate of $1/r^2$, this process results in an inverse square ratio between the fixed width of a particle and the width of the thin shell at that location. Also note that since these are distinct processes it is at least conceivable that their magnitudes are quite different. Thus, it is at least conceivable that one process can be identified with the gravitational force and the other with the electric force.

To show that the gravitational force is attractive and not also repulsive, as with the electric force, an asymmetry with respect to the relative widths of the inner and outer portions of a thin shell (which I develop in more detail with my discussion of two-step theories next) can be pointed to as constituting a “standing condition” whose presence is requisite for the occurrence of the process. Admittedly invoking the necessity of such a “standing condition” here is *ad hoc* inasmuch as it is not clear why such a condition should just apply to one side of a thin shell and not also to the other side. Possibly a non-*ad hoc* strategy in this context would be to claim that, rather than pulling the respective shells further apart,

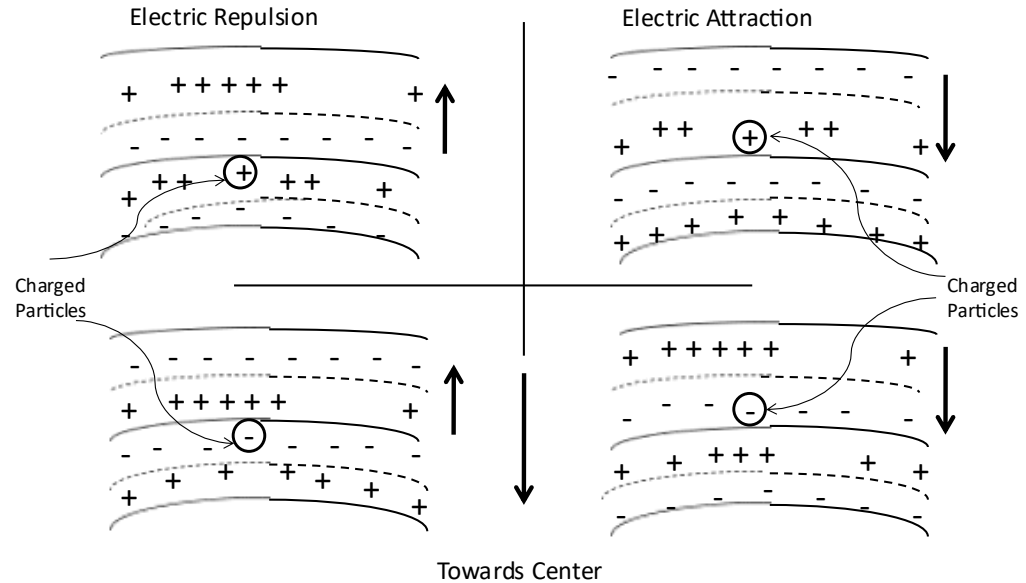


Figure 5.2 Asymmetry in thin shell widths (showing only a small portion of the shells)

instead the causal connection with the outermore thin shells actually draws the whole shell system inward closer together, perhaps by changing the intensities of the oscillations between the shell halves in the immediate region of the particle in question. Admittedly such a move is sketchy but hopefully somebody else can develop it in more detail.

In any event, I now move on to my discussion of two-step processes. Admittedly these processes have the demerit of being more complex in character than the one-step processes. However, at least the models which I consider also have the merit of being less *ad hoc* since they include a remnant factor which varies at an r^{-5} rate which in principle could account for the 10^{-39} ratio in relative strengths

between the gravitational and electric fields. To motivate these accounts, I wish first to elaborate on the asymmetry in thin shell widths just mentioned and which is illustrated in Figure 5.2. Notice that in all four cases the inner shell possesses a greater width than the outer shell to compensate for the fact that the spherical area being encompassed by the inner shell is less. It can be postulated that this greater width (in some manner whose nature I will not speculate about) serves as one factor resulting in the attractive forces towards the center of the shells being slightly stronger than the repulsive forces away from these centers. Still, some quantitative analysis can be done on the nature of this residual factor as I will now show.

Since each portion of a thin shell is $\frac{1}{2}$ the volume of a complete shell, the ratio between the two is the same as that occurring between adjacent whole shells. If n is the ordinal position of a thin shell and shell volumes are normalized to 1, the width of the n^{th} thin shell W_n is thus given by

$$W_n = \sqrt[3]{\frac{3}{4\pi} V_n} - \sqrt[3]{\frac{3}{4\pi} (V_n - 1)} \quad (5-3)$$

For successive thin shells, their difference $W_n - W_{n-1}$ will thus in turn be given by

$$W_n - W_{n-1} = \left[\sqrt[3]{\frac{3}{4\pi} V_n} - \sqrt[3]{\frac{3}{4\pi} (V_n - 1)} \right] - \left[\sqrt[3]{\frac{3}{4\pi} (V_n - 1)} - \sqrt[3]{\frac{3}{4\pi} (V_n - 2)} \right]$$

or

$$\sqrt[3]{\frac{3}{4\pi} V_n} - 2\sqrt[3]{\frac{3}{4\pi} (V_n - 1)} + \sqrt[3]{\frac{3}{4\pi} (V_n - 2)} \quad (5.4)$$

This can be shown to decrease at approximately a $1/r^5$ factor when the radius of a shell is doubled by the following considerations. As I showed in Chapter Two, the difference in actual widths of thin shells decreases approximately at an inverse square rate. In particular, if the radius doubles the widths decrease by a factor of approximately $\frac{1}{4}$. Since the radius of a sphere is directly proportional to the cube

A	B	C	D	E
V_n	$R(V_n)$	$B_n - B_{n+1}$	$C_n - C_{n+1}$	D_n / D_{n+1}
511	4.95949130997209	.00323301375716	.00000420693844	32.000276
512	4.96272432372925	.00322880681872		
513	4.96595313054797			
4095	9.92457206258504	.0008077871444	.0000001314656899	32.0008008
4096	9.92537984972948	.00080765567871		
4097	9.92618750540819			
32767	19.8504201739495	.0002019310051	.0000000041082	32.01661146
32768	19.8506221049546	.0002019268969		
32769	19.8508240318515			
262143	39.7009185408561	.00005048195199	.0000000001283146408	
262144	39.7009690228081	.00005048182368		
262145	39.7010195046318			

Figure 5.3 Computation showing that as radius doubles ratio of the difference in thin shell widths is 2^{-5} . V_n is the enclosed volume of the n^{th} shell and $R(V_n)$ is the overall radius of that shell.

root of its volume, for successive thin shells whose volumes have been normalized to 1, their successive radii will increase at a rate proportional to $\sqrt[3]{N}$ respective shells further apart, instead the causal connection with the outermore thin shells actually draws the whole shell system inward closer together, perhaps by changing the intensities of the oscillations between the shell halves in the immediate region of the particle in question. Admittedly such a move is sketchy but hopefully somebody else can develop it in more detail.

