

QUANTUM COLD-CASE MYSTERIES REVISITED

By Lane Davis

seattle.truth@gmail.com

©2010*

Based On The Works And Teachings Of Frank Znidarsic

Abstract

The current understanding of our inherently quantized world has failed to resolve many of the paradoxes first encountered by the pioneers of quantum theory. There has been zero progress toward the understanding of the transitional quantum state. There has been zero progress toward the understanding of why the electron does not continue to radiate energy once it has reached the ground state orbital. There has been zero progress toward postulating a hypothesis for why Planck's constant arises in virtually all quantum mechanical equations, or why the fine structure constant arises in comparing various intrinsic lengths. Just as mid-20th century physicists discovered that elementary particles were not necessarily elementary, new insights have given rise to formulas which ordain that some of the fundamental constants are not necessarily fundamental. After one hundred years of abject mystery, the first true look at the underlying causes for quantum nature is beginning to emerge.

Introduction

The term quantum physics is a misnomer. The correct phrase is quantum *mechanics*; as *physics* is the study of causation. Newton was the first to find mathematical relationships in the motions of celestial bodies, and used that math to express the relationships between various forces and energies; making him the first human to prove through impeccable logic that the objects in the heavens were not at all magical. In finding these causations for the motions in the heavens and providing mathematical proofs to back him up, the study of physics was born. A new technological golden age for mankind would ensue.

But all was not quaint in the laboratories of yesteryear. By the end of the 19th century, many unexplainable phenomena were discovered; poking holes in Newtonian physics as if the experimental scientists were dueling with the theoretical physicists. The photo-electric effect was a mystery. Blackbody radiation was nonsensical. Spectral emissions were unexplainable. Very few physicists truly cared though, as the Newtonian regime could still calculate most of the practical applications of physics with adequate accuracy.

However, the world of academia was soon turned on its head when Max Planck discovered a numerical relationship between spectral emissions over 100 years ago. The relationships between the spectral lines were multiples of a specific number; a constant which would soon prove to be successfully interjected ad-hoc into hundreds of physics equations, and which would eventually produce an immense number of accurate predictions.

The first successful adaptations of Planck's constant were in Einstein's photoelectric equations, which described the energy contained in a photon. Soon, the constant would be interjected into equations of all types to render the basic aspects of quantum theory, including the semi-intuitive, yet flawed inception of Bohr's solar-system-esque model of the atom.

Bohr's model was revolutionary, in that it vehemently defied all aspects of Newtonian physics. Where Newtonian physics told us that perpetual motion was *not* possible, quantum mechanics told us that not only was it *possible*; but that the most abundant commodity in the universe (i.e., the atom) would always function with perpetual motion as long as it is left undisturbed. Newtonian physics dictated that an electron should constantly emit electromagnetic radiation and continue to lose energy until eventually smacking into the nucleus. The experimental physicists had proved that this was not the case, and were able to organize a successful coup d'état in the scientific community. The theoretical physicists were eventually forced to admit that atoms *do* actually exist, *despite* seeming to break all the known laws of physics. And instead of sentencing all the atoms in the universe to lengthy terms in a quantum-prison for their impervious disenfranchisement from commonly known universal laws, the then-currently trending theories were tossed out the window, and Newtonian physics was relegated to the status of *classical* physics.

However, the real irony in this theoretical demotion is that a new successor was never truly crowned. There was no commencement ceremony. The sole successor in sight explained only what an atom

was *doing*, but never even came close to touching the explanations for *how* or *why* atoms were doing what they do. New theories were never written in academic journals to explain the *causations* for the inherent quantum mechanical nature that was observed, or why the positive and negative charges within an atom never radiated energy past a certain point and never crashed into each other. Only later was it inferred that the uncertainty principle worked like a form of magic to keep them apart; to keep the charges from emitting electromagnetic energy to an energy level lower than the ground state.

Nevertheless, the new kid on the block — dubbed quantum mechanics — was reluctantly given the title of king, even though it was not immediately appointed lordship over its scientific serfs, and never truly acquired its kingsly capstone. The great minds behind the quantum revolution had a serious problem with the fact that the quantum mechanical models could not explain the underlying causations for *why* atoms work. This troubled them deeply. The inherent paradoxes ran rampant, and pushed some early quantum pioneers near insanity.

Even before the legends of early 20th century physics bowed out from this universe, they were already being lambasted by the new generation of physicists as *senile old men* who were clutching onto their ignorant preconceptions; ideals wherein they assumed that everything in the universe should have an underlying cause.

Logical reasoning was replaced by a form of quantum-indoctrination. Yes, it was obvious that everything atomic *was* in reality quantized, but within a single generation, the world of physics had begun to completely ignore the fact that not a single PhD could explain the physical causations behind *why* an atom is stable, *why* an electron doesn't continue to radiate energy below the ground state orbital and eventually crash into the nucleus, *or why* an electron doesn't drop straight from a top energy level to the ground state energy level and emit a high energy photon in one fell swoop. Furthermore, the fine structure constant remained obscured; lost somewhere in Feynman's dark alley. Much like a high profile murder case which wasn't solved within the first few years, all of these mysteries began to gather dust, and soon they became forgotten cold-case quantum files. The fact that a true physical theory for quantum causation was never found to replace classical physics would prove to be conveniently ignored by the upper echelons of science; for admitting that one doesn't know much about the most abundant substance in the universe does not look good on one's résumé.

Luckily, there is no statute of limitation for unsolved mysteries in physics. Although no professional physicists are concerned with the fact that they can't explain the physics of an atom, and instead they confer that they can only explain what an atom actually *does*, that doesn't mean that everyone in the world has continued to feign ignorance. There exists a small underground movement which attempts to reconcile these oft-shelved-mysteries, which has coalesced on the world's first never-ending Copenhagen convention that is open to all comers; the World Wide Web. This unassuming movement, founded by the theories and mathematical relationships discovered by Frank Znidarsic, attempts to prove through the *only* language native to the universe itself — mathematics — that the quantum nature of atoms is not at all magical, and it doesn't break the laws of physics. Finally, a reconciliation of classical and quantum systems can begin to ensue. In this paper it will be shown that Newton *still* holds the crown in many regards, after 100 years under a quantum-quasi-rule by fiat.

Fine Structure: The First Clue

The key to unlocking some of the greatest paradoxes in the quantum world rests on admitting what the world of quantum academia does not know. It seems to know everything about the static atomic state, but will freely admit that absolutely nothing is known about the transitional atomic state. The only reason Einstein was able to advance his understanding of $E = mc^2$ and disseminate it to the world is because he was the first human in history to assume that the speed of light is constant, finite, and relativistic. Frank Znidarsic was the first human to assume that the speed of quantum transition is not only constant and finite, but also far less than the speed of light; rendering the relativistic paradigm negligible for the transitional regime.

The fact that a retarded speed of transition initially seems to break all known laws of physics should be put aside for a moment, although I advise the reader to stay critical of every argument presented. Just remember, the fact that an atom exists without the electron crashing into the nucleus seemed to break all the laws of physics only a century ago, but that didn't stop QED from it's ascent to power. Of course, these counter-intuitive mysteries will be explained as the paper progresses. However, through a mathematical examination of the speed of transition, one will see that this number cannot be an accident; it produces far too many accurate calculations for quantum aspects of the utmost importance, and explains some of the greatest quantum conundrums. One, two, or possibly three of these formulations could be explained as coincidence. Any more than that pushes

the limit of credulity past where one could believe – in good faith – that it is all an accident. Once the descent down the rabbit hole has commenced, the conclusion is distinct; it must be the hole belonging to the Energizer© bunny. It just keeps going and going™.

The first assumption that was made to find the speed of transition (c_t) was that it was a ratio of the fine structure constant (α) and the speed of light (c). Eventually the number of 1,093,846 m/s (c_t) was reached as the expression of:

$$c_t = \frac{c\alpha}{2}$$

This new application of the fine structure constant gave new meaning to its dimensionless property (meaning that it has no units associated with it; it is simply a number). The dimensionless aspect emerged because the units cancel out; because it is a ratio of two speeds. Of course, this relationship is inherently meaningless in and of itself unless we can make accurate calculations and predictions with this new found c_t .

The Energy Contained In A Photon

One of the greatest mysteries in quantum mechanics is why the energy contained in photon is proportional to its frequency. The energy in a *classical* wave is a product of the frequency squared multiplied by the amplitude squared. But the energy contained in an electromagnetic wave (quanta of light; A.K.A. the photon) is considered to be only a function of the frequency. Assuming c_t to be 1,093,846 m/s reconciles this dilemma. For future reference, there is an index of all the values for the nomenclature used in this paper on the last page.

The world of physics knows that the speed of light (c) is equal to the frequency (f) of the photon times the wavelength (λ): $c = f\lambda$. Because the frequency of the emitted wave always matches the frequency of the wave emitter in a classical system, the logical assumption was made that the speed of transition (c_t) is equal to the frequency of the photon in free space (traveling at the speed c); which would remain constant throughout the act of transition, multiplied by the compacted transitional wavelength (λ_t). This is expressed mathematically as:

$$c_t = f\lambda_t$$

Alongside the standard formulation for speed of a photon in vacuo, it looks very familiar.

$$c = f\lambda$$

Next, let us solve for the transitional wavelength (λ_t) by dividing both sides of the above equation ($c_t = f\lambda_t$) by the frequency (f). The result is:

Equation 1:

$$\lambda_t = \frac{c_t}{f}$$

This compressed wavelength of the transitional photon (λ_t) is analogous to the wavelength of a tsunami as it scrunches up (and simultaneously increases its amplitude) when it is slowed down by entering shallower waters as it nears shore. In the middle of the ocean, the wave's amplitude; i.e., the height of the wave, is very small. As it is slowed down by the shallower waters, the wavelength is shortened tremendously. This is what is happening to the transitional photon as it is slowed down to the speed of transition (c_t).

