

Chapter 6

Scalar Waves



There is wide confusion about what “scalar waves” are both in serious and less serious literature on electrical engineering. This chapter explains that this type of wave is a longitudinal wave of potentials. It has been shown that a longitudinal wave is a combination of a vector potential with a scalar potential. There is a full analogue to acoustic waves. Transmitters and receivers for longitudinal electromagnetic waves are discussed. Scalar waves were found and used at first by Nikola Tesla in his wireless energy transmission experiments. The SW is the extension of the Maxwell’s equation part that we call the More Complete Electromagnetic equation, as described herein.

6.1 Introduction

It is the purpose of this chapter to discuss a new unified field theory based on the work of Tesla. This *unified field* and particle theory describes quantum and classical physics, mass, gravitation, the constant speed of light, neutrinos, waves, and particles—all of which can be explained by vortices [1]. In addition, we discuss unique and various recent inventions and their possible modes of operation in order to convince those studying this of their value for hopefully directing a future program geared toward the rigorous clarification and certification of the specific role the electroscalar domain, and how it might play a role in shaping a future consistent with classical electrodynamics. Also, by extension, to perhaps shed light on inconsistencies that do exist within current conceptual and mathematical theories in the present interpretation of relativistic quantum mechanics. In this regard, we anticipate that by incorporating this more expansive electrodynamic model, the source of the extant problems with gauge invariance in quantum electrodynamics and the subsequent unavoidable divergences in energy/charge might be identified and ameliorated.

Not only does the electroscalar domain have the potential to address such lofty theoretical questions surrounding fundamental physics, but another aim in this chapter

is to show that the protocol necessary for generating these field effects may not be present only in exotic conditions involving large field strengths and specific frequencies involving expensive infrastructures such as the *large hadron collider* (LHC). As recent discoveries suggest, however, SWs may be present in the physical manipulation of everyday objects. We also will explain that nature has been and may be engaged in the process of using *scalar longitudinal waves* (SLWs) in many ways yet unsuspected and undetected by humanity. Some of the modalities of SW generation we will investigate include chemical bond-breaking, particularly as a precursor to seismic events (i.e., illuminating the study and development of earthquake early warning systems); solar events (i.e., related to eclipses); and sunspot activity and how it affects the Earth's magnetosphere. Moreover, this overview of the unique aspects of the electroscalar domain suggests that many of the currently unexplained anomalies—for example, overunity power observed in various energy devices and exotic energy effects associated with *low-energy nuclear reactions* (LENRs)—may find some basis in fact.

In regard to the latter, cold fusion or LENR fusion-type scenarios, the *electroscalar wave* might be the actual agent needed to reduce the nuclear Coulomb barrier, thereby providing the long sought for viable theoretical explanation of this phenomenon [2]. Longitudinal electrodynamic forces (e.g., in exploding wires) actually may be because of the operation of electroscalar waves at subatomic levels of nature. For instance, the extraordinary energies produced by Ken Shoulder's charge clusters (i.e. Particles of like charge repel each other - that is one of the laws describing the interaction between single sub-atomic particles) perhaps may be because of electroscalar mechanisms.

Moreover, these observations, spanning as they do many cross-disciplines of science, beg the question as to the possible universality of the SLW—that is, the concept of the longitudinal electroscalar wave, not present in current electrodynamics, may represent a general key, overarching principle, leading to new paradigms in other sciences besides physics. This idea also will be explored in the chapter, showing the possible connection of scalar-longitudinal (i.e., electroscalar) wave dynamics to biophysical systems. Admittedly, we are proposing quite an ambitious agenda in reaching for these goals, but we think you will see that recent innovations may have proved equal to the task of supporting this quest.

6.2 Descriptions of Transverse and Longitudinal Waves

As you know from a classical physics point of view, typically the following are three kinds of waves—the *soliton wave* is an exceptional case and should be addressed separately—and wave equations that we can talk about:

1. Mechanical waves (i.e., waves on string)
2. Electromagnetic (EM) waves (i.e., $\vec{E}$ and $\vec{B}$ fields from Maxwell's equations to deduce the wave equations, where these waves carry energy from one place to another)
3. Quantum mechanical waves (i.e., using Schrödinger equations to study particles' movements)

The second one is our subject of interest in terms of the two types of waves involved in EM waves: (1) transverse waves and (2) *longitudinal pressure waves* (lpws), also known as *scalar longitudinal waves* (slws).

From the preceding two waves, the SLW is of interest in directed energy weapons (DEWs) [3] and here is why. First, we briefly describe SLWs and their advantages for DEW purposes as well as communication within non-homogeneous media such as seawater with different *electrical primitivity* ϵ and *magnetic permeability* μ at various ocean depths (see Chap. 4 of this book).

A wave is defined as a disturbance that travels through a certain medium. The medium is material through which a wave moves from one to another location. If we take as an example a slinky coil that can be stretched from one end to the other, a static condition then develops. This static condition is called the wave's *neutral condition* or equilibrium state.

In the slinky coil the particles are moved up and down then come into their equilibrium state. This generates disturbance in the coil that is moved from end one to the other. This is the movement of a *slinky pulse*, which is a *single disturbance in medium* from one to another location. If it is done *continuously* and in a *periodical manner*, then it is called a *wave*, also known as an *energy transport medium*. They are found in diverse shapes, show a variety of behaviors, and have characteristic properties. On this basis, they are classified mainly as longitudinal, transverse, and surface waves. Here we discuss the properties of LWs and provide examples. The movement of waves is parallel to the medium of the particles in them.

1. Transverse Waves

For TWs the medium is displaced perpendicular to the wave's direction of propagation. A ripple on a pond and a wave on a string are visualized easily as TWs (Fig. 6.1).

TWs cannot propagate in a gas or a liquid because there is no mechanism for driving motion perpendicular to the propagation of the wave. In summary, a transverse wave (TW) is a wave in which the oscillation is perpendicular to the direction of wave propagation. Electromagnetic waves (and secondary waves, S-waves or shear waves, sometimes called elastic S-waves) in general are TWs.

2. Longitudinal Waves

In a LW the displacement of the medium is parallel to the propagation of the wave. A wave in a slinky is a good visualization. Sound waves in air are LWs (Fig. 6.2).

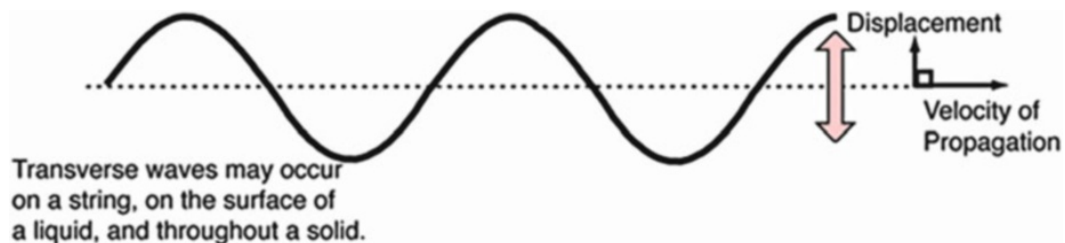


Fig. 6.1 Depiction of a transverse wave

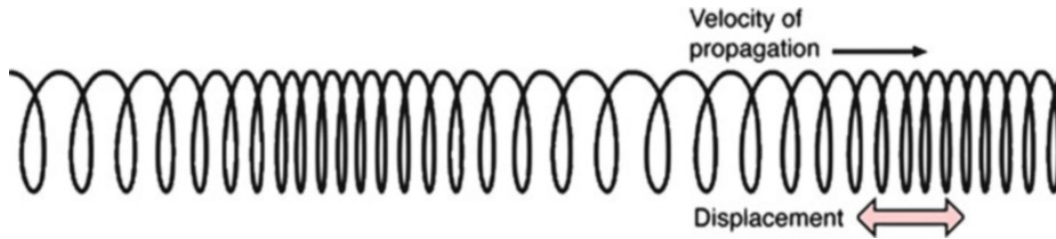


Fig. 6.2 Depiction of a longitudinal wave

In summary, an LW is a wave in which the oscillation is in an opposite direction to the direction of wave propagation. Sound waves—that is, primary waves, or P-waves, in general—are LWs. On the other hand, a wave with a motion that occurs through the particles of the medium oscillating about their mean positions in the direction of propagation is called an LW.

We use our knowledge to expand the subject of the *longitudinal wave* before we go deeper into the subject of the *scalar longitudinal wave*; for an LW the particles of the medium vibrate in the direction of wave propagation. An LW proceeds in the form of compression and rarefaction, which is stretch and compression in the same direction as the wave moves. For an LW at places of compression the pressure and density tend to be maximal, whereas at places where rarefaction takes place, the pressure and density are minimal. In gases only an LW can propagate; LWs are known as *compression waves*.

An LW travels through a medium in the form of compressions or condensations, C, and rarefaction, R. A compression is a region of the medium in which particles are compressed (i.e., particles come closer); in other words, the distance between the particles becomes less than the normal distance between them. Thus, volume temporarily decreases and, therefore the density of the medium increases in the region of compression. A *rarefaction* is a region of the medium in which particles are rarefied (i.e., particles get farther apart than what they normally are). Thus, volume temporary increases and, consequently, the density of the medium decreases in the region of rarefaction.

The distance between the centers of two consecutive rarefactions and two consecutive compressions is called *wavelength*. Examples of LWs are sound waves, tsunami waves, earthquake P-waves, ultra-sounds, vibrations in gas, and oscillations in springs internal water waves, waves in slinky, and so on.

(a) *Longitudinal waves*

Examples of the various types of waves are:

1. Sound wave
2. Earthquake *P*-wave
3. Tsunami wave
4. Waves in a slinky
5. Glass vibrations

- 6. Internal water waves
- 7. Ultra-sound
- 8. Spring oscillations

(b) *Sound waves*

Now the question is: *Are sound waves longitudinal?* The answer is *Yes*. A sound wave travels as an LW in nature. It behaves as a TW in solids. Through gases, plasma, and liquids the sound travels as an LW. Through solids the wave can be transmitted as a TW or an LW.