One way to try to handle the $1/r^5$ rate here is to just postulate that the magnitude of the remnant diminishes at the $1/r^5$ rate for a while until the 10^{-39} ratio

between the width of the remnant and the total width of the thin shell is reached and then $1/r^2$ "kicks in." It might be thought that such a theory would be refutable, at least in principle, if the strength of the force can be measured before the threshold ratio for the transition to the $1/r^2$ rate kicks in. However, if it is assumed that the width of an initial thin shell is on the same order of magnitude as an electron orbital (say an angstrom) and taking note that the radius must double three times for each order of magnitude, along with the fact that the width of the residue decreases by a factor of $1/32$ each time that the overall radius doubles, it can be seen that the overall radius will be less than a micron when it reaches the 10^{-39} ratio. Still, such a proposal would appear to be completely *ad hoc* with respect to the issue of why there would be a sudden switch from a $1/r^5$ rate of decrease to a $1/r^2$ rate. I will not consider this sort of model further, but instead will now consider models which postulate that the gravitational force is solely a residual effect of the electrical force without changing the rate of decrease in the strength of the gravitational field. Unlike the foregoing models, these models have a composite character possessing both residual and augmenting factors.

I will now assume that there is no change in the rate of decrease in the strength of the gravitational field. Thus, I am postulating that the gravitational force is a result of a residual effect of the electrical force and that the force itself remains electrical in character. One merit of such an account which deserves emphasis is that of relative simplicity. This is because it does not assume a separate field besides the electric one. That is, the residual factor does not itself constitute a distinct entity from the electric field.

Since the magnitude of this residual factor *per se* is proportional to $1/r^5$, a compensating augmentation factor proportional to r^3 is required in order to account for the inverse square rate of the observed gravitational force. One manner to motivate the r^3 factor is to invoke the enclosed volume of a given thin shell inasmuch as this volume is directly proportional to r^3 . The product of the enclosed volume with the remnant would then result in the gravitational force. A merit of

this account is that it only appeals to one factor. Also, postulating the enclosed volume as compensating for the $1/r^5$ difference does not just constitute a return to the original r^{-2} case since it only applies to the attractive central force. An obvious demerit of the account though is that it appeals to elements outside the thin shell *per se*, although perhaps these elements could be invoked to help explain the role of advanced waves in polarization entanglement which I discussed in Section 3.2.

If the r^3 factor is to be accounted for in terms of factors within the shell system *per se*, one possible way to do this is by invoking the product of the area of the shell (which is directly proportional to r^2) by the maximum circumference of a circle located on the surface of the shell; i. e. a great circle (which is directly proportional to r). One merit of such an account is that it just appeals to properties of the shells *per se*. Obvious demerits though are that it appeals to two factors (the area and the great circle circumference) and that there is no obvious way to motivate the account.

A closely-related variant on the foregoing accounts involves a "counting factor" which can be motivated as follows. It can first be observed that, by the original defining hypothesis, the volume of each thin shell is a constant. Thus, if this volume is normalized to one, the total enclosed volume will increase in integral increments. This also means that the total enclosed volume will correspond to the ordinal position of a closed shell. Motivated by the preceding consideration, one possible scenario to account for a compensating factor involves postulating that in the shell formation process each time a new shell is created, an equally-weighted factor is added to the multiplicative factor. It can be pointed out that such an equally-weighted factor (suitably normalized, by the unit of electric charge of 1.6×10^{-19} coulombs together with the appropriate mass to charge ratio), would in effect serve as a counter inasmuch as it would increase by one unit at each repetition. An obvious issue for such an account is to raise the question as to why it would apply to the case of gravitational forces and not also to the case of electrical forces.

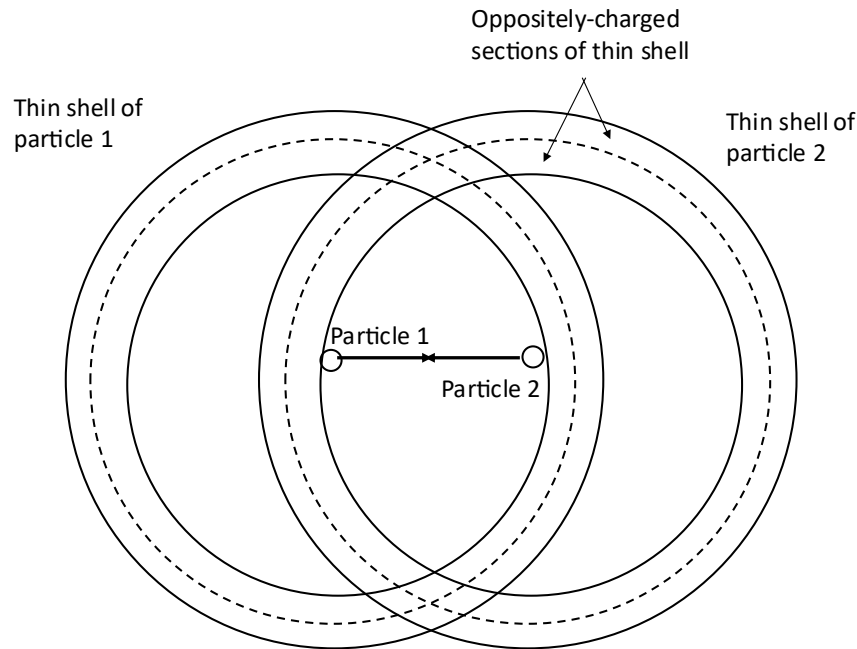


Figure 5.4 Diagram for a “push” theory of gravitation

Depending upon whether the causal direction goes from any of the r^3 factors to the r^5 remnant factors which constitutes the difference between the widths of adjacent thin shells or vice versa from these remnants to any of the r^3 factors, two possible types of models arise. One possible model for the former case of causal direction would be to hold that the "inertia" (resistance to change of position) from the entire r^3 factor is localized (without being absorbed) where a thin shell encounters a charged particle so as to enhance the magnitude of the attractive force. A possible model for the latter case of causal direction would be to hold that “friction” (or some other form of resistance possibly also connected with properties of ideal liquids) from the 10^{-5} remnant factor serves as to modulate the effect of the r^3 factor on a charged particle. Of course, either of these accounts would also have to explain why these ideal liquid properties would apply in the case of gravity but not also in the case of electromagnetism I find the “inertial” models where the r^3

factor constitutes the causal factor (or possibly a reciprocal causal factor if both the r^3 factor and the r^{-5} remnant factor are taken to exist) to be more plausible than the “friction” models where the r^3 factor constitutes the effect inasmuch as it is more straightforward how the resultant r^{-5} remnant factor would then constitute a central attractive force. Also, the nature of the remnant gravitational force must be motivated independently of electrical forces; which, presumably, are completely accounted for in terms of oscillations of the thin shells per se. The remnant factor must also interact with the enclosed volume in such a manner that the location of the enclosed volume determines the direction of the force where its magnitude is directly proportional to the product of the enclosed volume and the remnant factor. One problem with this account is that the width of the outer section of a thin shell is greater than the width of the inner section, and thus it is not clear why the resulting force would be an attractive one.