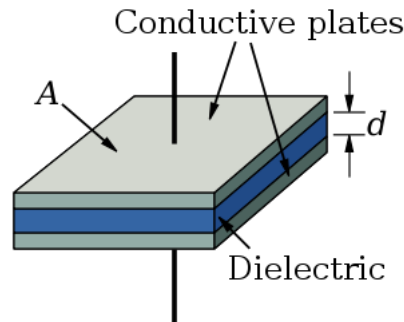
At this point, you might be wondering what the point of calculating this transitional wavelength of the photon is, seeing as it's impossible for us to ever measure it experimentally. But watch what solving for this transitional wavelength (λ_t) allows us to do:

First, let me introduce the next equation; an equation used to calculate the capacitance (C) of a two-plate capacitor, found in *any* electrical engineering textbook.

Equation 2:

$$C = \frac{\varepsilon_0 A}{d}$$

Figure 1:



For clarification, ϵ_o represents the electrical permittivity of free space; i.e., the measure of how much of an electric charge space is able to hold, A represents the area of each plate, and d stands for the distance between the two plates.

And of course, area equals length times width:

$$A = LW$$

A photon in free space is known to have a wavelength, but standard science does not consider its *wavewidth* to be relevant. The width of the wave of the uncollapsed photon is known to exist though, as the *two-slit experiment* proves. This experiment shows that an uncollapsed photon is able to be two places at once; the wave function travels simultaneously through both slits, then interferes with itself and spreads out, eventually crashing onto the backdrop and seemingly randomly showing up at a single spot. Furthermore, what is dubbed the *single-slit experiment*, wherein a beam of photons is passed through a very narrow single slit, which in turn causes the laser-dot to expand its width on the backdrop in a counter-intuitive manner, is also evidence that a photon has a wave-width.

The only question that remains is whether the wavefunction of the photon still retains any of its width during its collapsed state incurred by the quantum transition that it undergoes during atomic absorption. However, we do know that even though the particle aspect of a photon might be considered to be a point particle, the wavefunction never truly disappears. The only question is how small the wavefunction gets, even when it is collapsed.

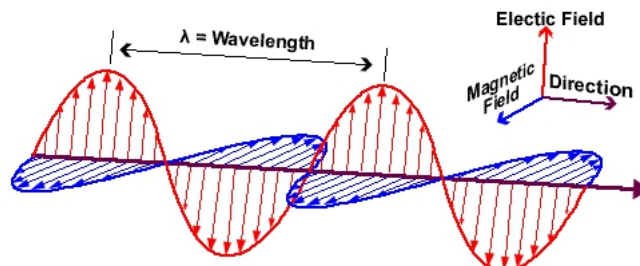
Inferring that this compacted wave-function-collapsed photon has a small width by using the logic listed above, we are almost ready to move on to the next set of equations. But what we are about to do with these equations requires a major paradigm shift in order to

understand. We are going to look at the photon in terms of the geometry of a three-dimensional capacitor. One might wonder how a *3D capacitor* has any relevance to a photon. Well, remember, the two-slit experiment and Heisenberg's uncertainty principle (which comes into play during the one-slit experiment) show that a photon's wavefunction must have a width. And we know it has a length.

But what *makes* a capacitor a capacitor? It's a positive and a negative charge separated by a distance, and the void in between the positive and negative charge differentials is filled by a dielectric; i.e., an insulator like glass, as was shown in Fig. (1).

And what exactly is a photon?

Figure 2:



A photon is an oscillating positive and negative charge. In the above figure, on the *y-axis* (vertical) above the zero line of the photon represents a positive charge. When the photon dips below the *y-axis*, its charge becomes negative. The moving electric field creates the magnetic field, and vice-verse, and that is why it is continually alternating. And one might notice, the photon is similar to a sine/cosine wave on a graph:

Figure 3:

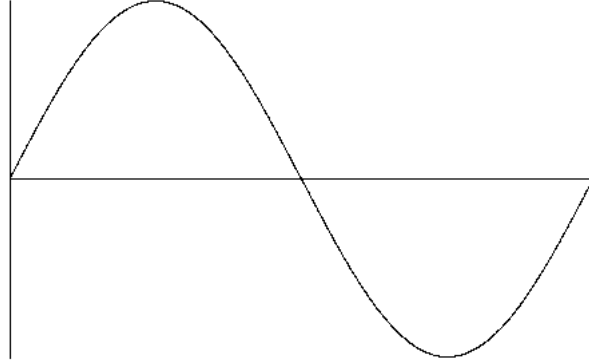


Fig. (3) represents a single wavelength on the horizontal axis. The vertical axis represents the positive and negative aspects of the wave. So, what is the distance between the positive and negative sides? It's not a whole wavelength, but *half* a wavelength. This means that when modeling the photon as a capacitor, the distance between the positive and negative sides will again be *half* a wavelength. In lieu of this knowledge, let us make the distance in the capacitor formula equal to half the transitional wavelength (λ_t) in the capacitance formula; Eq. (2).

Equation 3:

$$C = \frac{\varepsilon_0 A}{.5\lambda_t}$$

As we know, area equals length times width. And we also know that the photon will collapse to the transitional wavelength, which was calculated in Eq. (1) by setting the transitional wavelength (λ_t) equal to the speed of transition divided by the frequency (which stays the same during transition and in free space). Let us now assume for a minute that the width of the collapsed photon during the transitional state becomes the same as the wavelength. Yes, these are big jumps of logic, but bear with me and watch what happens. Because length times width equals area, and we are assuming width also to be equal to the wavelength, that gives us:

Equation 4:

$$C = \frac{\varepsilon_0 \lambda_t^2}{.5\lambda_t}$$

When this equation is reduced it simply leaves us with:

Equation 5:

$$C = 2\varepsilon_0\lambda_t$$

And because $\lambda_t = \frac{c_t}{f}$, as was stated previously, we can substitute $\frac{c_t}{f}$ in for λ_t , leaving us with:

Equation 6:

$$C = \frac{2\varepsilon_0c_t}{f}$$

This brings us to the next equation which can be found in any electrical engineering textbook:

Equation 7:

$$E = \frac{Q^2}{2C}$$

This equation is simply stating that the energy of the system (E) is equal to the charge of the system (Q) squared over two times the capacitance (C). Since we already solved for the capacitance of the transitional photon in Eq. (6), we can substitute $\frac{2\varepsilon_0c_t}{f}$ into Eq. (7) for the capacitance, yielding:

Equation 8:

$$E = \frac{Q^2f}{4\varepsilon_0c_t}$$

Its important to note that the charge of a photon is equal to the elementary charge ($1.60217646 \times 10^{-19}$ coulombs); the charge intrinsic to the proton an electron as well, which is denoted as e . Substituting e in for Q in the equation and rearranging the f in the numerator off to the side leaves us with:

Equation 9:

$$E = \left[\frac{e^2}{4\varepsilon_0c_t} \right] f$$

It just so happens that the terms within the brackets equals Planck's constant, which is denoted as h . Substituting h into the equation leaves us with $E = hf$, which is Einstein's photo-electric equation; the first-ever correct application of Planck's constant in history, for which Einstein was awarded the Nobel prize in 1921. What I have just shown with these equations is that in reality, Planck's constant, which is the fundamental increment of action in the quantum world, is actually an *aggregate* constant. Yes, it provides us with an accuracy in calculations that is uncanny, but in reality Planck's constant is completely empirical. This means that it is derived purely from experimentation alone. We have been able to formulate it in a few simple steps from some of its actual constituent constants: e , ϵ_0 , and c_t . I specifically said *some* of it's constituent constants, because in reality there are more ways to formulate Planck's constant from different sets of fundamental constants, as will be shown soon. The form of Planck's constant in Heisenberg's Uncertainty Principle arises out of the formulation of Planck's constant in Eq. (9), because the frequency of the photon (hence the length of it's wavepacket) and the energy of the photon are related by the conversion factor of $\frac{e^2}{4\epsilon_0 c_t} = h$.

But most importantly, Eq. (9) solves quantum cold-case file numero uno; *why* the energy level of a photon is proportional to its frequency. As the frequency gets higher, the transitional wavelength (λ_t) also gets smaller in order to maintain equivalence with the speed of transition (c_t). As λ_t gets smaller, the volume encompassed by the capacitor gets smaller. The simplest way to think of the capacitor would be as a box, and inside that box we have a charge. The voltage of said capacitor is most easily explained as the *pressure of the charge*. Because in the case of the photon, the charge always stays the same (e), as the volume of the box that you put the charge into gets smaller (which is caused by the decreasing wavelength/width), the voltage increases. In an electrical system, voltage can be thought of as the amplitude. So in this sense, the energy of a photon is coming from the voltage of the transitional state, which is the amplitude. As the frequency gets higher, the transitional wavelength gets smaller, making a smaller volume for the charge to be in, increasing the pressure of the charge, which means that the voltage is increased. This is why the energy of a photon is related to the frequency, and also goes to show that the wavefunction of the photon never truly vanishes. In fact, the energy of the photon is purely dictated by its collapsed wavefunction aspects. The only information that is not rendered absolutely by the wavefunction is the *exact* location of the photon's particle aspect, although the location becomes somewhat localized. Eq. (9) gives us the energy level for any frequency of photon imaginable by formulation through simple textbook equations, and simultaneously produces

Planck's constant as an aggregate. If you are still skeptical that this could all be a coincidence at this point; well, that is a healthy stance to take, but you won't be disappointed as the paper progresses.