A material medium is mandatory for the propagation of the sound waves. They mostly are longitudinal in common nature. The speed of sound in air is 332 m/s at normal temperature and pressure. Vibrations of an air column above the surface of water in the tube of a resonance apparatus are longitudinal. Vibrations of an air column in organ pipes are longitudinal. Sound is audible only between 20 Hz and 20 KHz. Sound waves cannot be polarized.

- (i) *Propagation of sound waves in air:* Sound waves are classified as LWs. Let us now see how sound waves propagate. Take a tuning fork, vibrate it, and concentrate on the motion of one of its prongs, say prong A. The normal position of the tuning fork and the initial condition of air particles is shown in Fig. 6.3a. As prong A moves toward the right, it compresses air particles near it, forming a compression as shown in Fig. 6.3b. Because of vibrating air layers, this compression moves forward as a disturbance.

As prong A moves back to its original position, the pressure on its right decreases, thereby forming a rarefaction. This rarefaction moves forward like compression as a disturbance. As the tuning fork goes on vibrating, waves consisting of alternate compressions and rarefactions spread in the air as shown in Fig. 6.3c,d. The direction of motion of the sound waves is the same as that of air particles, thus they are classified as LWs. The LWs travel in the form of compressions and rarefactions.

The main parts of the sound wave follow, along with descriptions:

1. *Amplitude:* The maximum displacement of a vibrating particle of the medium from the mean position. A shows amplitude in $y = A \sin(\omega t)$. The maximum height of the wave is called its amplitude. If the sound is more, then the amplitude is more.
2. *Frequency:* The number of vibrations made per second by the particles and is denoted by f , which is given as $f = 1/T$ and its unit is Hz. We also can get the expression for angular frequency.
3. *Pitch:* It is characteristic of sound with the help of which we can distinguish between a *shrill* note and a note that is *grave*. When a sound is shriller, it is said to be of higher pitch and is found to be of greater frequency, as $\omega = 2\pi f$. On the other hand, a grave sound is said to be of low pitch and is of low frequency. Therefore the pitch of a sound depends on its frequency. It should be made clear that pitch is not the frequency but changes with frequency.

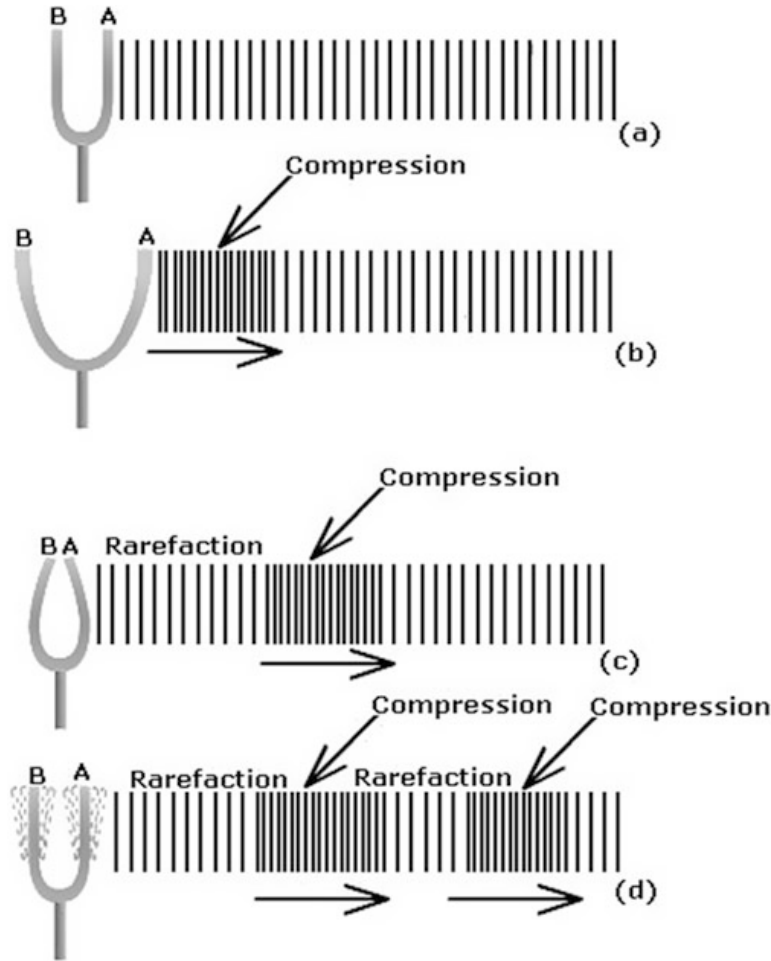


Fig. 6.3 Tuning fork

4. *Wavelength*: The distance between two consecutive particles in the same phase or the distance traveled by the wave in one periodic time and denoted by lambda, Λ .
5. *Sound wave*: This is a LW with regions of compression and rarefaction. The increase of pressure above its normal value may be written as:

$$\sum p = \sum p_0 \sin \omega \left(t - \frac{c}{v} \right) \quad (6.1)$$

where

$\sum p$ = increase in pressure at x position at time t

$\sum p_0$ = maximum increase in pressure

$\omega = 2\pi f$ where f is frequency

If $\sum p$ and $\sum p_0$ are replaced by P and P_0 , then Eq. 6.1 has the following form:

$$P = P_0 \sin \omega \left(t - \frac{c}{v} \right) \quad (6.2)$$

- (ii) *Sound intensity*: Loudness of sound is related to the intensity of it. The sound's intensity at any point may be defined as the amount of sound energy passing per unit time per unit area around that point in a perpendicular direction. It is a physical quantity and is measured in W/m^2 in S.I. units.

The sound wave falling on the ear drum of the observer produces the sensation of hearing. The sound's sensation, which enables us to differentiate between a loud and a faint sound, is called *loudness*, and we can designate it by symbol L . It depends on the intensity of the sound I and the sensitivity of the ear of the observer at that place. The lowest intensity of sound that can be perceived by the human ear is called the *threshold of hearing*, denoted by I_0 . The mathematical relation between intensity and loudness is:

$$L = \log \frac{I}{I_0} \quad (6.3)$$

The intensity of sound depends on:

Amplitude of vibrations of the source
 Surface area of the vibrating source
 Distance of the source from the observer
 Density of the medium in which sound travels from the source
 Presence of other surrounding bodies
 Motion of the medium

- (iii) *Sound reflection*: When a sound wave gets reflected from a rigid boundary, the particles at the boundary are unable to vibrate. Thus, the generation of a reflected wave takes place, which interferes with the oncoming wave to produce zero displacement at the rigid boundary. At the points where there is zero displacement, the variation in pressure is at a maximum. This shows that the phase of the wave has been reversed, but the nature of the sound wave does not change (i.e., on reflection the compression is reflected as compression and rarefaction as rarefaction. Let the incident wave be represented by the given equation:

$$Y = a \sin (\omega t - kx) \quad (6.4)$$

Then the Eq. 6.4 of reflected wave takes the form

$$Y = a \sin (\omega t + kx + \pi) = -a \sin (\omega t + kx) \quad (6.5)$$

Here in both Eqs. 6.4 and 6.5 the symbol of a is basically the designation of the amplitude of the reflected wave.

A sound wave also is reflected if it encounters a rarer medium or free boundary or low-pressure region. A common example is the traveling of a sound wave in a narrow open tube. On reaching an open end, the wave gets

reflected. So the force exerted on the particles there because of the outside air is quite small and, therefore, the particles vibrate with increasing amplitude. Because of this the pressure there tends to remain at the average value. This means that there is no alteration in the phase of the wave, but the ultimate nature of the wave has been altered (i.e., on the reflection of the wave the compression is reflected as rarefaction and vice versa).

The amplitude of the reflected wave would be a' this time and Eq. 6.4 becomes:

$$y = a' \sin (\omega t + kx) \quad (6.6)$$

- (c) *Wave interface*: When listening to a single sine wave, amplitude is directly related to loudness and frequency and directly related to pitch. When there are two or more simultaneously sounding sine waves, wave interference takes place. There are basically two types of wave interference: (1) constructive and (2) destructive.
- (d) *Decibel*: A smaller and practical unit of loudness is a decibel (dB) and is defined as follows:

$$1 \text{ Decibel} = \frac{1}{10} \text{ bel} \quad (6.7)$$

In dB the loudness of a sound of intensity I is given by

$$L = 10 \log \left(\frac{I}{I_0} \right) \quad (6.8)$$

- (e) *Timber*: Timber can be called the property that distinguishes two sounds and makes them different from each other, even when they have the same frequency. For example, when we play violin and guitar on the same note and same loudness, the sound is still different. It also is denoted as *tone color*.
- (f) *S-waves*: An *S-wave* is a wave in an elastic medium in which the restoring force is provided by shear. *S-waves* are divergence-less,

$$\nabla \cdot \vec{u} = 0 \quad (6.9)$$

where $\vec{u}$ is the displacement of the wave and comes in two polarizations: (1) SV (vertical) and (2) SH (horizontal).

The speed of an *S-wave* is given by:

$$v_s = \sqrt{\frac{\mu}{\rho}} \quad (6.10)$$

where μ is the shear modulus and ρ is the density.

- (g) *P-waves*: Primary waves also are called P-waves. These are compressional waves and are longitudinal in nature. They are a type of pressure wave. The speed of P-waves is greater than other waves. These are called the primary waves

because they are the first to arrive during an earthquake. This is because of great their velocity. The propagation of these waves knows no bounds and thus can travel through any type of material, including fluids.

P-waves, also called pressure waves, are longitudinal waves (i.e. the oscillation occurs in the same direction, and the opposite direction of wave propagation). The restoring force for *P*-waves is provided by the medium's bulk modulus. In an elastic medium with rigidity or shear modules being zero ($\mu = 0$), a harmonic plane wave has the form:

$$S(z, t) = S_0 \cos(kz - \omega t + \phi) \quad (6.11)$$

where S_0 is the amplitude of displacement, k is the wave number, z is the distance along the axis of propagation, ω is the angular frequency, t is the time, and ϕ is a phase offset. From the definition of bulk modulus (K), we can write:

$$K = -V \frac{dP}{dV} \quad (6.12)$$

where V is the volume and dP/dV is the derivative of pressure with respect to volume. The bulk modulus gives the change in volume of a solid substance as the pressure on it is changed, then we can write:

$$\begin{aligned} K &\equiv -V \left(\frac{dP}{dV} \right) \\ &\equiv \rho \left(\frac{\partial P}{\partial \rho} \right) \end{aligned} \quad (6.13)$$

Consider a wavefront with surface area A , then the change in pressure of the wave is given by the following relationship:

$$\begin{aligned} dP &= -K \frac{dV}{V} = -K \frac{A[S(z + \Delta z) - S(z)]}{A\Delta z} \\ &= -K \frac{S(z + \Delta z) - S(z)}{\Delta z} = -K \frac{\partial S}{\partial z} \end{aligned} \quad (6.14)$$

where ρ is the density. The bulk modulus has units of pressure.