One alternative for getting around this last point is illustrated in Figure 5.4 and involves the point that for a thin shell outside of particle, the wider inner section of the shell is adjacent to the particle. The claim would then be that this section “pushes” the impinged-on particle inwards towards the particle on which the thin shell is centered so as to draw the particles together. Since presumably there would also be an opposite “push” from the inner shell the resultant force would then be proportional to the difference in widths between the two sections and hence, as I have just shown, would be proportional to r^{-5} . One of the just-discussed processes linking the thin shell in question with its enclosed volume would have to then be invoked to create the inverse square nature of the gravitational force. It should be emphasized that this process would not involve electrical forces which would in this case be repulsive in nature.

Admittedly, none of the just discussed models based on the width of remnants from the difference between inner and outer shell widths, at least on the surface, appear to be very credible. Instead, they appear to be rather contrived in various ways and thus to be *ad hoc*. In view of this *ad hocness*, it would seem that

alternative accounts should also be considered. One such account involves whole shells including interior points; i. e. what in topology are called "balls." I now turn to a discussion of two variants of that account.

The first variant of the account involves the difference in overall radii from the center of adjacent thin shells with an associated force whose magnitude is postulated to be 10^{-39} that of the electrical one. Since this difference is the same as the width of a thin shell it also varies at an inverse square rate from the center. However, it involves two successive shells both of which include the entire radius. In the first variant of the account it is held in this account that the interior shells of a thin shell system contribute to the whole effect. A virtue of this variant is that it does not have a composite nature like the latter variant. A negative feature though is that it is not at all obvious how the interior shell regions can causally fulfill such a function inasmuch as there would be a violation of Descartes' principle of contact action if it is postulated that there is an immediate causal influence from these interior regions.

In the second variant of the account it is held that each thin shell is linked (possibly by ideal liquids) with a whole shell including interior regions. A negative feature of this variant is clearly its *ad hoc* character. Positive features of the variant include that it avoids complications with interior structures and also, as with the appeal to enclosed volumes of shells earlier in this section, it might help to explain the advanced waves invoked in my discussion of polarization entanglement. In particular, the model would appear to imply a rigid structure for these interior regions which would be required for the instantaneous action of the advanced waves.

The foregoing single-factor account may appear to be neither more nor less *ad hoc* than the multiple-factor account since both accounts arbitrarily bring in an *ad hoc* factor whose nature is not specified. However, it can still be observed that in spite of this point concerning the multiple-factor account, the asymmetry in shell widths in the residual R^{-5} factor is not *ad hoc* since it involves an actual difference

in shell widths. I am much more confident that this is one of the factors than I am about the nature of the R^3 compensating factor. Also, since the multiple factors constitute independent parameters their magnitudes can be assigned arbitrarily so that their product matches a predetermined amount such as that of the 10^{-39} ratio between the strengths of the gravitational and electrical forces. Thus, even though the multiple-factor account is still incomplete and needs much more work, it at least suggests a program for an explanation of gravity including both its inverse square nature and possibly also the ratio between the strength of the gravitational force and the electrical force. Obviously, a lot more details are required in order to flesh this out though.

5.3 Discussion of Experimental Evidence

It is true that some empirical evidence has been cited as favoring the gravitational theory of Einstein's general theory of relativity over Newton's theory; including an alleged gravitational red shift, the movement of the perihelion of the planet Mercury, and purported shifts in the direction of light by massive bodies, as evidenced by positional shifts in starlight observable during solar eclipses and so-called "gravitational lensing" resulting in Einstein rings. The strength of this empirical evidence is debatable though. For example, it is very difficult to distinguish between a Doppler red shift and a gravitational one, at least for astronomical sources, as Bruno Bertotti et al. (1962) point out. The results of experiments with terrestrial sources, such as the experiment utilizing the Mössbauer effect by Robert Pound and Glen Rebka (1960) in an elevator shaft at Harvard University, have also been disputed concerning the degree of precision possible with them; see for example the discussion of Alessandro Cacciani et al., 2006. Similarly, it turns out that Gerber's (1898) theory can also account for the shift in Mercury's orbit. Also, an alternative account of shifts in starlight by massive bodies can be given in terms of the refraction of the light by stellar coronas, such as that of the sun – see the discussions of Edward Dowdye (2007).

It is sometimes alleged that general relativity is required in order to properly synchronize the clock system of the global positioning system (GPS) due to the difference between the strength of the gravitational field at the surface of the earth and at the height of the satellites of the system. However, the so-called “relativistic correction” here can be better accounted for by a combination of two non-relativistic effects. One is the Sagnac effect (the phase shift of a rotating interferometer which I discussed in Section 3.1), used with respect to the earth-centered inertial frame of reference (ECI frame) together with an absolute reference frame for time as discussed by Ronald Hatch, 1995. The other is a correctional term for the gravitational field. Of course, an account is required here for the existence of the gravitational field, but at least the rudiments of such an account have been given in this chapter. It should also be pointed out that while no correction for the emitted signal from a satellite is given for the motion of the satellite in the GPS system, this is probably still consistent with an emission theory of light, such as the one given in this book, due to the extinction effect whereby light in the upper ionosphere in which the satellite is traveling is constantly obtaining new sources of emission in this realm.

It can be also noted that it is possible to account for the Sagnac effect without appealing to relativity. For example, Franco Selleri (1996) has shown that the Sagnac effect can be explained by postulating absolute simultaneity in a primitive rate of rotation and noting that special relativity predicts no shift for an observer located on the rotating platform while this is in fact the case. Also, I have given an account of the Sagnac effect in Section 3.1 in the context of an emission theory of light. I should also emphasize that it would be anthropomorphic to think that the ECI frame is the basic reference frame for the whole universe, and this suggests the wisdom of testing the GPS corrections (or any other phenomena based on the ECI frame) near astronomical bodies other than the earth should this ever become feasible.

Appendix A

Are Complex Numbers Essential to Quantum Mechanics?

In the book I have sedulously avoided the use of complex numbers in my treatment of quantum phenomena. This is in marked contrast to most standard treatments. In fact, it is sometimes held that the usage of complex numbers in quantum mechanics is essential and not just a useful shortcut in the mathematics. For example, Bohr (1928), Schrödinger (1927/1982, p. 171) Feynman (1961, Vol. III, pp. 1-6, 7-5), Roger Penrose (1991, pp. 79, 236; 2004, sec. 21.6) and Sunny Auyang (1995, p. 74) have made this claim. Certainly, complex numbers are ubiquitous in standard formulations of quantum mechanics. For example, they occur in the time-dependent Schrödinger equation and in Dirac and Von Neumann's state vector approach they occur in both the state vectors themselves and often also with the operators on them. For example, they occur with the rendition of the Heisenberg uncertainty relation in terms of the commutator $[P_x, X] = P_x X - X P_x = i\hbar$ where P_x and X are the respective momentum operator in the x direction and the position operator. Also, in this regard Roy Glauber (1963) has asserted that they are an essential element of the electric field operator, and he claims that different predictions, including correlations between photons, are made when using the operator as opposed to the classical field.