The Speed Of Mechanical Waves In The Nucleus

When an atom emits or absorbs as a photon, 100% of the potential energy from the two interacting charges (the positive proton and the negative electron) gets converted into electromagnetic radiation. To explain what I mean by potential energy of the charges, imagine being up at the top floor of a sky scraper. The elevation that is experienced gives one a potential energy. The potential energy can be converted into kinetic energy (motion), by base-jumping off of the roof. This same sort of situation applies to a proton/electron. The farther apart they are, the higher their energy state is (but the energy is just a potential). When they drop down to be closer to each other, their energy state drops as well, similar to the potential energy of the base-jumper as they plummet towards the ground. But the difference between the charges in the atom and a base-jumper is that instead of converting the potential energy into kinetic energy as the distance of separation decreases, it is converted into *electromagnetic* energy, which radiates away from the atom.

But in the realm of known laws of electrical engineering, 100% energy transfers require a 100% impedance match. Impedance is a measure of resistance, and can be applied not only to electrical systems, but to mechanical ones as well (such as a clutch in a manual transmission car). Impedance matching is required to transfer electrical power from power stations to your house. The job of the massive transformers that are seen near power stations (or the smaller ones going from the power line to your house) is to match the impedance of one set of lines to the next set of lines. If this doesn't happen, the energy from the lower impedance line will bounce off the line with the higher impedance and the energy will not transfer efficiently. In order to get *all* of the energy to transfer, the impedances must be matched 100%. Because we know that 100% of the electric energy transfers from the electrostatic potential energy to the energy contained in the photon that it emits (ignoring the *very* minute loss from entropy), we must also assume that there is a 100% impedance matching going on within the state of transition. That is, unless we want to rewrite *all* the known laws of electrical engineering.

So, if the line with the highest impedance is the limiting factor in transferring electricity efficiently through different types of power lines, what is the limiting factor; the aspect with highest impedance within an atom? Well, remember that when a photon is created, the energy comes from the potential energy from *both* of the charges, not simply the electron by itself. The nucleus is what becomes the limiting factor, and in a strange way. The nucleus is made up of protons and neutrons, but it is counter-intuitive in that the electric fields of the protons are expelled to the edge of the nucleus, similar to how the magnetic field of a superconductor is expelled (i.e., the Meisner effect). Yes, it sounds strange, I know, but this is proven by the fact that when you add a neutron to a nucleus (like deuterium, which is an isotope of hydrogen that has one proton and one neutron) it lowers the Coulombic barrier of the entire nucleus. And by Coulombic barrier, I mean the energy at which the repulsive electrostatic force experienced by two atoms becomes equal to the attractive strong nuclear force. When the Coulombic energy level is construed as a distance of separation; right at this point, the two protons feel no net force, but if moved *any* closer together, the strong force will suck the two protons together (which is the process of fusion, in a nutshell). If it wasn't true that the electrical field was expelled to the edge of the nucleus, then fusion with deuterium nuclei would not take high energies at all, because one side of the deuterium atom would have the same Coulombic barrier as regular hydrogen (the side with the proton on it), and the other side of the atom (the side with the neutron on it, which gives no repulsive electrostatic force) would have a very low Coulombic energy barrier. This is not the case. The energy required to fuse two deuterium atoms is indeed lower than two standard hydrogen isotopes, but the lowered Coulombic energy barrier is experienced by the entire nucleus, so we know that the electrostatic forces are expelled to the edge of the nucleus. Another example of this charge expulsion would be a nucleus halo.

The speed at which the electrostatic force is expelled from the nucleus becomes the limiting factor. The electrostatic force is expelled through mechanical waves. This makes electrostatic transfer in the nucleus analogous to how heat is transferred within materials. Heat is a form of electromagnetic radiation, but heat within a material is dissipated through phonons; mechanical vibrations that form waves of mechanical motion. Phonon research has been key to developing microprocessors which can dissipate heat efficiently. The electrical field of the nucleons is dissipated to the edge of the nucleus in the same fashion; through mechanical vibrations.

Now, allow me to retouch on how we initially calculated the value of c_t . Remember, it was from the equation $c_t = \frac{c\alpha}{2}$. The speed of transition equals the speed of light times the fine structure constant over two. And after formulating the energy level of all photons, I stated that Planck's constant only arises as an aggregate, and that in and of itself it is not a basic constant; as is currently thought by the world of academia. The problem here is that the fine structure constant is actually defined in terms of Planck's constant. So if we are defining c_t in terms of the fine structure constant (α), and the fine structure constant is in turn defined by Planck's constant (h), then we are simply defining c_t in terms of h . This is not a good scenario for one who is trying to prove that Planck's constant is completely irrelevant; this is what is called circular mathematics. Luckily, there is a way out of this vicious cycle. Because it is claimed that the speed of transition becomes the limiting factor when matching the impedances of the electrical fields to transfer the energy 100% efficiently in order to emit or absorb a photon, and it is also claimed that the protons expel their electrical field through mechanical waves similar to how phonons transfer heat, then all that has to be done is to calculate the speed of mechanical wave propagation in the nucleus. If this number does indeed match what is calculated by $c_t = \frac{c\alpha}{2}$, then the conundrum will be resolved and the math will no longer be circular.

Calculating the speed of mechanical waves within a material is a science that has been around for quite some time. It is used the most in non-destructive materials testing, known as NDT. It can be done by using Hooke's law, and by knowing the masses and forces within the medium. Hooke's law is expressed mathematically as $F = -kx$, which simply means that the force (F) experienced by a spring is equal to the spring constant (k) times the distance that the spring has been displaced (x). The negative sign just means that the force of the spring pushing back is in the opposite direction of the force that is pulling/pushing it. The negative sign is an application of Newton's third law, which states: "Every action has an equal and opposite reaction".

Coulomb's equation for two interacting charges expresses the energy between them.

Equation 10:

$$E = \frac{Q^2}{4\pi\epsilon_0} \left(\frac{1}{x} \right)$$

The Q stands for charge, and x is for the distance apart, and ϵ_0 is the permittivity of free space.

The standard equation for the energy contained in a spring is:

Equation 11:

$$E = \frac{1}{2}kx^2$$

So, because two protons are interacting charges and we want to analyze them in terms of Hooke's law, we can look at the two interacting charges as if they were a spring. Because potential energy is related to force, the energy contained in a spring is dependent upon the same two variables; the spring constant of that spring, and the amount of displacement that the spring is experiencing from its equilibrium position.

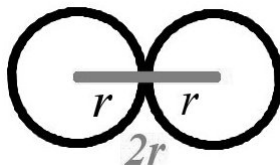
The spring constant (k) is the only part of the equation that is a little tricky to grasp. It can be formulated in a variety of ways, but for our application we will be defining k as the maximum force experienced by the spring divided by the maximum distance that has been displaced. Because $F = -kx$, this means that the force is always directly proportional to the displacement in a linear fashion. So if we are solving for the spring constant in terms of F_{max} , we have to also use the maximum displacement (x_{max}). This is written mathematically as:

Equation 12:

$$k = \frac{F_{max}}{x_{max}}$$

It turns out that the maximum electrostatic force experienced by two protons being pushed together is at the Coulombic barrier. This displacement is at 1.409 femtometers (r_c) and the maximum force experienced by the charges is 29.053 Newtons. If the two protons are pushed together any farther than this, they get sucked together by the overbearing attractive strong force. So, $F_{max} = 29.053\text{N}$. But inside the nucleus, the particles interact a little differently. They don't get sucked in until they touch. They get spaced out to the Fermi spacing (A.K.A. momentum spacing), which is 1.36 fm (r_n). The displacement of the charges at the Fermi spacing becomes twice the Fermi spacing ($2r_n = 2 \times 1.36\text{fm}$) because if you push two circles until they touch, the distance separating them becomes twice the radius.

Figure 4:



The next equation from physics 101 is the equation that expresses the frequency of a simple harmonic oscillator if you want to express a longitudinal wave as opposed to transverse.