6.2.1 Pressure Waves and More Details

As we mentioned in the preceding, the pressure waves present the behavior and concept of LWs, thus many of the important concepts and techniques used to analyze TWs on a string, as part of mechanical waves components, also can be applied to longitudinal pressure waves.

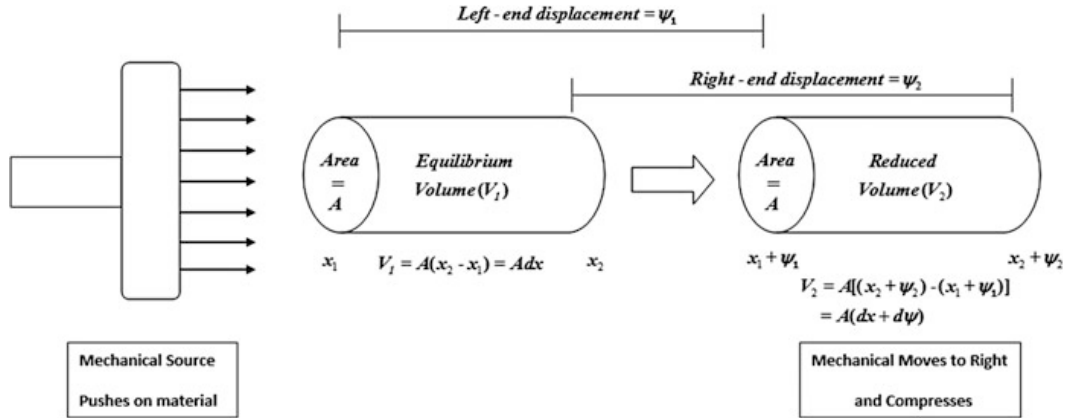


Fig. 6.4 Displacement and compression of a segment of materials

You can see an illustration of how a pressure wave works in Fig. 6.4. As the mechanical wave source moves through the medium, it pushes on a nearby segment of the material, and that segment moves away from the source and is compressed (i.e., the same amount of mass is squeezed into a smaller volume, so the density of the segment increases). That segment of increased density exerts pressure on adjacent segments, and in this way a pulse wave (if the source gives a single push) or a harmonic wave (if the source oscillates back and forth) is generated by the source and propagates through the material.

The “disturbance” of such waves involves three things: (1) the longitudinal displacement of material, (2) changes in the density of the material, and (3) variation of the pressure within the material. So pressure waves also could be called “density waves” or even “longitudinal displacement waves,” and when you see graphs of wave disturbance in physics and engineering textbooks, you should make sure you understand which of these quantities is being plotted as the displacement of the wave.

As you can see in Fig. 6.4, we are still considering one-dimensional wave motion (i.e., the wave propagates only along the x -axis). But pressure waves exist in a three-dimensional medium, so instead of considering linear mass density μ (as we did for the string in the previous section), in this case it is the volumetric mass density ρ that will provide the inertial characteristic of the medium. Nonetheless, just as we restricted the string motion to small angles and considered only the transverse component of the displacement, in this case we will assume that the pressure and density variations are small relative to the equilibrium values and consider only longitudinal displacement (i.e., the material is compressed or rarefied only by changes in the segment length in the x -direction).

The most straightforward route to finding the wave equation for this type of wave is very similar to the approach used for TWs on a string, which means you can use Newton’s second law to relate the acceleration of a segment of the material to the sum of the forces acting on that segment. To do that, start by defining the pressure (P) at any location in terms of the equilibrium pressure (P_0) and the incremental change in pressure produced by the wave (dP):

$$P = P_0 + dP \quad (6.15)$$

Likewise, density (ρ) at any location can be written in terms of equilibrium density (ρ_0) and the incremental change in density produced by the wave ($d\rho$):

$$\rho = \rho_0 + d\rho \quad (6.16)$$

Before relating these quantities to the acceleration of material in the medium using Newton's second law, it is worthwhile to familiarize yourself with the terminology and equations of volume compressibility. As you might imagine, when external pressure is applied to a segment of material, how much the volume (thus the density) of that material changes depends on the nature of the material. To compress a volume of air by 1% requires a pressure increase of about 1000 pascals (Pa or N/m²), but to compress a volume of steel by 1% requires a pressure increase of more than one billion Pa. The compressibility of a substance is the inverse of its "bulk modulus" (usually written as K or B , with units of Pa), which relates an incremental change in pressure (dP) to the fractional change in density ($d\rho/\rho_0$) of the material:

$$K \equiv \frac{dP}{d\rho/\rho_0} \quad (6.17)$$

or

$$dP = K \frac{d\rho}{\rho_0} \quad (6.18)$$

With this relationship in mind, you are ready to consider Newton's second law for the segment of material being displaced and compressed (or rarefied) by the wave. To do that consider the pressure from the surrounding material acting on the left and on the right side of the segment, as shown in Fig. 6.5. Notice that the pressure (P_1) on the left end of the segment is pushing in the positive x -direction and the pressure on

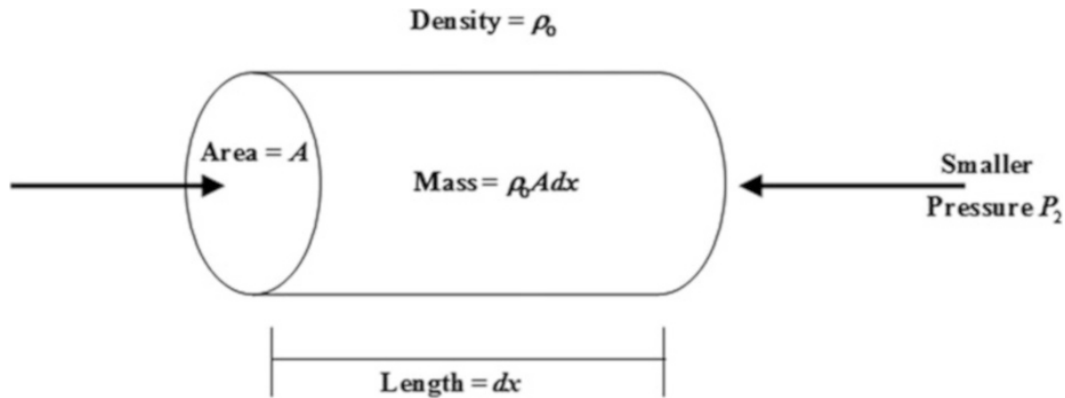


Fig. 6.5 Pressure on a segment of material

the left end of the segment is pushing in the negative χ -direction. Setting the sum of the χ -direction forces equal to the acceleration in the χ -direction gives:

$$\sum F_x = P_1 A - P_2 A = m a_x \quad (6.19)$$

where m is the mass of the segment.

If the cross-sectional area of the segment is A and the length of the segment is dx , the volume of the segment is $A dx$, and the mass of the segment is this volume times the equilibrium density of the material:

$$m = \rho_0 A dx \quad (6.20)$$

Notice also that the pressure on the right end of the segment is smaller than the pressure on the left end because the source is pushing on the left end, which means that the acceleration at this instant will be toward the right.

Using the symbol ψ to represent the displacement of the material because of the wave, the acceleration in the χ -direction can be written:

$$a_x = \frac{\partial^2 \psi}{\partial t^2} \quad (6.21)$$

Substituting these expressions for m and a_x into Newton's second law, Eq. 6.19 gives:

$$\sum F_x = P_1 A - P_2 A = \rho_0 A dx \frac{\partial^2 \psi}{\partial t^2} \quad (6.22)$$

Writing the pressure P_1 at the left end as $P_0 + dP_1$ and the pressure P_2 at the right end as $P_0 + dP_2$ means that

$$\begin{aligned} P_1 A - P_2 A &= (P_0 + dP_1) A - (P_0 + dP_2) A \\ &= (dP_1 - dP_2) A \end{aligned} \quad (6.23)$$

But the change in dP (i.e., the change in the overpressure, or under-pressure, produced by the wave) over the distance dx can be written:

$$\text{Change in over-pressure} = dP_2 - dP_1 = \frac{\partial(dP)}{\partial x} dx \quad (6.24)$$

which means

$$-\frac{\partial(dP)}{\partial x} dx A = \rho_0 A dx \frac{\partial^2 \psi}{\partial t^2} \quad (6.25)$$

or

$$\rho_0 \frac{\partial^2 \psi}{\partial t^2} = - \frac{\partial(dP)}{\partial x} \quad (6.26)$$

But $dP = d\rho K/\rho_0$, so

$$\rho_0 \frac{\partial^2 \psi}{\partial t^2} = - \frac{\partial \left[\left(\frac{K}{\rho_0} \right) d\rho \right]}{\partial x} \quad (6.27)$$

The next step is to relate the change in density ($d\rho$) to the displacements of the left and right ends of the segment (ψ_1 and ψ_2). To do that note that the mass of the segment is the same before and after the segment is compressed. That mass is the segment's density times its volume ($m = \rho V$), and the volume of the segment can be seen to be $V_1 = A dx$ before compression (see Fig. 6.4) and $V_2 = A(dx + d\psi)$ after compression. Thus,

$$\begin{aligned} \rho_0 V_1 &= (\rho_0 + d\rho) V_2 \\ \rho_0 (A dx) &= (\rho_0 + d\rho) A (dx + d\psi) \end{aligned} \quad (6.28)$$