However, if quantum states reconstrued realistically and not just as part of a calculating device for observables (as held under the positivist philosophy which predominated when modern quantum mechanics was first formulated), it is very hard to see how the usage of complex numbers can be truly fundamental. In particular, when physical properties are either measured quantitatively or are even indirectly computed, the magnitudes of these properties can inherently only be characterized by real numbers. Thus, if quantum properties are to be construed

physically they must also be characterizable by real and not complex numbers. I deliberately do not address issues concerning the merits of hidden variable theories of quantum mechanics other than to note that traditional refutations of these theories are directed at local hidden variable theories and not global ones such as those given in this book and by Bohm (1993).

In modern quantum mechanics expectation values for observables are represented by the product of a state vector with its Hermitian conjugate operator. This product results in a real number, even though each individual factor is expressed as a complex number, since Hermitian matrices are self adjoint (i. e., the transpose of the complex conjugate of each element of the matrix equals the original matrix). Thus, the resulting expectations values are real. This is fine if, as under the Copenhagen interpretation, we are only interested in dealing with observables. However, as just noted, it is very hard to see how the quantum states themselves can then be interpreted realistically under these renditions.

A somewhat analogous point to the foregoing involves a comparison of the Schrödinger and Heisenberg approaches to quantum mechanics. Schrödinger, with his wave mechanics, has the time-dependent term (the eigenfunction expressed as an exponential $e^{-iEt/\hbar}$) in the state vector and Heisenberg, with his matrices, has the time-dependent term included in the operators. In effect then, quantum mechanics only requires a time-dependent occurrence once in the product of the operator and state vector, and it is arbitrary in whether it occurs with the operator or the state vector. As in the case of Ψ and its complex conjugate Ψ^* a product thus must also be taken here in order to create the expectation value (eigenvalue) of an observable.

One move that can be made towards avoiding complex numbers in a realist interpretation of quantum mechanics involves utilizing Euler's identity whereby exponential functions can be rendered trigonometrically as $e^{ix} = \cos(x) + i\sin(x)$. It might seem that this just is a compact notation for expressing two orthogonal waves; in particular since when multiplied by the complex conjugate the identity is rendered as $\cos^2(x) + \sin^2(x)$. When this is applied

to quantum mechanics both waves come into play, since the phase difference, although not the absolute phase matters. Both Penrose (1991, p. 236) and Auyang (1995, p. 74) assert that accounting for this phase difference requires the usage of complex numbers, but it is not at all clear to me why this is the case. In particular, I suggest in Section 3.1 that probability amplitudes for both the sine and cosine waves be first individually vectorially summed and the resultants then squared. If the two resultants (which can be identified with energy density fields and not force fields) are then summed, it may be possible to get around Penrose and Auyang's point. Admittedly though this issue requires further analysis, particularly in respect to avoiding complex numbers both in probability amplitudes themselves and in their associated operators.

A series of other moves questioning the necessity of complex numbers for quantum mechanics are alluded to by Eckard Blumschein (2018). These include issues concerning the foundations of Fourier analysis such as claiming that only the real values (i. e., the cosine functions) and not also the imaginary values (the sine functions) are necessary and claiming that the temporal integration need only be taken over the past (corresponding to the real values and not also the future (corresponding to the the imaginary values) as with an integration from positive to negative infinity. While I find such suggestions to be intriguing I do not evaluate them here.

Complex numbers are also introduced to factor expressions of the form $x^2 + y^2$ into the form $(x + iy) \times (x - iy)$. It can be conceded that there is no general algebraic solution to the equation $x^2 + y^2 = z^2$. Similarly, a square root for $x^2 + y^2$, cannot be expressed in terms of x and y alone (due to the cross term $2xy$). Nevertheless, it should be emphasized that there will still be roots for particular numerical here. In other words, for each possible place value which can be substituted for the variables x and z respectively, there will be a place value for y which will make the equation come out true. However, it should also be noted in this case that the place value for y will typically be irrational, and also that there

are no standardly-defined functions for characterizing these relationships. Thus, as with the case of the Euler identity, the usage of complex numbers here would appear to be a matter of convenience rather than one of necessity.

A connection can be made between the preceding points and properties of the electric and magnetic fields. In particular, in the case of light, it can be noted that the energy density of the electromagnetic fields is given by $E^2 + B^2$, where E and B are the respective magnitudes of the $\mathbf{E}$ and $\mathbf{B}$ fields. When this is factored into $(E + iB)(E - iB)$ a parallel can be noted with the probability amplitudes Ψ and its complex conjugate Ψ^* so as to construct a photon wave function; see Iwo Bialynicki-Bibula (1996). However, it can be pointed out that an alternative to factoring the whole expression here would be to add the separate quadratic parts, which can be construed as energy densities as I discuss in Section 1.1 and in Chapter Three.

In closing, something should be said with respect to the subject of the ontological status of complex numbers. I believe that imaginary number were named "imaginary" with good reason; -1 does not have a square root. Also, complex numbers are not just ordered pairs of real numbers, as is sometimes claimed. This is only true if special rules for multiplying the ordered pairs are included. In particular, $(x_1, y_1)(x_2, y_2)$ is defined as equaling $(x_1x_2 - y_1y_2, x_1y_2 + x_2y_1)$.

I want to emphasize that the preceding remarks are not meant to discourage the use of complex numbers in science as a useful shortcut for making calculations such as in factoring certain quadratic equations or in solving many classes of differential equations. However, it must be possible, at least in principle, to cash out this usage in terms of functions not containing complex numbers such as with trigonometric functions. For example, consider the usage in electrical engineering of the complex impedance $Z = R + jX$ where j is used instead of i as the symbol for an imaginary number so as not to be confused with the symbol for current. In this context j is used as a useful shortcut for making calculations by characterizing the counterclockwise rotation of a vector in the complex plane. Similarly, in the

treatment of two-photon absorption in chemistry the coefficient β is proportional to the imaginary portion of the third order non-linear optical susceptibility $\chi^{(3)}$. The physical properties here clearly are not imaginary even if the functions used for characterizing them are expressed in terms of complex numbers.