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

The only thing that might confuse you about this equation is the funny looking w-looking thing; omega (ω). It represents *angular frequency*, that is to say; how many radians per second something spins. The real number of revolutions per second (e.g., the number of revolutions a tire on a car makes per second); the standard frequency (f), is simply the angular frequency divided by 2π . The angular frequency isn't something you can really measure by simply observing; it's just a mathematical formulation that is useful in doing calculations. To find the real frequency, which is what we want to find out in day-to-day life, we just have to divide ω by 2π . And it just so happens that any introductory physics textbook will also tell us that:

Equation 13:

$$\omega = \sqrt{\frac{k}{m}}$$

Therefor, by combining Eq. (13) with Eq (12), we can say that:

Equation 14:

$$f = \frac{1}{2\pi} \sqrt{\frac{k}{m}}$$

The next important thing to note is that speed equals frequency times displacement ($v = fx$). To understand this, imagine a speaker cone oscillating back and forth as it plays a solid note. The average speed at which the speaker cone is moving equals the displacement of the speaker times the frequency at which it is oscillating. So if we just calculated the frequency of the system in Eq (14), all we have to do is multiply it by the displacement, to find our speed of wave propagation. That gives us:

Equation 15:

$$v = \left(\frac{1}{2\pi}\right) \sqrt{\frac{k}{m}}(x)$$

And also because $k = \frac{F_{max}}{x_{max}}$, we can replace the k in Eq. (14) with $\frac{F_{max}}{x_{max}}$. And the distance of twice the Fermi spacing ($2r_n$) is plugged in to both of the displacements (x). As for the value for the mass (m), we will be using the average mass of the nucleons (m_n) which is 1.6737×10^{-27} kg. This yields:

Equation 16:

$$v_n = \left(\frac{1}{2\pi}\right) \sqrt{\frac{\left(\frac{F_{max}}{2r_n}\right)}{m_n}}(2r_n)$$

The velocity of wave propagation in the nucleus *is* equal to the speed calculated by $c_t = \frac{c\alpha}{2}$. They are one and the same ($v_n = c_t$).

The equation can be simplified and reduced as:

Equation 17:

$$c_t = \left(\frac{r_c}{\pi}\right) \sqrt{\frac{F_{max}}{2r_n m_n}}$$

Could this be another coincidence? Let's explore some further calculations of the basic quantum aspects to erode all skepticism.

The Orbital Radii Of Hydrogen

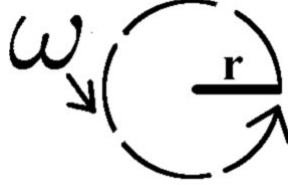
Next up: Let's formulate the orbital radii of the hydrogen atom without Planck's constant. We will start off by simply calculating the ground state radius. First, we have to set the angular velocity of the transitional electron equal to c_t .

Equation 18:

$$c_t = \omega r$$

You already know what the omega (ω) means. And the r represents radius. So what this equation is saying is that the angular speed experienced at the edge of something rotating is equivalent to how many radians per second it rotates multiplied by the radius of the object, which in this case equals c_t .

Figure 5:



And recall Eq. (13): $\omega = \sqrt{\frac{k}{m}}$. In this case the spring constant (k) is going to be the spring constant of the electron (k_{-e}) and the mass (m) is going to be the mass of the electron (m_{-e}), which is $9.109 \times 10^{-31} \text{kg}$. Therefor we can state:

Equation 19:

$$c_t = \sqrt{\frac{k_{-e}}{m_{-e}}} r_c$$

You might have noticed that I denoted the radius as the Coulombic radius (r_c) again, and this will be explained shortly. The radius (r_c) also happens to be the radius of the proton halo experienced during the third Zemach moment of the proton, but this is semi-coincidental.

You might already have suspicions for what we are going to do with k_{-e} . We are going to write the spring constant in terms of $\frac{F_{max}}{x_{max}}$ again. And when we model an orbiting system in terms of a spring, the x_{max} simply becomes the radius (r). And because we don't know what r is until we solve for it, we will call it r_x for now. Substituting $\frac{F_{max}}{r_x}$ in for k gives us:

Equation 20:

$$c_t = \sqrt{\left(\frac{F_{max}}{r_x}\right) \frac{r_c^2}{m_{-e}}} r_c$$

Solving for r_x yields:

Equation 21:

$$r_x = \frac{F_{max} r_c^2}{c_t^2 m_{-e}}$$

Doing the actual calculation gives us the ground state Bohr radius of the hydrogen atom ($a_0 = .529 \times 10^{-10} \text{ m}$). But wait, there's more; much more. If we throw a factor of n (which stands for the orbital level that we want to solve for) into Eq. (20), it gives us: $c_t = \sqrt{\left(\frac{F_{max}}{r_x}\right) n} r_c$. Solving for r_x again yields:

Equation 22:

$$r_x = n^2 \left[\frac{F_{max} r_c^2}{c_t^2 m_{-e}} \right]$$

By plugging any positive integer in for n , it gives us *every* single orbital radii for a hydrogen atom. And by plugging in a factor of Z (which equals the atomic number) into the equation, it gives us the radii of the ground state $1s$ orbital of every atom on the periodic table:

Equation 23:

$$r_x = n^2 \left[\frac{F_{max} r_c^2}{c_t^2 m_{-e} Z} \right]$$

But maybe this is just another coincidence? To show that it's not, let me introduce the electron's ugly sister; the muon. They are very similar in most regards. They are both leptons and both carry a negative charge which is equal to elementary charge $-e$. The only difference for our intents and purposes is that the muon is about 207 times more massive. But despite its increased mass, the muon is still able to be captured by a proton (like an electron is captured to form an atom) which creates an exotic type of hydrogen called *muonic hydrogen*. If you replace the mass of the electron (m_e) with the mass of the muon ($m_\mu = 1.884 \times 10^{-28} \text{kg}$) in Eq. (22) or (23), the terms in the brackets equals the ground state radius of muonic hydrogen (256 fm).

However, this is simply the *calculated* Bohr radius. It is the same number that is achieved by using the standard equation for Bohr radius involving Planck's constant: $a_0 = \frac{h}{2\pi m_e c \alpha}$, but in either sets of equations; the one I have modeled in Eq. (22) or the one currently in physics textbooks that utilizes Planck's constant, to get the number to match with experimental data we have to use an ancient formulation for two orbiting masses invented by Newton himself; the reduced mass.

By reduced mass I don't mean that the particle actually weighs any different. It is simply a formulation which allows us to look at two orbiting masses as one system. To describe what I mean, imagine the moon orbiting the earth. The earth pulls on the moon with its gravity. But the moon also pulls back on the earth with its own gravity. As the moon tugs on the earth, it actually displaces the earth from where its orbit would be if the moon were non-existent. As the earth is displaced slightly, the moon actually orbits farther away from the center of the system than it should. To calculate the correct distance from the moon to the center of the orbiting system (but not the distance to the earth itself), we need to use the reduced mass of the moon. The reduced mass is calculated by taking into account the mass of the orbiter, and the object that is orbiting around. The reduced mass is written as μ instead of m . The equation is pretty simple:

Equation 24:

$$\mu = \frac{m_1 m_2}{m_1 + m_2}$$

So to calculate the reduced mass of the muon (μ_μ), we just have to plug in the mass of the muon in for m_1 and the mass of the proton in for m_2 , which yields $1.69289416 \times 10^{-28} \text{kg}$. If we plug the reduced muon mass (μ_μ) back into Eq. (22) it gives us the true ground state

radius of muonic hydrogen, and of course if we plug a positive integer in for n it gives us *all* of the orbital radii of muonic hydrogen. And using the reduced muon mass and plugging a positive integer in for the Z in Eq. (23) it yields the orbital radii of the $1s$ orbital of every muonic atom imaginable, although none others have been produced in the laboratory at this point in time except for muonic helium.

And if you think using the reduced mass is an ad-hoc solution, or strange; don't. Actually, even when using the standard formulation for the ground state hydrogen radius (denoted as a_0), which incorporates Planck's constant and the mass of the electron; even then, to be *entirely* accurate, one is *also* supposed to use the reduced mass of the electron. But because the difference in mass between the proton and the electron is so huge (the proton is approximately 1860 times more massive than the electron), the number calculated without using the reduced electron mass is only off by about .1%, so the reduced electron mass (μ_{-e}) is usually disregarded. However, these formulations that allow us to calculate all the orbital radii of all the elements plus all the radii of the muonic elements show that this can't be attributed to chance alone. If one was so inclined to extend the model further, one could also calculate the ground state radius of positronium (an exotic "atom" made of an electron and positron) if the reduced mass is taken into account. The coincidences are sure starting to build up, aren't they?

Now let me explain what Eq. (20-23) actually mean in terms of theory. It is a little complicated. To start off; simple harmonic motion can be modeled in the form of a spring. *Any* simple harmonic motion. An orbiting mass is a form of simple harmonic oscillation if you use its position on the y-axis as the oscillation that you are modeling. Usually a simple harmonic oscillator (like a mass on a spring) can be modeled in terms of phase-space, which turns the position of the mass on the spring into what looks like an orbiting mass. The angular frequency of the orbiting mass is – of course – omega (ω). This is what gives rise to the equation $\omega = \sqrt{\frac{k}{m}}$ which we have seen already. To explain what I mean by modeling the position of a mass on a spring in terms of phase space, allow me to introduce a pair of images into the court's record:

Figure 6:

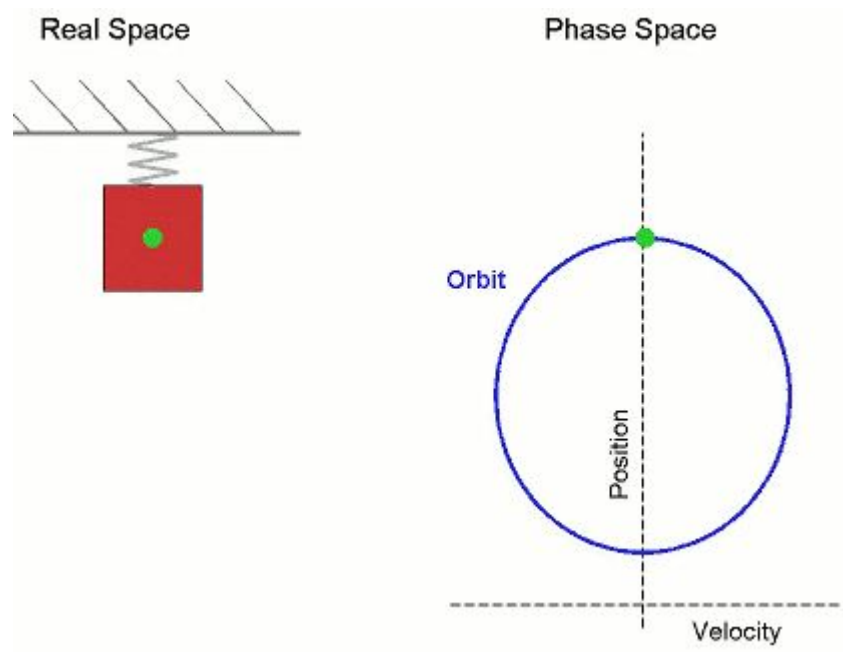
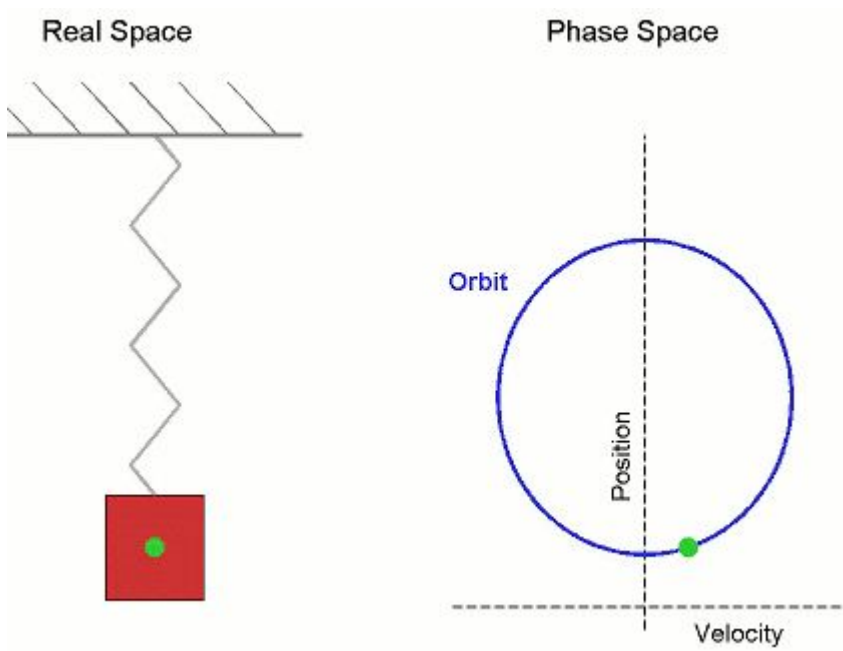


Figure 7:



Both of these figures are images of the same mass and spring but at different moments in time. In Fig. (6) the spring is *compressed* all the way. In Fig. (7) the spring is extended *nearly* all the way. The position of the green dot in the middle of the mass on the spring can be mapped into *phase space* which turns it into an orbit by expressing the dot's position on the y-axis and the dot's velocity on the x-axis. But we are mapping the orbit of a *real* orbiter; the electron. We are modeling the harmonic motion of something that's *really* orbiting, as if it were a mass on the spring. So in our model, the *phase* space becomes the *real* space, and the formulation of the spring becomes what I like to call *spring space*. The spring isn't real, but it allows us to model the simple harmonic motion of the orbiting electron in a manner that has a basis in reality.

And because the spring isn't real, the F_{max} that effects its spring constant is *not real either*. It is simply a constant that is intrinsic to all quantum orbiters; whether they be electrons or muons. The formulation of F_{max} as a mathematical construct should not be surprising, seeing as the angular momentum of the electron carried at the Bohr radius as a multiple of Planck's constant is just a mathematical construct as well. But remember, the spring constant is variable by radius, because $k = \frac{F_{max}}{x_{max}}$. So, the spring constant changes for each orbital level, which means that each orbital level effectively becomes its own "spring". And also remember, this F_{max} is not a *real* restoring force. It is simply a constant that is intrinsic to the angular frequency that is apparent in each of the orbital levels. The constant value of F_{max} in all of the static orbitals when their simple harmonic motion is modeled as a spring is what gives rise to Planck's constant, because remember; in every Bohr orbital the electron carries an integer multiple of Planck's constant. This is because the period of the harmonic motion of the electron equals the inverse of $\frac{\sqrt{\frac{k}{m}}}{2\pi}$ which is $2\pi\sqrt{\frac{m}{k}}$, or more simply; $\frac{2\pi}{\omega}$. And if you're thinking that Planck's constant must actually be *hiding* in our set of formulas somewhere as an aggregate constant (just like we saw with the photon equations), seeing as we are able to calculate the *exact same* values for the orbital radii as the standard equations, you would be right.

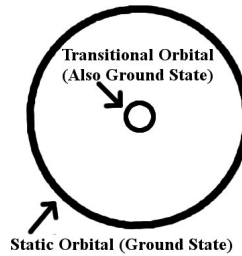
And if you think that this explanation of the F_{max} constant is rather abstract and shouldn't apply to the ground state hydrogen orbital radius (a_0); don't worry. A discovery that I made is that a_0 can be expressed in terms of c_t with no need for the F_{max} *at all*.

$$a_0 = \frac{c^2 r_c}{2c_t^2}$$

You will see why this expression for the ground state hydrogen orbital is extremely pertinent in a minute, and before that I'll explain how to find Planck's constant as an aggregate in these last equations. But first let me finish the explanation of what's going on with these equations in terms of physical theory.

We saw earlier that the wave function of the photon collapses to a transitional state wherein the wavelength gets scrunched up. Well, in reality, electrons do not orbit like the earth does around the sun. They are in what's called an electron cloud, where Schrödinger's wave equation describes the wavefunction of the electron in this cloud; which allows us to calculate the probability of where an electron is at any point in time. During the state of transition, the wavefunction of the electron also collapses into a transitional subset (where its radius starts out as integer multiples of r_c), which is directly related to the parent state in a few ways.

Figure 8:



The first way in which it is related is that the angular frequency associated with the Bohr model in the static radii is translated directly into the transitional radii. This means that it completes just as many revolutions per second once it drops down to the integer multiple; the factor of n of r_c (where n corresponds to the orbital level of the parent state). But the angular velocity changes once the electron drops down to this subset. The angular velocity of the electron/muon is *always* exactly c_t , no matter what integer multiple of the radius r_c it is experiencing. This means that at the higher radii, the angular frequency is lower to keep the velocity constant at c_t . This adds another level of impedance matching to the transitional system. The electron experiences a constant angular momentum as it switches from one orbital to another as it absorbs or emits a photon in its collapsed transitional wavefunction. Of course, in reality, the electron is not truly orbiting or spiraling down or up, just like in reality the electron isn't orbiting around the nucleus like a solar system; the description of the Bohr model. But this description allows one to *visualize* it,

and the Bohr radius has a real significance to reality, because the point at which the electron cloud's wavefunction predicts the highest probability of finding the electron directly matches the radius predicted by Bohr. In fact, nothing in terms of physical mathematics intrinsic to the universe is a coincidence at all; if it *seems* to be a coincidence, it simply means that not everything about the system is fully understood yet. As the electron switches from one multiple of r_c down to a lower one, the angular speed/angular momentum stays *constant*, so as the radius becomes tighter, the angular frequency becomes larger. This is analogous to one of those whirlpool-type money donation boxes at the mall where a coin is dropped through the slot. The angular momentum of the coin stays constant as it rolls, but as the radius that it experiences gets smaller, in order to conserve the angular momentum, the angular frequency must get higher. A simpler analogy would be an ice skater spinning with their arms and a leg stretched out. As the skater brings their outstretched limbs in closer to their body, they spin much faster because the angular momentum is conserved. The transitional state of the hydrogen atom never drops below the radius r_c , which also happens to be the radius of the proton halo at the third Zemach moment, as predicted by Miller and Cloët, and verified experimentally by Friar and Sick. With atoms other than hydrogen, the electron collapses into a wave function that is actually *inside* the nucleus. This does *not* break any laws of physics, especially the Pauli exclusion principle. Actually, even in the *static* atomic state, electrons normally appear inside the nucleus. This is part of quantum tunneling, and although the appearance of the electron within the nucleus is common, it only is captured by a neutron to flip spins into a proton when the spin orbit force is thrown out of whack by an imbalance in the spin states of the nucleons, and even then it is relatively rare. My point is that the appearance inside the nucleus – although it defies common sense – is indeed a common occurrence.