The change in displacement ($d\psi$) over distance dx can be written:

$$d\psi = \frac{\partial \psi}{\partial x} dx \quad (6.29)$$

so

$$\begin{aligned} \rho_0 (A dx) &= (\rho_0 + d\rho) A \left(dx + \frac{\partial \psi}{\partial x} dx \right) \\ \rho_0 &= (\rho_0 + d\rho) \left(1 + \frac{\partial \psi}{\partial x} \right) \\ &= \rho_0 + d\rho + \rho_0 \left(\frac{\partial \psi}{\partial x} \right) + d\rho \left(\frac{\partial \psi}{\partial x} \right) \end{aligned} \quad (6.30)$$

Because we are restricting our consideration to the cases in which the density change ($d\rho$) produced by the wave is small relative to the equilibrium density (ρ_0), the term $d\rho(\partial\psi/\partial x)$ must be small compared with the term $\rho_0(d\psi/dx)$. Thus, to a reasonable approximation we can write:

$$d\rho = -\rho_0 \frac{\partial \psi}{\partial x} \quad (6.31)$$

which we can insert into Eq. 6.27, giving the following:

$$\begin{aligned}\rho_0 \left(\frac{\partial^2 \psi}{\partial t^2} \right) &= - \frac{\partial \left\{ \left(\frac{K}{\rho_0} \right) \left[-\rho_0 \left(\frac{\partial \psi}{\partial x} \right) \right] \right\}}{\partial x} \\ &= \frac{\partial \left[K \left(\frac{\partial \psi}{\partial x} \right) \right]}{\partial x}\end{aligned}\tag{6.32}$$

Rearranging makes this into an equation with a familiar form of wave equation in one dimension:

$$\rho_0 \frac{\partial^2 \psi}{\partial t^2} = K \frac{\partial^2 \psi}{\partial x^2}\tag{6.33}$$

or

$$\frac{\partial^2 \psi}{\partial x^2} = \left(\frac{\rho_0}{K} \right) \frac{\partial^2 \psi}{\partial t^2}\tag{6.34}$$

As in the case of TWs on a string, you can determine the phase speed of a pressure wave by comparing the multiplicative term in the classical wave equation of Eq. 6.35 with that in Eq. 6.34.

$$\frac{\partial^2 \psi}{\partial x^2} = \frac{1}{v^2} \frac{\partial^2 \psi}{\partial t^2}\tag{6.35}$$

Setting these factors equal to one another gives the result:

$$\frac{1}{v^2} = \frac{\rho_0}{K}\tag{6.36}$$

or

$$v = \frac{K}{\rho_0}\tag{6.37}$$

As expected, the phase speed of the pressure wave depends both on the elastic (K) and on the inertial (ρ_0) properties of the medium. Specifically, the higher the bulk modulus of the material (i.e., the stiffer the material), the faster the components of the wave will propagate (because K is in the numerator), and the higher the density of the medium, the slower those components will move (because ρ_0 is in the denominator).

6.2.2 What Are Scalar Longitudinal Waves?

Scalar longitudinal waves (SLW) are conceived of as LWs because they are sound waves. Unlike the TWs of electromagnetism, which move up and down perpendicularly to the direction of propagation, LWs vibrate in line with the direction of propagation. Transverse waves can be observed in water ripples: the ripples move up and down as the overall waves move outward, such that there are two actions—one moving up and down and the other propagating in a specific direction outward.

Technically speaking, SWs have magnitude but no direction because they are imagined to be the result of two EM waves that are 180° out of phase with one another, which leads to both signals being canceled out. This results in a kind of “pressure wave.” Mathematical physicist James Clerk Maxwell, in his original mathematical equations concerning electromagnetism, established the theoretical existence of SWs. After his death, however, later physicists assumed these equations were meaningless because SWs had not been empirically observed, and they had not been repeatedly verified among the scientific community at large.

Vibrational or subtle energetic research, however, has helped advance our understanding of SWs. One important discovery states that there are many different types of them, not just those of the EM variety. For example, there are vital SWs (corresponding with the vital or “Qi” body, described next), emotional SWs, mental SWs, causal SWs, and so forth. In essence, as far as we are aware, all “subtle” energies are made up of various types of scalar waves.

Qi Body

Qi can be interpreted as the “life energy” or “life force” that flows within us. Sometimes, it is known as the “vital energy” of the body. In traditional Chinese medicine (TCM) theory, *qi* is the vital substance constituting the human body. It also refers to the physiological functions of organs and meridians.

Some of the general properties of SWs (of the beneficial kind) include that they:

- Travel faster than the speed of light
- Seem to transcend space and time
- Cause the molecular structure of water to become coherently reordered
- Positively increase immune function in mammals
- Are involved in the formation process in nature

More details about SLW applications are discussed in next section.

6.2.3 *Scalar Longitudinal Wave Applications*

The possibility of developing a means of establishing communication through something non-homogeneous is looking very promising via use of the *More Complete Electrodynamics* (MCE) theory [4]. This theory reveals that the SLW, which is created by a *gradient-driven* current, has *no* magnetic field and, thus, is *not* constrained by the skin effect. The SLW is slightly attenuated by the non-linearities in electrical conductivity as a function of electrical field magnitude. The SLW does not interfere with classical *transverse electromagnetic* (TEM) transmission or vice versa. By contrast, TEM waves are severely attenuated in conductive media because of magnetically driven eddy currents that arise from the skin effect. Consequently, only very-low and ultra-low frequency TEM waves can be successfully used for long-distance underwater communications. The SLW also has immediate implications for the efficient redesign and optimization of existing TEM-based electronic technologies because both TEM and SLW are created simultaneously with present electronic technologies.

The goal of application of SLW-based (≥ 150 kb/s) digital data propagation over distances of many kilometers (km) to address strategic, tactical, surveillance, and undersea warfare missions of an organization such as the Navy. With this goal in mind, the optimization of SLW underwater-antenna design will be guided by development of a first principles SLW simulator from the MCE theory because all existing simulators model only circulating current-based TEM waves.

A proof-of-principle demonstration of the prototype antenna through freshwater will be conducted in-house, followed by controlled tests at typical government underwater test range(s). These tests would include characterization of wave attenuation versus frequency, modulation bandwidth, and beam-width control. The deliverable will be an initial prototype for SLW communications over tactical distances or more, followed by a field-deployable prototype, as dictated by Navy performance needs.

The unique properties of the SLW lead to more sophisticated application areas, with implications for ocean surveillance systems, underwater imaging, energy production, power transmission, transportation, guidance, and national security. This disruptive technology has the potential to transform communications, as well as electrodynamic applications in general.

As far as *low energy nuclear reactions* are concerned, it is certain that most of us have heard of scalar electrodynamics. Nevertheless, we probably have many questions about this phenomenon. Because up to now it has been mostly shrouded in mystery, we may even wonder whether electrodynamics exists at all; and if it exists, do we need exotic conditions to produce and use it? In addition, will it require a drastic transformation in our current understanding of classical electrodynamics? How much of an impact will it have on future modes of power generation and conversion? Which applications: weaponry, medical, or a low-energy fusion-driven source of energy—a D + D reaction? Most of these were mentioned in the preceding.

There is also a possibility of applying a scalar electrodynamic wave (SEW) in applications such as developing and demonstrating of an *all-electronic* (AE) engine that replaces *electromechanical* engines for vehicle propulsion. As far as other applications of SLW are concerned, a few are discussed next.

6.2.3.1 Medical Applications for Scalar Longitudinal Waves

Not all SWs, or subtle energies, are beneficial to living systems. Electromagnetism of the 60 Hz AC variety, for example, emanates a secondary longitudinal-scalar wave that is typically detrimental to living systems. To use the SLW as an application in biofield technology effectively, however, we need to cancel the detrimental aspect of wave scale and transform it into a beneficial wave; therefore, this innovative approach qualifies as a medical application of a SLW, where we can approach biomedical people to suggest such an invention and to ask for funding as well. Currently, there seems to be an interest in a biofield approach application of SLWs.

6.2.3.2 A Genuine Application of SLW for Low-Temperature Fusion Energy

In the case low-temperature fusion interactions of $D + D$, by lowering the nuclear potential barrier for purposes of “cold fusion,” for lack of better words, we know that in low-energy heavy-ion fusion, the term “Coulomb barrier” commonly refers to the barrier formed by the repulsive Coulomb and the attractive “nuclear” (nucleus-nucleus) interactions in a central (S-wave) collision. This barrier frequently is called a *fusion barrier* (for light and medium mass heavy-ion systems) or capture barrier (for heavy systems). In general, there is a centrifugal component to such a barrier (non-central collisions).

Experimenters may use the term *Coulomb barrier* to define the nominal value of the “Coulomb barrier distribution” when either coupled-channel effects operate or (at least) a collision partner is deformed because the barrier features depend on orientation. To this author’s knowledge, the terminology “transfer barrier” has not been used much. In my view it could be applied to the transfer of charged particles/clusters. There is a vast literature on methods for calculating Coulomb barriers. For instance, the double-folding method is broadly used in the low-energy nuclear physics community. Based on this technique, there is a potential, called the “Sao Paulo potential,” because it was developed by theorists in Sao Paulo, Brazil.

The Coulomb barrier is calculated theoretically by adding the nuclear and Coulomb contributions of the interaction potential. For fusion there are other contributions coming from the different degrees of freedom such as the angular momentum (centrifugal potential), the vibrational and rotational states in both interacting nuclei, in addition to the transfer contribution. This is an area for which we could approach the DOE or NRC with some Requests for Proposals—for example, the Idaho National Laboratory (INL).

6.2.3.3 Application of SLW for Directed-Energy Weapons

Scalar beam weapons were originally invented in 1904 by Nicola Tesla, an American immigrant from Yugoslavia (1856 or 1857–1943). Since his death in 1943, many nations have secretly developed his beam weapons, which now have been further refined to be so powerful that just by satellite one can create the following: a nuclear-like destruction, an earthquake, a hurricane, a tidal wave, as well as cause an instant freeze—killing every living thing instantly over many miles. An SLW also can cause intense heat, like a burning fireball over a wide area; induce hypnotic mind control of a whole population; or even remotely read anyone on the planet's mind.