Appendix B

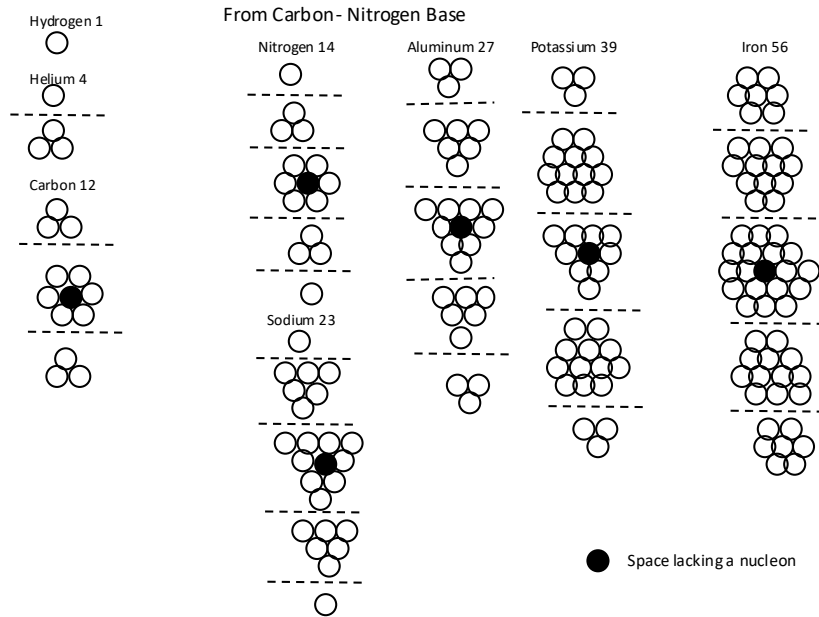
Proposed Nuclear Structures

The following is a set of speculative proposed nuclear structures for common elements. As far as I can tell it is consistent with the remainder of the model although it is not developed in the actual text. The model is based on construing nuclei in terms of tightly packed spheres, and where thus the overall radius R is proportional to $N^{1/3}$ where N is the nucleon number. There are obvious parallels here with so-called "liquid drop" models of the nucleus.

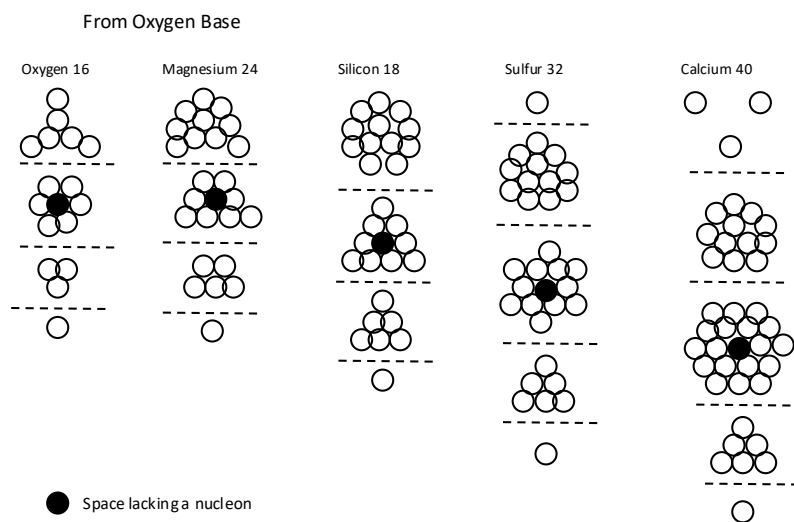
In the model the proposed structure of the helium nucleus is the tetrahedron which possesses a 4-fold symmetry. Nuclei with higher nucleon numbers than helium build on this tetrahedral structure with various symmetries of their own. Instead of postulating a tri-alpha process (as originally suggested by Fred Hoyle) for the formation of nuclei involved in the carbon-nitrogen-oxygen cycle, I postulate a quadri-alpha process whereby four helium nuclei (alpha particles, nucleon number 4) combine to form an oxygen nucleus (nucleon number 16). Given the extremely short half-lives involved with the formation of these nuclei (respectively $8 = 10^{-17}$ s for beryllium 8 and 2×10^{-16} s for the so-called Hoyle state of carbon 12) along with the fact that the carbon-nitrogen-carbon sequence is cyclical in character, I do not see that there is any empirical evidence that the cycle does not begin with oxygen. It can also be pointed out that under my model when each of four helium nuclei combine by joining along one of their 2-nucleon edges to form an oxygen nucleus this results in a "gap" in the center of the proposed structure. It can be observed that subsequent nuclei, starting with carbon (nucleon number 12) also possess this gap.

I propose two separate bases for nuclear bases with nucleon numbers of 12 and above - one off of a carbon (and subsequent nitrogen) base and one off of an oxygen base. The oxygen base builds off of the four-tetrahedron base forming the

oxygen nucleus. Subsequent nuclei with an "oxygen base" build off of this structure. With respect to the carbon - nitrogen base I leave it as an open question as to whether the carbon nucleus forming this base originates from the oxygen nucleus (by losing the four nucleon protuberances that project outward), or instead is due to some other mechanism. Notice that the carbon through nitrogen series of nuclear models builds on the carbon - nitrogen base and not on the oxygen base. It can also be observed that the outer shell of the iron (the ninth most common element in the universe) nucleus constitutes a "perfect shell" surrounding the carbon (the fourth most common element) nucleus, which itself constitutes a perfect shell.



Nuclear structures



Nuclear structures continued

Bibliography

Afshar, Shahriar, 2005, "Violation of the Principle of Complementarity, and its Implications" *Proceedings of SPIE* vol. 5866, pp. 229-244.

Alley, Carroll and Shih, Yanhua, 1987, "Foundations of Quantum Mechanics in the Light of New Technology," Physical Society of Japan, Tokyo, p. 47.

Aristotle, 1941, *The Basic Works of Aristotle*, edited by Richard McKeon, New York, Random House.

Aspect, Alain, Dalibard, Jean, and Roger, Gérard, 1982, "Experimental Test of Bell's Inequalities Using Time-Varying Analyzers," *Physical Review Letters* 49, pp. 1804-1807.

Assis, Andre, 1992, "Deriving Gravitation from Electromagnetism," *Canadian Journal of Physics* 70, pp. 330-340.

Auyang, Sunny, 1995, *How is Quantum Field Theory Possible?*, Oxford, Oxford University Press.

Batchelor, G. K., 1967, *An Introduction to Fluid Dynamics*, Cambridge, Cambridge University Press.

Beal, Alasdair, 2024, "Special Relativity and the Lorentz Equations: Errors in Einstein's 1905 Paper," *Physics Essays* 37, pp. 46-54.

Bell, John, 1987, *Speakable and Unsayable in Quantum Mechanics*, Cambridge, Cambridge University Press.

Becker, Richard, 1964, *Electromagnetic Fields and Interactions*, New York, Dover.

Beckmann, Petr and Mandics, P., 1965, "Test of the Constancy of Electromagnetic Radiation in High Vacuum," *Radio Science*, 69D, pp. 623-628.