The Hidden Planck Aggregate

Now let me show you where Planck's constant is hiding in our formulation for the ground state hydrogen radius (a_0) that we made in the steps above. To do this the easy way, we're going to cheat a little. Don't worry though; our results will be produced from scratch in later equations. Let us start with the standard formulation for the ground state Bohr radius (a_0) which is from any quantum mechanical textbook:

$$a_0 = \frac{h}{2\pi m_e c \alpha}$$

And we already know that:

$$\alpha = \frac{2c_t}{c}$$

So we can plug $\frac{2c_t}{c}$ in for α :

$$a_o = \frac{h}{2\pi m_{-e} c \left(\frac{2c_t}{c} \right)}$$

The c 's cancel leaving:

$$a_o = \frac{h}{4\pi m_{-e} c_t}$$

Solve for h yielding:

$$h = 4\pi m_{-e} c_t a_o$$

And it's been shown in Eq (21) that:

$$\frac{F_{max} r_c^2}{c_t^2 m_{-e}} = a_o$$

So we can plug in $\frac{F_{max} r_c^2}{c_t^2 m_{-e}}$ for a_o in the above equation for h :

$$h = 4\pi m_{-e} c_t \left(\frac{F_{max} r_c^2}{c_t^2 m_{-e}} \right)$$

Reducing yields:

Equation 25:

$$h = \frac{4\pi F_{max} r_c^2}{c_t}$$

Now, if we substitute Eq. (25) back into the original textbook equation for $a_o = \frac{h}{2\pi m_{-e} c \alpha}$ it gives us:

Equation 26:

$$a_o = \frac{2F_{max} r_c^2}{m_{-e} c \alpha c_t}$$

And substituting $\frac{2c_t}{c}$ in for α again and reducing *brings us right back to where we started* with Eq. (21):

$$a_0 = \frac{F_{max} r_c^2}{c_t^2 m_{-e}}$$

And knowing that $h = \frac{4\pi F_{max} r_c^2}{c_t}$, Eq. (21) can be rearranged as:

$$a_0 = \left[\frac{4\pi F_{max} r_c^2}{c_t} \right] \left(\frac{1}{4\pi m_{-e} c_t} \right)$$

Or, in terms that *do not require* the electron mass to be measured experimentally:

$$a_0 = \left[\frac{4\pi F_{max} r_c^2}{c_t} \right] \left(\frac{c^2}{8\pi F_{max} c_t r_c} \right)$$

Of course, the terms in brackets equals Planck's constant. See; I told you that it was hidden in there *somewhere*. It's just obscured primarily because the factors of π cancel each other out. This again shows that Planck's constant is nothing but an *aggregate* of other fundamental constants. And to actually show the connection, if we substitute Eq. (21) for a_o in Eq. (26) and solve for α and reduce, the result is – as expected:

$$\alpha = \frac{2c_t}{c}$$

Fine Structure Constant Revisited

This next section is this author's most exhaustive personal contribution to quantum mathematics.

The universe works in mysterious ways, but in ways that are intimately intertwined. The fine structure constant arises in multiple instances in atomic calculations. Richard Feynman described it as “a number that all *real* physicists should have up on their wall to worry about”. The fine structure constant (α) is the factor between the ground state Bohr radius of hydrogen (a_0), the electron's reduced Compton wavelength ($\frac{\lambda_{-e}}{2\pi}$), and also the classical electron radius (r_{-e}). If you multiply a_o by α you get $\frac{\lambda_{-e}}{2\pi}$. If you multiply $\frac{\lambda_{-e}}{2\pi}$ by α you get r_{-e} . This can be written mathematically as:

$$a_0 \alpha^2 = \frac{\lambda_{-e}}{2\pi} \alpha = r_{-e}$$

In a similar fashion, the force constant of 29.05N (F_{max}) and the radial constant of (r_c) are intrinsic to the most basic quantum calculations. They also arise in many different ways. For instance, F_{max} is a real measurable maximum of electrical force in terms of two protons at a distance of separation of r_c (right before they get sucked together by the strong force), whilst the F_{max} that comes into play in the spring constant of the electron orbitals is purely mathematical and describes the quantum of momentum for the electron. The fine structure constant arises in mysterious ways because it is simply an aggregate of other fundamental constants, just as Planck's constant arises in all of the quantum mechanical equations because *it too* is an aggregate of the fundamental constants. Allow me to demonstrate this with some equations that I formulated through basic algebra and using basic textbook equations. The reader is encouraged to check the math for themselves (the table of values for the nomenclature is on the last page). The fine structure constant arises in all of these equations because it is an aggregate that can be expressed in a multitude of ways. The most basic form is one that is very familiar by now:

$$\alpha = \frac{2c_t}{c}$$

But it can also be expressed in a much more complicated form:

$$\alpha = \frac{c_t e^2}{8\pi\epsilon_0 r_c^2 c F_{max}}$$

To formulate the above equation for α , we start by taking the standard equation for Bohr radius (written with the unreduced h).

$$a_0 = \frac{4\pi\epsilon_0 h^2}{4\pi^2 m_e e^2}$$

And plugging in for h with our known value of $\frac{4\pi F_{max} r_c^2}{c_t}$ (which was shown in the previous section, and will be also shown in the next) gives us:

$$a_0 = \frac{\epsilon_0 \left(\frac{4\pi F_{max} r_c^2}{c_t} \right)^2}{\pi m_e e^2}$$

When the terms in () are squared, it yields:

$$a_0 = \frac{16\pi^2 \epsilon_0 F_{max}^2 r_c^4}{\pi c_t^2 m_e e^2}$$

And it was already shown in Eq (26) that:

$$a_0 = \frac{2F_{max}r_c^2}{m_{-e}c\alpha c_t}$$

So because $a_0 = a_0$, and we have two different formulas, then we can set them equal.

$$\frac{2F_{max}r_c^2}{m_{-e}c\alpha c_t} = \frac{16\pi\epsilon_0 F_{max}^2 r_c^4}{c_t^2 m_{-e} e^2}$$

Solving for α and reducing gives us:

$$\alpha = \frac{c_t e^2}{8\pi\epsilon_0 r_c^2 c F_{max}}$$

And the standard equation for the classical radius found in textbooks is expressed as:

$$r_{-e} = \frac{e^2}{4\pi\epsilon_0 m_{-e} c^2}$$

This author made a mathematical discovery which allowed insight into α to progress:

$$m_{-e} = \frac{2F_{max}r_c}{c^2}$$

So that means we can substitute $\frac{2F_{max}r_c}{c^2}$ in for m_{-e} in the above equation $\left(r_{-e} = \frac{e^2}{4\pi\epsilon_0 m_{-e} c^2}\right)$, which after reducing gives us:

$$r_{-e} = \frac{e^2}{8\pi\epsilon_0 F_{max} r_c}$$

But the *simplest* way to express the classical radius of the electron is just:

$$r_{-e} = 2r_c$$

All of the basic constants can be expressed in complex *and* simple ways; it should be apparent by now. An interesting side note is that: $E = mc^2$ and $m = \frac{E}{c^2}$ so the rest energy of an electron can be expressed as:

$$E_{-e} = 2F_{max}r_c$$

This also means that we can say the energy contained in an electron is $E_{-e} = r_{-e}F_{max}$, and the mass of an electron is $m_{-e} = \frac{r_{-e}F_{max}}{c^2}$. We can also define the electron's mass without F_{max} as: $m_{-e} = \frac{e^2}{8\pi\epsilon_0 r_c c^2}$. So that means the conversion factor between the classical electron radius and its mass is $\frac{F_{max}}{c^2}$. Lots more coincidences, right?

So now let's see exactly how α is a factor in the conversion from the Compton wavelength (λ_{-e}) to classical electron radius (r_{-e}). Another standard equation involving Planck's constant (h) from any textbook:

$$m_{-e}c\lambda_{-e} = h$$

So we can plug in our value $\frac{4\pi F_{max}r_c^2}{c_t}$ from the last segment in for h :

$$m_{-e}c\lambda_{-e} = \frac{4\pi F_{max}r_c^2}{c_t}$$

Solving for λ_{-e} gives us:

$$\lambda_{-e} = \frac{4\pi F_{max}r_c^2}{c_t c m_{-e}}$$

And we already know that:

$$m_{-e} = \frac{2F_{max}r_c}{c^2}$$

Plugging in for m_{-e} and reducing:

$$\lambda_{-e} = \frac{2\pi c r_c}{c_t}$$

Dividing by 2π gives us:

$$\frac{\lambda_{-e}}{2\pi} = \frac{2\pi c r_c}{2\pi c_t} = \frac{c r_c}{c_t}$$

So if we divide $\frac{\lambda_{-e}}{2\pi} = \frac{c r_c}{c_t}$ by $r_{-e} = \frac{e^2}{8\pi\epsilon_0 F_{max} r_c}$, it should give us the formula for $\alpha = \frac{c_t e^2}{8\pi\epsilon_0 r_c^2 c F_{max}}$ that we already solved for above. Let's multiply the reciprocal of $\frac{c r_c}{c_t}$, which is $\frac{c_t}{c r_c}$ by our formula for $r_{-e} = \frac{e^2}{8\pi\epsilon_0 F_{max} r_c}$ (flipping and multiplying is the same as dividing):

$$\left(\frac{c_t}{c r_c}\right) \left(\frac{e^2}{8\pi\epsilon_0 F_{max} r_c}\right) = \frac{c_t e^2}{8\pi\epsilon_0 r_c^2 c F_{max}} = \alpha$$

Our formulation for α was indeed proven. And if this is true, we should also be able to take the reduced Compton wavelength ($\frac{\lambda_{-e}}{2\pi} = \frac{(\frac{2\pi cr_c}{2\pi})}{\frac{c_t}{c_t}} = \frac{cr_c}{c_t}$) times the fine structure constant ($\alpha = \frac{c_te^2}{8\pi\epsilon_0r_c^2cF_{max}}$) and be able to formulate the same classical electron radius (r_{-e}) that we have already solved for ($\frac{e^2}{8\pi\epsilon_0F_{max}r_c}$). This can be written mathematically as:

$$\left(\frac{cr_c}{c_t}\right)\left(\frac{c_te^2}{8\pi\epsilon_0r_c^2cF_{max}}\right) = \frac{cr_cc_te^2}{8\pi\epsilon_0r_c^2cc_tF_{max}}$$

And after reducing that leaves us with:

$$\frac{e^2}{8\pi\epsilon_0r_cF_{max}} = r_{-e}$$

Just like was expected. And this is why $\frac{\lambda_{-e}}{2\pi}\alpha = r_{-e}$. It's because we can formulate them in new ways that show us the true constituents of the fine structure constant. In this instance: $\lambda_{-e} = \frac{2\pi cr_c}{c_t}$, $r_{-e} = \frac{e^2}{8\pi\epsilon_0r_cF_{max}}$, and $\alpha = \frac{c_te^2}{8\pi\epsilon_0r_c^2cF_{max}}$. When you do the math straight across with the factor of 2π , it yields the correct formula and number every time. $\alpha = 0.007297$.