Because of the nature of a pressure wave's behavior and its ability to carry tremendous energy, SLWs can remove something right out of its place in time and space faster than the speed of light, without any detectable warning, by crossing two or more beams with each other. Moreover, any target can be aimed at or even right through to the opposite side of the Earth. If either of the major scalar weapons' armed countries (e.g., U.S. or Russia) were to fire a nuclear missile to attack the other, it possibly may not even reach the target because it could be destroyed with scalar technology before it even left its place or origin. The knowledge via radio waves that it was about to be fired could be eavesdropped on and the target could be destroyed in the bunker if fired at from space by a satellite.

Above 60 Hz AC frequency, this wave can be very detrimental in nature. A scalar beam can be sent from a transmitter to the target, coupled with another sent from another transmitter, and as they cross an explosion can be initiated. This interference grid method could enable scalar beams to explode the missile before launch, as well as en route by knowing the right coordinates. If the target does manage to launch, what are known as Tesla globes, or Tesla hemispheric shields, can be sent to envelop a missile or aircraft. These are made of luminous plasma, which emanates physically from crossed scalar beams and can be created in any size, even more than 100 miles across.

Initially detected and tracked as it moves on the scalar interference grid, a continuous electromagnetic pulse (EMP) Tesla plasma globe could kill the electronics of the target. More intensely hot Tesla "fireball" globes could vaporize the missile. Tesla globes also could activate a missile's nuclear warhead en route by creating a violent low-order nuclear explosion. Various parts of flying debris can be subjected to smaller and more intense Tesla globes where the energy density to destroy is more powerful than the larger globe first encountered. This can be done in pulse mode with any remaining debris given maximum continuous heating to vaporize metals and materials. If anything still rains down on Russia or America, either could have already made a Tesla shield over the targeted area to block it from entering its airspace.

Other useful aspects of SLWs in military applications is for a community in the U.S. that believes the SWs are realizable in nature by a mathematical approach. In recent conferences sponsored by the *Institute of Electrical and Electronics Engineers* (IEEE), these were discussed openly and a *Proceedings* report of the conference exists. Dedicated to Nicola Tesla and his work, the conference's papers

presented claims that some of Tesla's work used SW concepts. Thus, there is an implied "Tesla Connection" in all of this. As was stated in the preceding, these are unconventional waves that are not necessarily a contradiction to Maxwell's equations as some have suggested; they might represent an extension to Maxwell's understanding at the time. If realizable, SLWs could represent a new form of wave propagation that may well penetrate seawater (knowing the permeability, permittivity of salt water, and consequently skin depth), resulting in a new method of submarine communications and possibly a new form of technology for *anti-submarine warfare* (ASW). This technology also helps those in the Naval Special Warfare (NSW) community, such as the Navy SEALs, to be able to communicate with each other even in murky water conditions.

The following are some mathematical notations and physics involved with this aspect of SLWs:

1. The SW, as it is understood, is not an *electromagnetic (EM) wave*. An EM wave has both electric ($\vec{E}$) and magnetic ($\vec{B}$) fields and power flow in EM waves is by means of the Poynting vector, as written in Eq. 6.38:

$$\vec{S} = \vec{E} \times \vec{B} \text{ Watts/m}^2 \quad (6.38)$$

The energy per second crossing a unit area with a normal that is pointed in the direction of $\vec{S}$ is the energy in the EM wave.

A SW has no time-varying $\vec{B}$ field. In some cases, it also has no $\vec{E}$ field. Thus, it has no energy propagated in EM wave form. It must be realized, however, that any vector could be added that may be integrated into zero over a closed surface and the *Poynting theorem* still applies. Thus, there is some ambiguity in even stating the relationship that is given by Eq. 6.38—that is, the total EM energy flow.

2. The SW could be accompanied by a vector potential $\vec{A}$, $\vec{E}$, and yet $\vec{B}$ remains zero in the far field. From EM theory, we can write as follows:

$$\begin{cases} \vec{E} = -\vec{\nabla}\phi - \frac{1}{c} \frac{\partial \vec{A}}{\partial t} \\ \vec{B} = \vec{\nabla} \times \vec{A} \end{cases} \quad (6.39)$$

In this case ϕ is the scalar (electric) potential and $\vec{A}$ is the (magnetic) vector potential. Maxwell's equations then predict the following mathematical relation:

$$\nabla^2 \phi - \frac{1}{c^2} \frac{\partial^2 \phi}{\partial t^2} = 0 \quad (\text{Scalar Potential Waves}) \quad (6.40)$$

$$\nabla^2 \vec{A} - \frac{1}{c^2} \frac{\partial^2 \vec{A}}{\partial t^2} = 0 \quad (\text{Vector Potential Waves}) \quad (6.41)$$

A solution appears to exist for the special case of $\vec{E} = 0$, $\vec{B} = 0$, and $\nabla \times \vec{A} = 0$, for a new wave satisfying the following relations:

$$\begin{cases} \vec{A} = \vec{\nabla} S \\ \phi = -\frac{1}{c} \frac{\partial S}{\partial t} \end{cases} \quad (6.42)$$

s then stratifies the following relationship:

$$\nabla^2 S - \frac{1}{c^2} \frac{\partial^2 S}{\partial t^2} \quad (6.43)$$

Note that quantity c represents the speed of light. Mathematically, s is a “potential” with a wave equation, one that suggests propagation of this wave even through $\vec{E} = \vec{B} = 0$ and the Poynting theorem indicates no EM power flow.

3. From paragraph 2 in the preceding, there is the suggestion of a solution to Maxwell’s equations involving a SW with potential s that can propagate without a *Poynting vector* EM power flow. The question arises, however, as to where the energy is drawn from to sustain such a flow of energy. Some suggest a vector that integrates to zero over a closed surface might be added in the theory. Another is the possibility of drawing energy from the vacuum, assuming net energy could be drawn from “free space.”

Quantum electrodynamics allows random energy in free space but conventional EM theory has not allowed this to date. Random energy in free space that is built of force fields that sum to zero is a possible approach. If so, these might be a source of energy to drive the s waves drawn from “free space.” A number of engineers and scientists in the community suggested, as stated earlier in a statement within this discussion, that, if realizable, the SW could represent a new form of wave propagation that could penetrate seawater or be used as a new approach for DEWs.

This author suggests considering another scenario where one may need to look at equations of the MCE theory and obtain new predictions for producing energy that way; thus, generate an SLW where the *Lagrangian density equation* for MCE can be defined as:

$$\mathcal{L} = -\frac{\epsilon c^2}{4} F_{\mu\nu} F^{\mu\nu} + J_\mu A^\mu - \frac{\gamma \epsilon c^2}{2} (\partial_\mu A^\mu) - \frac{\epsilon c^2 k^2}{2} (A_\mu A^\mu) \quad (6.44)$$

where the Lagrangian density equation written in terms of the potentials $\vec{A}$ and ϕ as follows:

$$\begin{aligned} \mathcal{L}_{EM} = \frac{\epsilon c^2}{2} & \left[\frac{1}{c^2} \left(\nabla \vec{\phi} + \frac{\partial \vec{A}}{\partial t} \right)^2 - (\nabla \times \vec{A})^2 \right] \\ & - \rho \vec{\phi} + \vec{J} \cdot \vec{A} - \frac{\epsilon c^2}{2} \underbrace{\left(\frac{1}{c^2} \frac{\partial \vec{\phi}}{\partial t} + \nabla \cdot \vec{A} \right)}_c^2 \end{aligned} \quad (6.45)$$

The proof of the equation was given in Chap. 5 of this book (see Eq. 5.126).

This is the area where there is a lot of speculation among scientists, around the community of EM and the use of the SW as a weapons application, and you will find a lot of good, as well as nonsense, approaches on the Internet by various authors. The present approach uses several approaches:

1. Acoustic signals that travel slowly (1500 m per second in seawater)
2. Blue-green laser light that has a typical range of 270 m and is readily scattered by seawater particulate
3. High-frequency radio waves that are limited to a range of 7–10 m in seawater at high frequencies
4. Extremely low-frequency radio signals that are long range (worldwide) but transmit only a few characters per second for one-way, bell-ring calls to individual submarines

The new feature of this proposed work is the use of a novel electrodynamic waves that have no magnetic field, and thus are not so severely constrained by the high conductivity of seawater, as regular radio waves are. We have demonstrated the low loss property of this novel (scalar–longitudinal) wave experimentally by sending a video signal through 2 mm of solid copper at 8 GHz.

If you have a background in physics or electrical engineering, you know that, unquestionably, our knowledge of the properties and dynamics of EM systems is believed to be the most solid and firmly established in all classical physics. By its extension, the application of quantum electrodynamics, describing accurately the interaction of light and matter at the subatomic realms, has resulted in the most successful theoretical scientific theory to date, agreeing with corresponding experimental findings to astounding levels of precision. Accordingly, these developments have led to the belief among physicists that the theory of classical electrodynamics is complete and that it is essentially a closed subject.

Nevertheless, at least as far back to the era of *Nikola Tesla*, there have been continual rumblings of discontent stemming from occasional physical evidence from both laboratory experimental protocols and knowledge obtained from observation of natural phenomena (e.g., the dynamics of atmospheric electricity, etc.). This may suggest that in extreme situations involving the production of high energies at specific frequencies, there might be some cracks exposed in the supposed impenetrable monolithic fortress of classical/quantum electrodynamics, implying possible key missing theoretical and physical elements. Unfortunately, some of these experimental phenomena have been difficult to replicate and produce on-demand. Moreover, some have been shown apparently to violate some of the established principles underlying classical thermodynamics.

On top of that, many of those courageous individuals promoting the study of this phenomenon have couched their understanding of the limited reliable experimental evidence available from the sources in a language unfamiliar to the legion of mainstream technical specialists in electrodynamics, preventing clear communication of these ideas. Also, the various sources that have sought to convey this information have at times delivered contradictory statements.

It is therefore no wonder that for many decades such exotic claims have been disregarded, ignored, and summarily discounted by mainstream physics. Because of important developments over the past 2 years, however, there has been a welcome resurgence of research in this area, bringing back renewed interest toward the certifying the existence of these formerly rejected anomalous energy phenomena. Consequently, this renaissance of the serious enterprise in searching for specific weaknesses, which currently plague a fuller understanding of electrodynamics, has propelled the proponents of this research to more systematically outline their ideas in a clearer fashion. The possible properties of these dynamics and how inclusion of them could change our current understanding of electricity and magnetism, as well as suggest implications for potential vast, practical ramifications, may change the disciplines of physics, engineering, and energy generation.