Beckmann, Petr, 1987, *Einstein Plus Two*, Boulder, Golden Press.

Bergman, David and Wesley, Paul, 1990, "Spinning Charged Ring Model of Electron Yielding Anomalous Magnetic Moment," *Galilean Electrodynamics*, 1, pp. 63-67.

Bertotti, Bruno; Brill, Dieter; and Krotkov, Robert, 1962, "Experiments on Gravitation" in *Gravitation: an Introduction to Current Research*, edited by Louis Witten, New York, John Wiley, pp. 1-48.

Bethell, Tom, 2009, *Questioning Einstein Is Relativity Necessary*, Pueblo West, CO, Vales Lake Publishing.

Bkialnicki-Bibula, Iwo, 1996, "Photon Wave Function," *Progress in Optics* 36, pp. 245-293.

Blumschein, Eckard, 2018, "Semi-fundamental Constructs," posted in Fundamental Questions Institute.

Bohm, David, 1951/1989, *Quantum Theory*, New York, Dover Publications.

Bohm, David, 1965, *The Special Theory of Relativity*, New York, W. A. Benjamin.

Bohm, David and Hiley, Basil, 1993, *The Undivided Universe*, New York, Routledge.

Bohr, Niels, 1913, "On the Constitution of Atoms and Molecules," *Philosophical Magazine* 26, 1, republished in Dirk Ter Haar, 1967, pp. 132-159.

Bohr, Niels, 1928, "The Quantum Postulate and the Recent Development of Atomic Theory," *Nature* 121. pp. 580-590.

Bohr, Niels, 1949, "Discussion with Einstein on Epistemological Problems in Atomic Physics" in *Albert Einstein: Philosopher – Scientist* edited by Paul Schilpp, the Library of Living Philosophers Vol 7, La Salle, IL, Open Court.

Bohr, Niels, 1958, *Atomic Physics and Human Knowledge*, New York, Wiley.

Bohr, Niels, Kramers, Hans; and Slater, John, 1924, "The Quantum Theory of Radiation," *Philosophical Magazine* 47, pp. 785-802.

Born, Max, 1920/1964, *Einstein's Theory of Relativity*, New York, Dover.

Boscovich, Roger, 1758/1966, *Theory of Natural Philosophy*, translated by J. M. Child, Cambridge, MA, M.I.T. Press.

Brecher, Kenneth, 1977, "Is the Speed of Light Independent of the Velocity of the Source," *Physical Review Letters* 39, pp. 1051-1054.

Cacciani, Alessandro; Briguglio, Runa; Massa, Fabrizio; and Rapex, Paolo, 2006, "Precise Measurement of the Gravitational Red Shift," *Celestial Mechanics and Dynamical Astronomy* 95, pp. 425-437.

Cramer, John, 1986, "The Transactional Interpretation of Quantum Mechanics," *Review of Modern Physics* 58, pp. 647-688.

Cross, William and Ramsey, Norman, 1950, "The Conservation of Energy and Momentum in Compton Scattering," *Physical Review* 80, pp. 929-936.

De Climont, Jean, 2014, "Electron Beams Magnetic Field is not a Result of Electron Motion but of their Intrinsic Magnetic Moment," *General Science Journal*.

Dedekind, Richard, 1901/1963, *Essays on the Theory of Numbers*, New York, Dover.

De Hilster, David, 2011, "The Neutrino: Doomed from Inception," *Proceedings of the Natural Philosophy Alliance* 8, pp. 148-151.

De Sitter, Willem, 1913, "A Proof of the Constancy of the Velocity of Light," *Proceedings of the Royal Netherlands Academy of Arts and Sciences* 15, pp. 1297-1298.

Descartes, René, 1644/1983, *Principia Philosophiae, Principles of Philosophy*, translated by Valentine Miller and Reese Miller, Dordrecht, Netherlands, D. Reidel.

Descartes, René, 1897, *Oeuvres de Descartes*, edited by Charles Adam and Paul Tannery, Paris, Léopold Cerf.

Deutsch, David, 1997, *The Fabric of Reality*, New York, Penguin Publishing.

Dingle, Herbert, 1972, *Science at the Crossroads*, London, Martin, Brian and O'Keefe.

Dirac, P. A. M., 1930/1981, *The Principles of Quantum Mechanics*, Oxford, Oxford University Press.

Dowdye, Edward, 2007, "Time Resolved Images from the Center of the Galaxy Appear to Counter General Relativity," *Astronomical Notes* 328, pp. 186-191.

Eberhard, Philippe, 1978, "Bell's Theorem and the Different Concepts of Reality," *Il Nuovo Cimento* 46B, pp. 392-419.

Eddington, Arthur, 1928, *The Nature of the Physical World*, New York, MacMillan.

Einstein, Albert, 1905, "Zur Elektrodynamik bewegter Körper," "On the Electrodynamics of Moving Bodies," *Annalen der Physik* 17, pp. 891-921.

Einstein, Albert, 1920/2006, *Relativity The Special and General Theory*, New York, Penguin.

Einstein, Albert, 1922/1974, *The Meaning of Relativity*, Princeton, Princeton University Press.

Einstein, Albert; Podolsky, Boris; and Rosen, Nathan, 1935, "Can Quantum-Mechanical Description of Physical Reality be Considered Complete?," *Physical Review* 47, pp. 777-780.

Everett, Hugh, 1957, "Relative State Formulation of Quantum Mechanics," *Review of Modern Physics* 29, pp. 454-462.

Fei, Hong-Bing; Jost, Bradley; Popescu, Sandu; Saleh, Bahaa; and Teich, Malvin, 1997, "Entangled Two-Photon Transparency," *Physical Review Letters* 78, pp. 1679-1682.

Feynman, Richard, 1964, *The Feynman Lectures on Physics*, Reading, MA, Addison Wesley.

Feynman, Richard and Hibbs, Albert, 1965, *Quantum Mechanics and Path Integrals*, New York, McGraw Hill.

Feynman, Richard, 1985, *QED: The Strange Theory of Light and Matter*, Princeton, Princeton University Press.

Fox, John, 1965, "Evidence against Emission Theories," *American Journal of Physics* 33, pp. 1-17.

French, Robert, 2008, "Wave Particle Unity and a Physically Realist Interpretation of Light," *Physics Essays* 21, pp. 196-199.

Galstyan, T. V. and V. Dranoyan, V., 1997, "Light-Driven Molecular Motor," *Physical Review Letters* 78, pp. 2760-2763.

Gerber, Paul, 1898, "Die räumliche und zeitliche Ausbreitung der Gravitation," *Zeitschrift für Mathematik* 43, pp. 93-104. There is an English translation as "The Spatial and Temporal Propagation of Gravity" at www.alternativephysics.org/gerber/Perihelion.hum.

Gerlach, Walter, and Stern, Otto, 1922, "Der Experimentelle Nachweis der Richtungsquantelung im Magnetfeld," *Zeitschrift für Physik* 9, pp. 349-352.