Let's do it the other way, just to make sure we get the right answer. We'll divide r_{-e} by α to see if we get $\frac{\lambda_{-e}}{2\pi}$. So let's multiply by the reciprocal of α to make it easier, and then reduce:

$$\left(r_{-e} = \frac{e^2}{8\pi\epsilon_0r_cF_{max}}\right) \times \left(\frac{1}{\alpha} = \frac{8\pi\epsilon_0r_c^2cF_{max}}{c_te^2}\right) = \frac{cr_c}{c_t} = \frac{\lambda_{-e}}{2\pi}$$

It looks like it works every way from Sunday. Now let's go from the Bohr radius (a_0) down to λ_{-e} and see how the factor of α fits in with this one.

We start off with the standard equation:

$$2\pi\alpha a_0 = \lambda_{-e}$$

And its already been established that:

$$a_0 = \frac{F_{max}r_c^2}{c_t^2m_{-e}}$$

And we can plug in $\frac{2F_{max}r_c}{c^2}$ for the m_{-e} in $a_0 = \frac{F_{max}r_c^2}{c_t^2 m_{-e}}$ which reduces to:

$$a_0 = \frac{c^2 r_c}{2c_t^2}$$

And we also already know that:

$$\lambda_{-e} = \frac{2\pi c r_c}{c_t}$$

So now the equation is all set up:

$$2\pi\alpha \left(\frac{c^2 r_c}{2c_t^2} \right) = \frac{2\pi c r_c}{c_t}$$

But to actually solve for α , we can set up the equation by flipping and multiplying:

$$\left(\frac{2\pi c r_c}{c_t} \right) \left(\frac{2c_t^2}{2\pi c^2 r_c} \right) = \frac{4\pi r_c c_t^2 c}{2\pi r_c c_t c^2}$$

And $\frac{4\pi r_c c_t^2 c}{2\pi r_c c_t c^2}$ reduces to simply:

$$\frac{2c_t}{c} = \alpha$$

Now let's see how the fine structure constant relates when going *straight* from the Bohr radius down all the way to the classical electron radius. The Bohr radius times α^2 equals the classical electron radius; expressed mathematically as:

$$\alpha^2 a_0 = r_{-e}$$

Solving for α^2 :

$$\alpha^2 = \frac{r_{-e}}{a_0}$$

Plug in our equation of $\frac{e^2}{8\pi\epsilon_0 F_{max} r_c}$ for r_{-e} , and our equation of $\frac{c^2 r_c}{2c_t^2}$ for a_0 :

$$\alpha^2 = \frac{\left(\frac{e^2}{8\pi\epsilon_0 F_{max} r_c} \right)}{\left(\frac{c^2 r_c}{2c_t^2} \right)}$$

Flipping and multiplying to divide it through:

$$\alpha^2 = \left(\frac{e^2}{8\pi\epsilon_0 F_{max} r_c} \right) \left(\frac{2c_t^2}{c^2 r_c} \right)$$

Combining and reducing yields:

$$\alpha^2 = \frac{e^2 c_t^2}{4\pi\epsilon_0 F_{max} r_c^2 c^2}$$

To solve for α , we could say:

$$\alpha = \sqrt{\frac{e^2 c_t^2}{4\pi\epsilon_0 F_{max} r_c^2 c^2}}$$

But the above equation does not really tell us much, because not all of the constants were squared, so we can't reduce it further. But there is a very simple way to show how both the constituent α 's are truly taking form.

Just remove $\alpha = \frac{2c_t}{c}$ from the equation $\left(\alpha^2 = \frac{e^2 c_t^2}{4\pi\epsilon_0 F_{max} r_c^2 c^2} \right)$ and the other value for alpha is produced.

So, one α equals:

$$\alpha = \frac{c_t e^2}{8\pi\epsilon_0 r_c^2 c F_{max}}$$

And the other α equals:

$$\alpha = \frac{2c_t}{c}$$

Both of these forms of α were already previously formulated. They fell right out of simply dividing one length by the other. Now you understand why the fine structure constant is a conversion factor between these three fundamental lengths, as I have proven the relationship through math. For the first time in history you can fully see the constituents that make up the fine structure constant and how they relate to the aggregate of Planck's constant and the speed of light (which determines the classical electron radius). I encourage others to share these equations and ideas and expand upon them; but if you do, please give credit where it is due. This section on fine structure constant relationships and the mass of the electron was discovered by this author, but of course, my work would be *non-existent* without the years of dedication and countless hours of teaching by Frank

Znidarsic. Frank is the one who discovered the speed of transition, $\alpha = \frac{2c_t}{c}$, the radial constant (r_c), the force constant (F_{max}), and using these was able to develop the most pertinent equations *of all*: The Energy Contained in a Photon, The Orbital Radii Of Hydrogen, and...

The Probability of Transition / Intensity of Spectral Emission

Next, we will formulate the equation for the probability of transition. Together with the energy of a photon and the orbital radii of an atom, this unassuming group forms the trio that defines the most important aspects of quantum reality; how much energy it takes to displace an electron and why the frequency of the photon is proportional to the energy level, the size of an atom after an electron is captured, how often the atom will give off photons, and what the intensity of the specific frequencies of spectral emissions are. In reality, atoms are a lot simpler than most people imagine. They don't really do a whole lot; mainly just emit and absorb photons. The rate at which photons are given off is extremely important to life here on earth; if it occurred too quickly, our sun would possibly be thrown out of hilt and intelligent life would maybe not have had ample time to evolve. The equation beknownst by the world of academia to calculate the probability of transition takes many blackboards full of equations to formulate. Frank Znidarsic taught me how to formulate a version without Planck's constant in a few easy steps. We start off with an expression that should be slightly familiar from the orbital radii equations:

$$(2\pi f)r = \sqrt{\frac{k_{-e}}{m_{-e}}}nr_c$$

Because both $(2\pi f)$ and $\left(\sqrt{\frac{k_{-e}}{m_{-e}}}\right)$ are equal to omega (ω), we are basically saying $\omega r = \omega r$, although they are entirely different ω 's and different r 's. This is representative of both the transitional state of the photon and that of the electron. Of course, n corresponds to the orbital radii. The spring constant of the electron in the transitional state is:

$$k_{-e} = \frac{F_{max}}{nr_c}$$

This is because the spring constant of the electron can always be expressed as $\frac{F_{max}}{r}$, in both the static and transitional states. We saw in the equations for the orbital radii that the r turned out to be the orbital of the static atomic state. In the above equation, we are describing the spring constant of the transitional atomic state. In the equation for orbital radii we saw that the radius of the transitional atomic state was a multiple of r_c . So the r in our spring constant for the description of the transitional orbital radii also needs to be a multiple of r_c , as this is the radius of the collapsed transitional wavefunction.

Substituting that into the equation for k_{-e} gives us:

$$(2\pi f)r = \sqrt{\frac{\left(\frac{F_{max}}{nr_c}\right)}{m_{-e}}}(nr_c)$$

Solving for r gives us:

$$r = \sqrt{\frac{F_{max}}{nr_cm_{-e}}}\left(\frac{nr_c}{2\pi f}\right)$$

Then when we square it, to get rid of the square root, we get:

$$r^2 = \left(\frac{F_{max}}{nr_cm_{-e}}\right)\left(\frac{n^2r_c^2}{4\pi^2f^2}\right)$$

Combining the terms in parenthesis:

$$r^2 = \frac{F_{max}nr_c}{4\pi^2f^2m_{-e}}$$

Next, we need to substitute in the Compton frequency of the electron (f_{-e}). The Compton frequency arises from combining the photoelectric equation ($E = hf$) and Einstein's rest-energy equation ($E = mc^2$). We will use the aggregate form of h that we solved for in the orbital equations $\left(\frac{4\pi F_{max}r_c^2}{c_t}\right)$.

$$m_{-e}c^2 = \left[\frac{4\pi F_{max}r_c^2}{c_t}\right]f_{-e}$$

We know that $m_{-e} = \frac{2F_{max}r_c}{c^2}$, so we can substitute that in for m_{-e} . Doing so and solving for f_{-e} yields:

$$f_{-e} = \frac{2F_{max}r_c c_t}{4\pi F_{max}r_c^2}$$

Reducing the above equation simply leaves us with:

$$f_{-e} = \frac{c_t}{2\pi r_c}$$

When we factor in this Compton frequency, it is replacing one of the f 's in the f^2 . So, it only leaves one f , which represents the frequency of the emitted photon (which is the same frequency for the transitional state as well as once it's vacated the atom). When we replace one of the f 's with our Compton frequency (f_{-e}), it gives us:

$$r^2 = \frac{2\pi F_{max}nr_c^2}{4\pi^2 c_t f m_{-e}}$$

Which reduces down to:

$$r^2 = \frac{F_{max}nr_c^2}{2\pi c_t f m_{-e}}$$

This expresses the exact probability of transition related to the frequency of the photon that is taught in quantum mechanical textbooks. The equation can be regrouped as:

$$r^2 = \left[\frac{4\pi F_{max}r_c^2}{c_t} \right] \left(\frac{n}{8\pi^2 m_{-e} f} \right)$$