It is the purpose of this book, and particularly this chapter, not only to report on these unique, various recent inventions and their possible modes of operation, but also to convince readers of their value for hopefully directing a future program geared toward the rigorous clarification and certification of the specific role the electroscalar domain might play in shaping a future consistent, classical electrodynamics. Also, by extension, to perhaps shed light on thorny conceptual and mathematical inconsistencies that do exist in the present interpretation of relativistic quantum mechanics. In this regard, it is anticipated that, by incorporating this more expansive electrodynamic model, the source of the extant problems with gauge invariance in quantum electrodynamics and the subsequent unavoidable divergences in energy/charge might be identified and ameliorated. Not only does the electroscalar domain have the potential to address such lofty theoretical questions surrounding fundamental physics, but it also aims to show that the protocol necessary for generating these field effects may not be present in exotic conditions relating to large field strengths and specific frequencies involving an expensive infrastructure such as the *large hadron collider* (LHC).

Insight into the incompleteness of classical electrodynamics can begin with the Helmholtz theorem, which states that any sufficiently smooth three-dimensional vector field can be uniquely decomposed into two parts. By extension, a generalized theorem exists that was certified through the scholarly work of physicist-mathematician Dale Woodside (2009) [5] (see Eq. 6.44 in the preceding) for unique decomposition of a sufficiently smooth, Minkowski 4-vector field (three spatial dimensions, plus time) into *four-irrotational* and *four-solenoidal* parts, together with the tangential and normal components on the bounding surface.

With this background, the theoretical existence of the electroscalar wave can be attributed to failure to include certain terms in the standard, general four-dimensional EM Lagrangian density related to the four-irrotational parts of the vector field. Here ϵ is electrical permittivity, not necessarily of the vacuum. Specifically, the electroscalar field becomes incorporated into the structure of electrodynamics when we get Eq. 6.44 in the preceding for $\gamma = 1$ and $k = (2\pi m c/h) = 0$. As we can see in this representation, as written in the equation, it is the presence of the third term that describes these new features.

We can see more clearly how this term arises by writing Lagrangian density in terms of the standard EM scalar (ϕ) (see Eq. 6.45) and magnetic vector potentials ($\vec{A}$), without the electroscalar representation included. This equation has zero divergence of the potentials (formally called solenoidal) consistent with classical electromagnetics, as we can see here. The second class of 4-vector fields has zero curl of the potentials (i.e., irrotational vector field), which will emerge once we add this scalar factor. Here this is represented by the last term, which is usually zero in standard classical EMs. The expression in the parentheses, when set equal to zero, describes what is known as the Lorentz condition that makes the scalar potential and the vector potential in their usual form mathematically *dependent* on each other.

Accordingly, the usual EM theory then specifies that the potentials may be chosen arbitrarily based on the specific so-called gauge that is chosen for this purpose. The MCE theory, however, allows for a non-zero value for this scalar-valued expression, essentially making the potentials *independent* of each other, where this new scalar-valued component (C in Eq. 6.45 that can be called *Lagrangian density*) is a dynamic function of space and time. It is this new idea of the independence of the potentials, out of which the scalar value (C) is derived, and from which the unique properties and dynamics of the *scalar longitudinal electrodynamic* wave arises.

To put all these in perspective, a more complete electrodynamic model can be derived from this last equation of the Lagrangian density. The Lagrangian expression is important in physics because invariance of the Lagrangian under any transformation gives rise to a conserved quantity. Now, as is well known, conservation of charge-current is a fundamental principle of physics and nature. Conventionally, in classical electrodynamics charged matter creates an $\vec{E}$ field. Motion of charged matter creates a magnetic $\vec{B}$ field from an electrical current that in turn influences the $\vec{B}$ and $\vec{E}$ fields.

Before, we continue further, let us write the following equations:

$$\vec{E} = -\nabla\phi - \frac{\partial\vec{A}}{\partial t} \quad \text{Relativistic Covariance} \quad (6.46)$$

$$\vec{B} = \nabla \times \vec{A} \quad \begin{array}{l} \text{Classical Fields } (\vec{B} \text{ and } \vec{E}) \\ \text{in terms of usual classical} \\ \text{potentials } (\vec{A} \text{ and } \vec{\phi}) \end{array} \quad (6.47)$$

$$C = \frac{1}{c^2} \frac{\partial\vec{\phi}}{\partial t} + \nabla \cdot \vec{A} \quad \text{Classical wave equations for } \vec{A}, \vec{B} \quad (6.48)$$

$$\nabla \times \vec{B} - \frac{1}{c^2} \frac{\partial\vec{E}}{\partial t} - \nabla C = \mu\vec{J} \quad \vec{E} \text{ and } \vec{\phi} \text{ without the use of a gauge} \quad (6.49)$$

$$\nabla \cdot \vec{E} + \frac{\partial C}{\partial t} = \frac{\rho}{\epsilon} \quad \begin{array}{l} \text{Condition (the MCE theory produces} \\ \text{cancellation of } \partial C / \partial t \text{ and } -\vec{\nabla} C \text{ in the} \\ \text{classical wave equation for } \vec{\phi} \text{ and } \vec{A}, \\ \text{thus eliminating the need for a gauge} \\ \text{condition)} \end{array} \quad (6.50)$$

These effects can be modeled by Maxwell's equations. Now, exactly how and to what degree do these equations change when the new scalar-valued C field is incorporated? Those of you who have knowledge of Maxwellian theory will notice that the two homogeneous Maxwell's equations—representing *Faraday's Law* and $\vec{\nabla} \cdot \vec{B}$, the standard Gauss's Law equation for a divergence-less magnetic field—are both unchanged from the classical model. Notice the last three equations incorporate this new scalar component that is labeled C .

This formulation, as defined by Eq. 6.48, creates a somewhat revised version of Maxwell's equations, with one new term, $-\vec{\nabla} C$, in Gauss's Law (Eq. 6.50), where ρ is the charge density, and one new term $(\partial C / \partial t)$ in Ampère's Law (Eq. 6.49), where J is the current density. We can see that these new equations lead to some important conditions. First, relativistic covariance is preserved. Second, unchanged are the classical fields $\vec{E}$ and $\vec{B}$ in terms of the usual classical potentials ($\vec{A}$ and $\vec{\phi}$). We have the same classical wave equations for $\vec{A}$, $\vec{\phi}$, $\vec{E}$, and $\vec{B}$ *without* the use of a gauge condition (and its attendant incompleteness) because the MCE theory shows cancellation of $\partial C / \partial t$ and $-\nabla C$, the classical wave equations for $\vec{\phi}$ and $\vec{A}$; and a SLW is revealed, composed of the scalar and longitudinal-electric fields.

A wave equation for C is revealed by use of the time derivative of Eq. 6.50, added to the divergence of Eq. 6.49. Now, as is known, matching conditions at the interface between two different media are required to solve Maxwell's equations. The divergence theorem on Eq. 6.51 will yield an interface matching in the normal component (“ $\hat{n}$ ”) of $\nabla C / \mu$, as shown in Eq. 6.51:

$$\frac{\partial^2 C}{\partial c^2 t^2} - \nabla^2 C \equiv \square^2 C = \mu \left(\frac{\partial \rho}{\partial t} + \nabla \cdot \vec{J} \right) \quad (6.51)$$

$$\left(\frac{\nabla C}{\mu} \right)_{1n} = \left(\frac{\nabla C}{\mu} \right)_{2n} \quad (6.52)$$

$$C = C_0 \exp[j(kr - \omega t)] / r \quad (6.53)$$

Note: The preceding, Eqs. 6.51 and 6.52, present a wave equation for *scalar factor* C matching condition in the normal component of $\nabla C / \mu$, spherically *symmetric wave solution* for C , and the operator $\square^2 = \frac{\partial^2}{\partial c^2 t^2} - \nabla^2$, called the *d' Alembert operator*.

The subscripts in Eq. 6.51 denote $\nabla C / \mu$ in medium 1 or medium 2, respectively; (μ) is magnetic permeability—again not necessarily that of the vacuum. In this regard, with the vector potential ($\vec{A}$) and scalar potential ($\vec{\phi}$) now stipulated as independent of each other, it is the surface charge density at the interface that produces a discontinuity in the gradient of the scalar potential, rather than the standard discontinuity in the normal component of $\vec{E}$ (see Hively [4]).

Notice, also from Eq. 6.51, the source for scalar factor C implies a violation of charge conservation on the non-zero right-hand side (RHS), a situation that we noted cannot exist in macroscopic nature. Nevertheless, this will be compatible with standard Maxwellian theory if this violation occurs at very short time scales, such

as in subatomic interactions. Now, interestingly, with the stipulation of charge conservation on large time scales, giving zero on the RHS of Eq. 6.51., longitudinal wave-like solutions are produced with the lowest order form in a spherically symmetric geometry at a distance (r), $C = C_0 \exp [j(kr - \omega t)]/r$. Applying the boundary condition $C \rightarrow 0$ as $r \rightarrow \infty$ is thus trivially satisfied. The C wave therefore is a *pressure* wave, similar to that in acoustics and hydrodynamics.

This is unique under the new MCE model because, although classical electrodynamics forbids a spherically symmetric TW to exist, this constraint will be absent under MCE theory. Also, an unprecedented result is that these longitudinal C waves will have energy but no momentum. But then again, this is not unlike charged particle–anti-particle fluctuations that also have energy but no net momentum.

Now that we are here so far, the question of why this constraint prohibiting a spherically symmetric wave is lifted in MCE can be resolved in the following sets of Eq. 6.54 for the wave equation for the vertical magnetic field:

$$\begin{cases} \frac{1}{c^2} \frac{\partial^2 \vec{B}}{\partial t^2} - \nabla^2 \vec{B} = \mu_0 (\nabla \times \vec{J}) \\ \nabla \times \vec{J} = 0 \rightarrow J = \nabla k \\ \text{Gradient-driven Current} \rightarrow \text{SLW} \end{cases} \quad (6.54)$$

The sets of this equation are established for the $\vec{B}$ wave equation, resulting in a gradient-driven current in MCE for generating the SLW.