Glauber, Roy, 1963, "The Quantum Theory of Optical Coherence," *Physical Review* 130, pp. 2529-2539.

Goudsmit, Samuel, and Uhlenbeck, George, 1926, "Spinning Electrons and the Structure of Spectra," *Nature* 117, 264-265.

Grangier, Philippe, Roger, Gérard, and Aspect, Alain, 1986, "Experimental Evidence for a Photon Anticorrelation Effect on a Beam Splitter: A New Light on Single-Photon Interferences," *Europhysics Letters* 1, pp. 173-179.

Greenberger, Daniel, 1983, "The Neutron Interferometer as a Device for Illustrating the Strange Behavior of Quantum Systems," *Reviews of Modern Physics* 55, pp. 875-904.

Hatch, Ronald, 1995, "Relativity and GPS," *Galilean Relativity* 6 52-57.

Heisenberg, Werner, 1927, "Über den Anschaulichen Inhalt der Quantentheoretischen Kinematik und Mechanik," *Zeitschrift für Physik* 43, pp. 172-198.

Heisenberg, Werner, 1962, *Physics and Philosophy*, New York, Harper.

Helmholtz, Hermann, 1858, "Über Integrale der hydrodynamischen Gleichungen welche den Wirbelbewegungen entsprechen," *Crelle's Journal für Mathematik* 55, pp. 4-55. English Translation "On Integrals of the Hydrodynamical Equations, which express Vortex-motion," *Philosophical Magazine and Journal of Science* 33 (1867), pp. 485-511.

Herzberg, Gerhard, 1937/1944, *Atomic Spectra and Atomic Structure*, New York, Dover Publications.

Herzog, Thomas, Kwiat, Paul, Weinfurter, Harald, and Zeilinger, Anton, 1995, "Complementarity and the Quantum Eraser," *Physical Review Letters* 75, pp. 3034-3037.

Hossenfelder, Sabine, 2018, *Lost in Math – How Beauty Leads Physics Astray*, New York, Basic Books.

Hübel, Hannes; Vanner, Michael; Lederer, Thomas; Blauensteiner, Bibiane; Lorünser, Thomas; Poppe, Andreas, and Zeilinger, Anton, 2007, "High-Fidelity Transmission of Polarization Encoded Qubits from an Entangled Source over 100 km of Fiber," *Optics Express* 15, 7853-7862.

Jackson, John David, 1962/1988, *Classical Electrodynamics*, Third Edition, New York, John Wiley.

Jaques, Vincent, Wu, E., Grosshans, Frederic, Troussart, Francois, Grangier, Philippe, Aspect, Alain, and Roch, Jean-Francois, 2007, "Experimental Realization of Wheeler's Delayed-Choice Gedanken Experiment," *Science* 315, pp. 966-968.

Jones, Robert Clark, 1941, "A New Calculus for the Treatment of Optical Systems," *Journal of the Optical Society of America* 31, pp. 488-493.

Kaufman, Richard and French, Robert, 2025, "Reconstructing Simultaneity Using Doppler-Based Inference: A Frame-Specific Convention within Special Relativity," *Physics Essays* 38.

Kant, Immanuel, 1786/2004, *Metaphysical Foundations of Natural Science*, Cambridge, Cambridge University Press.

Kim, Yoon-Ho and Shih, Yanhua, 1999, "Experimental Realization of Popper's Experiment: Violation of the Uncertainty Principle?" *Foundations of Physics* 29, pp. 1849-1861.

Kim, Yoon-Ho, Yu, R., Kulik, S, Shih, Y and Scully, M., 2000, "A Delayed Choice Quantum Eraser," *Physical Review Letters* 84, pp. 1-5.

Klyshko, David, 1988, "Combined EPR and Two-slit Experiments: Interference of Advanced Waves," *Physics Letters A* 132, pp. 299-304.

Kochen, Simon and Specker, Ernst, 1967, "The Problem of Hidden Variables in Quantum Mechanics," *Journal of Mathematics and Mechanics* 17, pp. 59-87.

Kwiat, Paul; Mattle, Klaus; Weinfurter, Harald; Zeilinger, Anton; Sergienko, Alexander, and Shih, Yanhua, 1994, "New High-Intensity Source of Polarization-Entangled Photon Pair," *Physical Review Letters* 75, pp. 4337-4341.

Lewis, Gilbert, 1916, "The Atom and the Molecule," *Journal of the American Chemical Society* 38, pp. 762-785.

Locke, John, 1690/1959, *An Essay Concerning Human Understanding*, New York, Dover Publications.

Lorentz, Hendrik, 1900, "Considerations on Gravitation," *Proceedings of the Royal Netherlands Academy of Arts and Sciences* 2, pp. 559-574.

Lorentz Hendrik, 1904, "Electromagnetic Phenomena in a System Moving with any Velocity less than that of Light," *Proceedings of the Academy of Sciences of Amsterdam* 6, pp. 809-831; reprinted in *The Principle of Relativity*, New York, Dover, pp. 11-34.

Lorentz, Hendrik, 1916, *The Theory of Electrons*, Second Edition, Leipzig, B. G. Teubner.

Lucas, Charles, 2013, *The Universal Force*, Vol. 1, commonsensescience.org.

Lucretius, c. 60 B. C. E./1951, *De Rerum Natura* (On the Nature of Things), translated by Ronald Latham, Harmondsworth, England, Penguin.

Lunenburg, Rudolf, 1947, *Mathematical Analysis of Binocular Vision*, Princeton, Princeton University Press.

Mandl, Franz and Shaw, Graham, 1993, *Quantum Field Theory Revised Edition*, New York, John Wiley.

Maudlin, Timothy, 2002, *Quantum Non-Locality and Relativity*, Oxford, Blackwell.

Maxwell, James, 1861, "On Physical Lines of Force," *Philosophical Magazine* 21, pp. 161-175.

Maxwell, James, 1873/1954, *A Treatise on Electricity and Magnetism*, New York, Dover Publications.

Menger, Karl, 1940, "Topology without Points," *Rice Institute Pamphlet* 27, pp. 80-107.

Menger, Karl, 1943, "What is Dimension?," *American Mathematical Monthly* 50, pp. 2-7.

Messiah, Albert, 1958/1999, *Quantum Mechanics*, New York, Dover Publications.

Michelson, Albert, 1913, "Effect of Reflection from a Moving Mirror on the Velocity of Light," *Astrophysical Journal* 37, 190-193.

Mizobuchi, Yutaka and Ohtaké, Yoshiyuki, 1992, "An Experiment to Throw More Light on Light," *Physics Letters A* 168, pp. 1-5.

Mojorana, Quirino, 1918, "On the Second Postulate of the Theory of Relativity; Experimental Demonstration of the Constancy of Light Reflected from a Moving Mirror" *Philosophical Magazine* 25, 163-174.