The terms in the brackets $\left[\frac{4\pi F_{max}r_c^2}{c_t} \right]$ are equal to Planck's constant. Lo and behold; it is the *exact same* version of Planck's constant that was produced by solving for it in the orbital radii equations. Now I've shown how to formulate it from scratch. Also, when you replace the terms in the brackets with the symbol for Planck's constant, it simply becomes:

$$r^2 = \frac{nh}{8\pi^2 m_{-e} f}$$

The above equation is the standard expression for the probability of transition which is written in the textbooks. Of course, we could also express the equation $r^2 = \frac{F_{max}nr_c^2}{2\pi c_t f m_{-e}}$ by substituting in $\frac{2F_{max}r_c}{c^2}$ for the mass of the electron (m_{-e}). After reducing it then becomes:

$$r^2 = \frac{nc^2r_c}{4\pi c_t f}$$

Just another coincidence, right?

And as a final side note, the Compton wavelength/frequency of the electron can also be formulated with the *other* fundamental constants in equations that are more complicated. We can do this by using the the original photoelectric equation that we formulated at the beginning of the paper. We set $E = \left[\frac{e^2}{4\varepsilon_0 c_t} \right] f$ equal to $E = mc^2$.

$$m_{-e}c^2 = \left[\frac{e^2}{4\varepsilon_0 c_t} \right] f_{-e}$$

We have to substitute $\frac{2F_{max}r_c}{c^2}$ in for m_{-e} again and solving for f_{-e} and then reducing gives us:

$$f_{-e} = \frac{8\varepsilon_0 c_t F_{max} r_c}{e^2}$$

And because $\lambda = \frac{c}{f}$, when we divide c by $f_{-e} = \frac{8\varepsilon_0 c_t F_{max} r_c}{e^2}$, it yields:

$$\lambda_{-e} = \frac{ce^2}{8\varepsilon_0 c_t F_{max} r_c}$$

So now we have two forms of the Compton wavelength and frequency, the others being: $\lambda_{-e} = \frac{2\pi cr_c}{c_t}$ and $f_{-e} = \frac{c_t}{2\pi r_c}$.

Comparison of these two forms also illuminates how π is related to the fundamental constants:

$$\pi = \frac{e^2}{16\varepsilon_0 r_c^2 F_{max}}$$

And if we substitute the version of $f_{-e} = \frac{8\varepsilon_0 c_t F_{max} r_c}{e^2}$ in for the Compton frequency into the probability of transition equation, we arrive at:

$$r^2 = \frac{e^2 nc^2}{64\pi^2 \varepsilon_0 c_t F_{max} r_c f}$$

And if one was feeling pugnacious, they could reformulate the ground state Bohr hydrogen radius also as:

$$a_0 = \frac{e^2 c^2}{32\pi\epsilon_0 F_{max} c_t^2 r_c}$$

Or with the factors of n and Z which produce all the $1s$ orbital radii of all the elements:

$$a_0 = \frac{(enc)^2}{32\pi\epsilon_0 F_{max} c_t^2 r_c Z}$$

But if you really want to go bananas, the Rydberg constant (R_∞) can also be formulated as:

$$R_\infty = \frac{c_t^3 e^2}{16\pi^2 \epsilon_0 r_c^3 c^3 F_{max}}$$

Or:

$$R_\infty = \frac{c_t^3 e^4}{256\pi^3 \epsilon_0^2 r_c^5 F_{max}^2 c^3}$$

Or:

$$R_\infty = \frac{16\epsilon_0 c_t^3 F_{max} r_c}{c^3 e^2}$$

Or simply:

$$R_\infty = \frac{c_t^3}{\pi r_c c^3}$$

Notice that *none* of these formulations for R_∞ require an experimentally measured mass of the electron; something that the world of academia can only do with a single equation: $R_\infty = \frac{\alpha^2}{2\lambda_{-e}}$.

Conclusion:

I have shown how the three major aspects of quantum nature can be formulated from the ground up in a few simple steps. Furthermore, I have shown that this is the only model which truly shows the reasoning for how the trio of intrinsic lengths (r_{-e} , λ_{-e} , and a_0) are all related by the fine structure constant, and I've also shown how the fine structure constant (α) and Planck's constant (h) are in reality *both* made up of *two* forms of constituent constants. The

two basic equations to describe each are both in terms of the fundamental constants. Even the mass of the electron can be expressed in multiple ways with these fundamental constants, and doesn't require experimental measurement. The immense amount of equalities listed in this paper show beyond any reasonable doubt that the new constants introduced by Znidarsic (c_t , F_{max} , r_c) are not coincidental. The relationships that they comprise are real and insurmountable. Furthermore, this model sheds light on *why* the electron cannot crash into the nucleus, as one would expect. The electron has to transition from parent state to daughter state through the transitional (collapsed) subset of the wavefunction. There *is* no collapsed transitional wavefunction smaller than ground state transitional orbital (r_c for the hydrogen atom), because that is the smallest multiple that it can be; only positive integers are allowed for the factor of n . Because there is no transitional state for it to transition to from the ground state, in essence; there is no daughter state for it to transition into during its transitional (collapsed) wavefunction progression. This essentially locks the electron into the ground state orbital (it can only jump higher, not lower), rendering it into a state of perpetual motion; that if left undisturbed, will stay in motion until the end of the universe. This paper is the first time in history that a theory for true causation has been given to explain this phenomenon, and although the novel theories presented in this paper would be non-existent without Frank Znidarsic's groundwork, the collapsed transitional wavefunction theory is this author's contribution that leads us to glean this understanding of the mechanism for why electrons cannot transgress past the ground state energy level. Furthermore, this model gives an explanation for why the electron transitions from one orbital to the next, instead of transitioning straight from an orbital level higher than 2, down to orbital level 1 (a_0). This is because the orbitals have to transition from parent state to daughter state via the transitional collapsed wavefunction subset. Each collapsed transitional wavefunction is modeled as its own individual spring. These separate springs are not congruent; they can only transfer to the immediately neighboring energy level through the transitional subset. Once one transition takes place, there is a discontinuity because of the different springs relating to the different static/transitional orbitals. A photon representing the jump from only a single energy level is released as the transition in the energy level of the atom is made. A new probability of transition is then related to daughter orbital, and a new corresponding transitional subset. After over one hundred years of complete mystery (compounded by willful ignorance) regarding the causations which give birth to the quantum mechanical nature that is undeniably observed, the first look past the simple *mechanics* of quantum nature to the explanations for causation; the first true look at quantum *physics* has taken root and begun to sprout.

Nomenclature: h = Planck's constant ($6.626069 \times 10^{-34} \text{J/s}$)
 α = The fine Structure Constant (0.0072973525)
 c_t = The speed of transition (1,093,845 m/s)
 c = The speed of light in vacuo (299,792,458 m/s)
 C = Capacitance
 f = Frequency
 f_{-e} = The Compton frequency of the electron ($1.23559 \times 10^{20} \text{Hz}$)
 λ = Wavelength
 λ_t = Transitional wavelength of the photon
 λ_{-e} = The Compton Wavelength of the electron ($2.42631 \times 10^{-12} \text{m}$)
 ε_0 = The permittivity of free space ($8.854 \times 10^{-12} \text{F/m}$)
 E = Energy
 e = Elementary charge ($1.602 \times 10^{-19} \text{C}$)
 Q = Charge
 A = Area
 d = distance
 r = radius
 r_c = The Coulombic and collapsed a_0 radius ($1.409 \times 10^{-15} \text{m}$)
 r_{-e} = Classical radius of the electron ($2.8179 \times 10^{-15} \text{m}$)
 r_n = The spacing of the nucleons/Fermi spacing ($1.36 \times 10^{-15} \text{m}$)
 F_{max} = The force constant for harmonic motion (29.053N)
 a_0 = The Ground state Bohr radius ($.529 \times 10^{-10} \text{m}$)
 m_{-e} = Mass of the electron, ($9.10938 \times 10^{-31} \text{kg}$)
 m_n = Average mass of nucleons ($1.6737 \times 10^{-27} \text{kg}$)
 n = Quantum number/orbital level (a positive integer)
 k = Spring constant
 ω = Omega (angular frequency)
 π = Pi (3.1416)
 Z = Atomic number (positive integer)

* This paper is ©2010 all rights reserved to Lane Davis, but permission is granted for reproduction and dissemination for non-profit purposes, as long as credit is given to Frank Znidarsic and this author. If you wish to republish this paper on a website, please inform me via email at seattle.truth@gmail.com. If you have questions regarding theory, please check out the 5+ hour video series explaining all of this in detail at QuantumTransition.com (where you can also find a forum for questions and discussion) and on YouTube at YouTube.com/Seattle4Truth.

** In honor of the late, great, Albert Einstein and his celebrated paper which announced the famous equation $E = mc^2$ to the world, this dissertation does not include any references.