Notice again that the source of the magnetic field on the RHS is a non-zero value of $\nabla \times \vec{J}$, which signifies solenoidal current density, as is the case in standard Maxwellian theory. When $\vec{B}$ is zero, so is $\nabla \times \vec{J}$. This is an important result. Thus, the current density is irrotational, which implies that $J = \nabla \kappa$. Here κ is a scalar function of space and time. Therefore, in contrast to closed current paths generated in ordinary Maxwell theory that result in classical waves that arise from a solenoidal current density ($\nabla \times \vec{J} \neq 0$), J for the SLW is gradient-driven and may be uniquely detectable.

We also can see from this result that a zero value of the magnetic field is a necessary and sufficient condition for this gradient-driven current. Now, because in linearly conductive media, the current density ($\vec{J}$) is directly proportional the electric field intensity ($\vec{E}$) that produced it, where σ is the conductivity, this gradient-driven current will then produce a longitudinal $\vec{E}$ -field. Based on calculations so far, we can establish a wave equation for the $\vec{E}$ solution for longitudinal $\vec{E}$ in MCE *spherically symmetric wave solutions* for $\vec{E}$ and $\vec{J}$ in *linearly conductive media*, as follows:

$$\frac{\partial^2 \vec{E}}{\partial c^2 t^2} - \nabla^2 \vec{E} = \left(\frac{\partial^2}{\partial c^2 t^2} - \nabla^2 \right) \vec{E} \equiv \square^2 \vec{E} = \mu \frac{\partial \vec{J}}{\partial t} - \frac{\nabla \rho}{\epsilon} \quad (6.55)$$

$$E = E_r \hat{r} \exp[j(kr - \omega t)]/r \quad (6.56)$$

$$\vec{J} = \sigma \vec{E} \rightarrow \square^2 \vec{J} = 0 \quad (6.57)$$

We also can see this from examining the standard vectorial wave equation for the electric field. The wave equation for $\vec{E}$ (Eq. 6.50) arises from the curl of Faraday's Law, use of $\nabla \cdot \vec{B}$ from Ampère's Law (Eq. 5.49), and the substitution of $\nabla \cdot \vec{E}$ from Eq. 6.50 with cancellation of the terms $\nabla(\partial C/\partial t) = (\partial/\partial t)\nabla C$. When the RHS of Eq. 13 is zero, the lowest order outgoing spherical wave is $E = E_r \hat{r} \exp[j(kr - \omega t)]/r$, where $\hat{r}$ represents the unit vector in the radial direction and r represents the radial distance. The electrical field is also longitudinal. Substitution of $\vec{J} = \sigma \vec{E}$ into $\square^2 \vec{E} = 0$ results in $\square^2 \vec{J} = 0$, meaning that the current density is also radial. The SLW equations for E and J are remarkable for several reasons.

First, the vector SLW equations for $\vec{E}$ and $\vec{J}$ are fully captured in one wave equation for the scalar function (κ), $\square^2 \kappa = 0$. Second, these forms are like $\square^2 C = 0$. Third, these equations have zero on the RHS for propagation in conductive media. This occurs because $\vec{B} = 0$ for the SLW, implying no back EM field from $\partial \vec{B}/\partial t$ in Faraday's Law, which in turn gives no circulating eddy currents. Experimentation has shown that the SLW is not subject to the skin effect in media with linear electric conductivity and travels with minimum resistance in any conductive media.

This last fact affords some insight into another related ongoing conundrum in condensed matter physics—the mystery surrounding high-temperature superconductivity (HTS). As we know, the physical problem of HTS is one of the major unsolved problems in theoretical condensed matter physics—in part, because the materials are somewhat complex, multilayered crystals. Here the MCE theory may provide an explanation on the basis of gradient-driven currents between (or among) the crystal layers. The new MCE Hamiltonian (Eq. 6.58) includes the SLW because of gradient-driven currents among the crystalline layers as an explanation for HTS. The electrodynamic Hamiltonian for MCE is written:

$$\mathcal{H}_{EM} = \left(\frac{\epsilon E^2}{2} + \frac{B^2}{2\mu} \right) + (\rho - \epsilon \nabla \cdot \vec{E}) \vec{\phi} - \vec{J} \cdot \vec{A} + \frac{C^2}{2\mu} + \frac{C \nabla \cdot \vec{A}}{\mu} \quad (6.58)$$

In conclusion we can build an antenna based on the preceding concept within a laboratory environment and use a simulation software such as Multi-Physics COMSOL[®] or ANSYS computer code to model such an antenna. We believe, however, that we have done adequate analysis in this chapter to show the field of electrodynamics (classical and quantum), although considered to be totally understood, with any criticisms of incompleteness on the part of dissenters essentially taken as veritable heresy; nevertheless, it needs reevaluation in terms of apparent unfortunate sins of omission in the failure to include an electrosalar component.

Anomalies previously not completely understood may get a boost of new understanding from the operation of electrosalar energy. We have seen in the three instances examined—the mechanism of generation of seismic precursor electrical

signals because of the movement of the Earth's crust, the ordinary peeling of adhesive tape, and irradiation by the special TESLAR chip—the common feature of the breaking of chemical bonds. In fact, we ultimately may find that any phenomenon requiring the breaking of chemical bonds, in either inanimate or biological systems, actually may be mediated by SWs.

Thus, we may discover that the scientific disciplines of chemistry or biochemistry may be more closely related to physics than is currently thought. Accordingly, the experimental and theoretical reevaluation of even the simplest phenomena in this regard, such as triboelectrification processes, is the absolute essence for those researchers knowledgeable of the necessity for this reassessment of EMs. As this author said in the introduction, it may even turn out that the gradient-driven current and associated SLW could be the umbrella concept under which many of the currently unexplained electrodynamic phenomena are discussed frequently at conferences, yielding a satisfying explanation.

The new SLW patent itself, which is the centerpiece of this chapter, is a primary example of the type of invention that probably would not have seen the light of day even 10 years ago. As previously mentioned, we are seeing more of this inspired breakthrough technology based on operating principles, formerly viewed with rank skepticism bordering on haughty derision by mainstream science, now surfacing to provide an able challenge to the prevailing worldview by reproducible corroborating tests by independent sources. This revolution in the technological witnessing of the overhaul of current orthodoxy is definitely a harbinger of the rapidly approaching time when many of the encrusted and equally ill-conceived, but still accepted, paradigms of science—thought to underpin our sentient reality—will fall by the wayside.

On a grander panoramic scale, our expanding knowledge gleaned from further examining the electroscalar wave concept, as applied to areas of investigation (e.g., cold fusion research, overunity power sources, etc.), explicitly will shape the future of society as well as science, especially concerning our openness to phenomena that challenge current belief systems.

To this point the incompleteness in established understanding of the properties of electrodynamic systems can be attributed to the failure to properly incorporate what can be termed the electroscalar force into the structural edifice of electrodynamics. Unbeknownst to most specialists in the disciplines mentioned, over the last decade in technological development circles, there has quietly, but inexorably, emerged bona fide physical evidence of the demonstration of the existence of SLW dynamics in recent inventions and discoveries.

As technology leads to new understanding, we are rapidly approaching a time in which these findings no longer can be pushed aside or ignored by orthodox physics, and physics must come to terms with their potential physical and philosophical impacts on our world. By the time you read this book, this author thinks you might agree with the fact that many could be on the brink of a new era in science and technology, the likes of which the current generation has never seen before. Despite what mainstream physics may claim, the study of electrodynamics is by no means a closed book. Further details are provided in the following sections.

6.3 Description of the $\vec{B}^{(3)}$ Field

During the investigation of the theory optically induced line shifts in *nuclear magnetic resonance* (NMR), people have come across the result that the anti-symmetric part of the intensity tensor of light is directly proportional in free space to an entirely novel, phase-free, magnetic field of light, which was identified as the $\vec{B}^{(3)}$ field, as defined in Eq. 6.59a. The presence of $\vec{B}^{(3)}$ in free space shows that the usual, propagating TWs of EM radiation are linked geometrically to spin field $\vec{B}^{(3)}$, which indeed emerges directly from the fundamental, classic equation of motion of a single electron in a circularly polarized light beam [6]:

$$\vec{B}^{(1)} \times \vec{B}^{(2)} = iB^{(0)} \vec{B}^{(3)} \quad (6.59a)$$

$$\vec{B}^{(2)} \times \vec{B}^{(3)} = iB^{(0)} \vec{B}^{(1)*} \quad (6.59b)$$

$$\vec{B}^{(3)} \times \vec{B}^{(1)} = iB^{(0)} \vec{B}^{(2)*} \quad (6.59c)$$

Note that the symbol * means conjugate form of the field, and superscripts (1), (2), and (3) can be permuted to give the other two equations, Eqs. 6.1; thus, the fields $\vec{B}^{(1)}$, $\vec{B}^{(2)}$, and $\vec{B}^{(3)}$ are simply components of the *magnetic flux density of free space* electromagnetism on a circular, rather than on a Cartesian, basis. In quantum field theory, longitudinal component $\vec{B}^{(3)}$ becomes the fundamental photomagnetic of light, and the operator is defined by the following relationship [7–12]:

$$\widehat{B}^{(3)} = B^{(0)} \frac{\widehat{P}}{\hbar} \quad (6.60)$$

where $\widehat{P}$ is the angular momentum operator of one photon. The existence of the longitudinal $\widehat{B}^{(3)}$ in free space is indicated experimentally by optically induced nuclear magnetic resonance (NMR) shifts and by several well-known phenomena of magnetization by light—for example, the *inverse Faraday effects*.

The core logic of Eq. 6.59a asserts that there exists a novel cyclically symmetric field algebra in free space, implying that the usual transverse solutions of Maxwell's equations are tied to the longitudinal, non-zero, real, and physical magnetic flux density $\vec{B}^{(3)}$, which we name the *spin field*. This deduction fundamentally changes our current appreciation of electrodynamics and therefore the principles on which the old quantum theory was derived—for example, the Planck Law [13] and the light quantum hypothesis proposed in 1905 by Einstein.