Muraviev, A., Konnov, D. and Vodopyanov, 2020, "Broadband High-Resolution Molecular Spectroscopy with Interleaved Mid-infrared Frequency Combs," *Nature/com/Scientific Reports*.

Muthukrishnan, Ashok; Agarwa, Girish; and Scully, Marlan, 2004, "Inducing Disallowed Two-Atom Transitions with Temporally Entangled Photons," *Physical Review Letters* 93, pp. 093002-1-4.

Nelson, Edward, 1985, *Quantum Fluctuations*, Princeton, Princeton University Press.

Newton, Isaac, 1687/1934, *Philosophia Naturalis Principia Mathematica, Mathematical Principles of Natural Philosophy*, Berkeley, University of California Press.

Omnès, Roland, 1994, *The Interpretation of Quantum Mechanics*, Princeton, Princeton University Press.

Paul, Harry, 1986, "Interference between Independent Photons," *Reviews of Modern Physics* 58, pp. 209-231.

Pauling, Linus, 1939/1960, *The Nature of the Chemical Bond*, Ithaca, Cornell University Press.

Peng, Tao, Chen, Hui, Shih, Yanhua, and Scully, Marlan, 2014, "Delayed-Choice Quantum Eraser with Thermal Light," *Physical Review Letters* 112, pp. 180401-1-5.

Peng, Tao; Simon, Jason; Chen, Hui; French, Robert and Shih, Yanhua, 2015, "Popper's Experiment with Chaotic-Thermal Light," *Europhysics Letters* 109 pp. 14003-1-6.

Penrose, Roger, 1991, *The Emperor's New Mind*, New York, Penguin Books

Penrose, Roger, 2004, *The Road to Reality*, New York, Random House.

Poincaré, Henri, 1889, *Théorie Mathématique de la Lumière, Mathematical Theory of Light*, Paris, Georges Carré.

Poincaré, Henri, 1906, "Sur la Dynamique de l'Électron," *Rendiconti del Circolo Matematico di Palermo* 21, pp. 129-176.

Poincaré, Henri, 1912/1963, *Dernières Pensées*, New York, Dover Publications.

Popper, Karl, 1934/1959, *Logik der Forschung, The Logic of Scientific Discovery*, New York, Harper.

Popper, Karl, 1982, *Quantum Theory and the Schism in Physics from the Postscript to the Logic of Scientific Discovery*, London, Routledge.

Pound, Robert and Rebka, Glen, 1960, "Apparent Weight of Photons," *Physical Review Letters* 4, pp. 337-341.

Quine, Willard, 1953, "On What There Is" in *From a logical point of view*, pp. 1-19, Cambridge, Harvard University Press.

Renninger, Mauritius, 1960, "Messungen ohne Störung Messobjekts" "Observations without Destroying the Measured Object," *Zeitschrift für Physik* 158 (1960), pp. 417-421.

Ritz, Walther, 1908, "Recherches critique sur l'Electrodynamique Générale," "Critical Researches on General Electrodynamics," *Annales de Chimie et de Physique* 13, pp. 145-275. There is an English translation at www.datasync.com/~rsf1/crit/1908a.htm

Russell, Bertrand, 1925/2009, *ABC of Relativity*, London, Routledge.

Roychoudhuri, Chandrasekhar, 2010, "Various Ambiguities in Generating and Reconstructing Laser Pulse Parameters," in *Laser Pulse Phenomena and Applications*, edited by F. J. Duarte, InTech, pp. 117-142.

Roychoudhuri, Chandrasekhar and Continuum, Femto, 2005, "If Superposed Light Beams do not Re-distribute each others Energy in the Absence of Detectors (Material Dipoles), Can an Indivisible Single Photon Interfere by/with Itself?," *Proceedings of SPIE*, vol. 5866, pp. 26-35.

Santilli, Ruggero, 1981, "An Intriguing Legacy of Einstein, Fermi, Jorday, and Others: The Possible Invalidation of Quark Conjectures," *Foundations of Physics* 11, pp. 383-472.

Schrödinger, Erwin, 1927/1982, *Collected Papers on Wave Mechanics*, New York, Chelsea Publishing Company.

Scully, Marland and Drühl, Kai, 1982, "Quantum Eraser: A Proposed Photon Correlation Experiment Concerning Observation and "Delayed Choice" in Quantum Mechanics" *Physical Review A* 25, pp. 2208-2212.

Selleri, Franco, 1996 "Noninvariant One-way Velocity of Light," *Foundations of physics* 26, pp. 641-664.

Shih, Yanhua, 1999, "Two-photon Entanglement and Quantum Reality" in *Advances in Atomic, Molecular, and Optical Physics*, Vol. 41, edited by Benjamin Bederson and Herbert Walther, Academic Press, Cambridge.

Shih, Yanhua; Sergienko, Alexander; Rubn, Morton; Kiess, T., and Alley, Carroll, 1994, "Two-photon Entanglement in type-II parametric down-conversion," *Physical Review A* 50, pp. 23-28.

Ter Haar, Dirk, 1967, *The Old Quantum Theory*, Oxford, Pergamon Press.

Thomson, J. J., 1881, "On the Electric and Magnetic Effects Produced by the Motion of Electrified Bodies," *Philosophical Magazine* 11, pp. 229-149.

Thomson, William, 1867, "On Vortex Atoms," *Philosophical Magazine* 34, pp. 15-24.

Tolman, Richard, 1910, "The Second Postulate of Relativity," *Physical Review* 31, pp. 28-40.

Tolman, Richard, 1912, "Some Emission Theories of Light," *Physical Review* 35, pp. 136-143.

Von Neumann, John, 1932/1955, *Mathematische Grundlagen der Quantenmechanik, Mathematical Foundations of Quantum Mechanics*, Princeton, Princeton University Press.

Wheeler, John, 1978, "The "Past" and the "Delayed-Choice" Double-slit Experiment," in A. R. Marlow, *Mathematical Foundations of Quantum Theory*, New York, Academic Press, pp. 9-48.

Whitehead, Alfred North, 1925, *Science and the Modern World*, New York, Macmillan.

Wiener, Otto, 1890, "Stehende Lichtwellen und die Schwingungsrichtung Polarisiten Lichtes," *Annalen der Physik* 34, pp. 203-243,

Wigner, Eugene, 1962, "Remarks on the Mind-Body Question" in *The Scientist Speculates*, edited by I. J. Good; New York, Basic Books, pp. 284-302.

Yilmaz, Hüseyin, 2010, "Lorentz Transformations and Wave Particle Unity," *Physics Essays* 23, pp. 334-336.

Yin, Juan et al., 2017, Satellite-based Entanglement Distribution over 1200 Kilometers, *Science* 356, pp. 1140-1144.

Zöllner, Johann, 1883, *Über die Natur der Cometen*, Leipzig., L. Staackmann.