The belated recognition of $\vec{B}^{(3)}$ implies that there is a magnetic field in free space that is associated with the longitudinal space axis, z , which is labeled (3) in the circular basis. Conventionally, the radiation intensity distribution is calculated using only two transverse degrees of freedom, right and left circular, corresponding to (1) and (2) on the circular basis. The $\vec{B}^{(3)}$ field of vacuum electromagnetism introduces a new paradigm of the field theory, summarized in the cyclically

symmetric equations linking it to the usual transverse magnetic plane wave components, $\vec{B}^{(1)} = \vec{B}^{(2)*}$ [6, 14, 15].

In January 1992 at Cornell University the $\vec{B}^{(3)}$ field was first and obliquely inferred from a careful reexamination of known magneto-optics phenomena [16, 17] that had previously been interpreted by convention through the conjugate product $\vec{E}^{(1)} = \vec{E}^{(2)*}$ of electric plane-wave components $\vec{E}^{(1)} = \vec{E}^{(2)*}$. In the intervening three-and-a-half years its understanding developed substantially into monographs and papers [6, 14, 15] covering several fundamental aspects of field theory.

The $\vec{B}^{(3)}$ field produces magnetization in an electron plasma that is proportional to the square root of the power density dependence of the circularly polarized EM radiation—conclusive evidence for the presence of the phase-free $\vec{B}^{(3)}$ in the vacuum. There are many experimental consequences of this finding, some of which are of practical utility (e.g., optical NMR). Nevertheless, the most important theoretical consequence is that there exist longitudinal components in free space of EM radiation—a conclusion that is strikingly reminiscent of that obtained from the theory of finite photon mass.

The two ideas are interwoven throughout this book. The characteristic square root light intensity dependence of $\vec{B}^{(3)}$ dominates and theoretically is observable at low cyclotron frequencies when intense, circularly polarized EM radiation interacts with a single electron, or in practical terms an electron plasma or beam. The magnetization induced in such an electron ensemble by circularly polarized radiation therefore is expected to be proportional to the square root of the power density (i.e., the intensity in W/m^2) of the radiation. This result emerges directly from the fundamental, classic equation of motion of one electron in the beam, the *relativistic Hamilton-Jacobi equation*.

To establish the physical presence of $\vec{B}^{(3)}$ in the vacuum therefore requires the observation of this magnetization as a function of the beam’s power density, a critically important experiment. Other possible experiments to detect $\vec{B}^{(3)}$, such as the optical equivalent of the *Aharonov-Bohm effect*, are suggested throughout the book.

More details about this section’s subject in can be found in the references listed at the end of this chapter. Further details are beyond the scope of this book; thus, we encourage readers to refer to [4, 6–12, 14, 17].

6.4 Scalar Wave Description

What is a “scalar wave” exactly? A *scalar wave* (SW) is just another name for a *longitudinal wave* (LW). The term “scalar” is sometimes used instead because the hypothetical source of these waves is thought to be a “scalar field” of some kind, similar to the Higgs field, for example. In general, the definition of the LW falls into the following description: “A wave motion, in which the particles of the medium

oscillate about their mean positions in the direction of propagation of the wave, is called longitudinal wave.”

For LWs the vibration of the particles of the medium are in the direction of wave propagation. An LW proceeds in the form of compression and rarefaction, which is the stretch and compression in the same direction as the wave moves. For an LW at places of compression the pressure and density tend to be maximum, while at places where rarefaction takes place, the pressure and density are minimum. In gases only, LWs can propagate. Longitudinal waves also are known as compression waves.

There is nothing particularly controversial about LWs in general. They are a ubiquitous and well-acknowledged phenomenon in nature. Sound waves traveling through the atmosphere (or underwater) are longitudinal, as are plasma waves propagating through space (also known as *Birkeland currents*). Longitudinal waves moving through the Earth’s interior are known as *telluric currents*. They can all be thought of as pressure waves of sorts.

Scalar and longitudinal waves are quite different from “transverse” waves. You can observe a *transverse wave* by plucking a guitar string or watching ripples on the surface of a pond. They oscillate (i.e., vibrate, move up and down or side-to-side) perpendicular to their arrow of propagation (i.e., directional movement), as shown in Fig. 6.6.

In modern-day electrodynamics (both classical and quantum), EM waves traveling in “free space” (e.g., photons in vacuum) generally are considered to be TW. Nonetheless, this was not always the case. When the preeminent mathematician James Clerk Maxwell first modeled and formalized his unified theory of electromagnetism in the late nineteenth century, neither the EM SW or LW nor the EM TW had been experimentally proved, but he had postulated and calculated the existence of both.

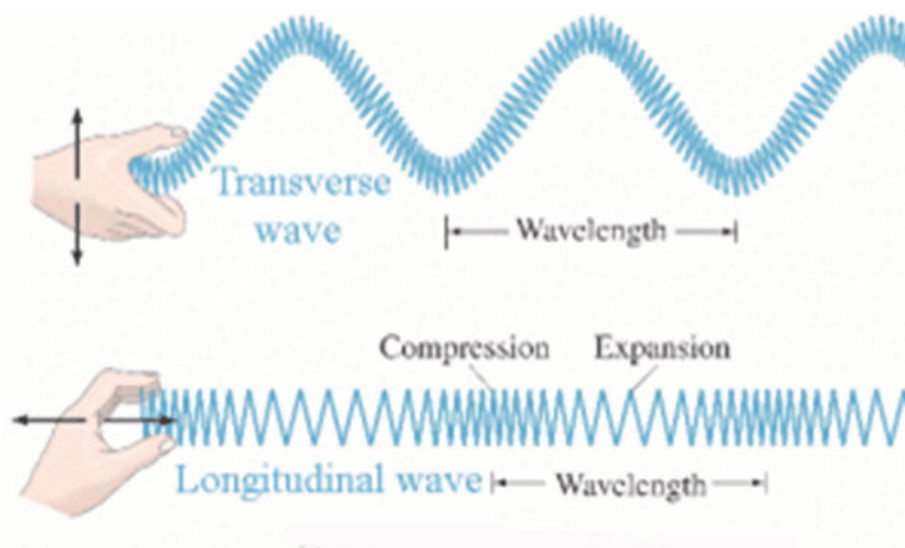


Fig. 6.6 Illustration of a transverse wave and a longitudinal wave

After Heinrich Hertz demonstrated experimentally the existence of transverse radio waves in 1887, theoreticians (e.g., Heaviside, Gibbs, and others) went about revising Maxwell's original equations; he was now deceased and could not object. They wrote out the SW/LW component from the original equations because they felt the mathematical framework and theory should be made to agree only by experimenting. Obviously, the simplified equations worked; they helped make the AC/DC electrical age engineerable. But at what expense?

Then in the 1889 Nikola Tesla, a prolific experimental physicist and inventor of the *alternating current* (AC) theory, threw a proverbial wrench in the works when he discovered experimental proof for the elusive electric SW. This seemed to suggest that SW/LW, opposed to TW, could propagate as pure electric waves or as pure magnetic waves. Tesla also believed these waves carried a hitherto-unknown form of excess energy he referred to as "radiant." This intriguing and unexpected result was said to have been verified by Lord Kelvin and others soon after.

Instead of merging their experimental results into a unified proof for Maxwell's original equations, however, Tesla, Hertz, and others decided to bicker and squabble over who was more correct. In actuality they all derived correct results. Nevertheless, because humans (even "rational" scientists) are fallible and prone to fits of vanity and self-aggrandizement, each side insisted dogmatically that they were correct and the other side was wrong.

The issue was allegedly settled after the dawn of the twentieth century when: (1) the concept of the mechanical (passive/viscous) ether was purportedly disproved by Michelson-Morley and replaced by Einstein's Relativistic Spacetime Manifold, and (2) detection of SW/LWs proved much more difficult than initially thought (mostly because of the waves' subtle densities, fluctuating frequencies, and orthogonal directional flow). As a result, the truncation of Maxwell's equations was upheld.

The SW and LW in free space, however, are quite real. Beside Tesla, empirical work carried out by electrical engineers (e.g., Eric Dollard, Konstantin Meyl, Thomas Imlauer, and Jean-Louis Naudin, to name only some) clearly have demonstrated their existence experimentally. These waves seem able to exceed the speed of light, pass through EM shielding (also known as Faraday cages), and produce overunity (more energy out than in) effects. They seem to propagate in a yet unacknowledged counterspatial dimension (also known as hyper-space, pre-space, false-vacuum, Aether, implicit order, etc.).

Because the concept of an all-pervasive material ether was discarded by most scientists, the thought of vortex-like electric and/or magnetic waves existing in free space, without the support of a viscous medium, was thought to be impossible. Nevertheless, later experiments carried out by Dayton Miller, Paul Sagnac, E. W. Silvertooth, and others have contradicted the findings of Michelson and Morley. More recently Italian mathematician-physicist Daniele Funaro, American physicist-systems theorist Paul LaViolette, and British physicist Harold Aspden have all conceived of (and mathematically formulated) models for a free space ether that is dynamic, fluctuating, self-organizing, and allows for the formation and propagation of SWs and LWs.

With the appearance of experiments on the non-classical effects of electrodynamics, authors often speak of EM waves not being based on oscillations of electric and magnetic fields. For example, it is claimed that there is an effect of such waves on biological systems and the human body. Even medical devices are sold that are assumed to work on the principle of transmitting any kind of information via “waves” that have a positive effect on human health. In all cases the explanation of these effects is speculative, and even the transmission mechanism remains unclear because there is no sound theory about such waves, often subsumed under the notion SWs. We have tried to give a clear definition of certain types of waves that can explain the observed effects [18].

Before analyzing the problem in more detail, we must distinguish between SWs that contain fractions of ordinary electric and magnetic fields and such waves that do not and therefore appear even more obscure. Often SWs are assumed to consist of longitudinal fields. In ordinary Maxwellian electrodynamics such fields do not exist, and EM radiation is said always to be transversal. In modern unified physics’ approaches (e.g., *Einstein-Cartan-Evans theory* [19, 20], however, it has been shown that the polarization directions of EM fields do exist in all directions of four-dimensional space. So in the direction of transmission, an ordinary EM wave has a longitudinal magnetic component, the so-called $\vec{B}^{(3)}$ field of Evans [21]. (See Sect. 6.4 for more details about the $\vec{B}^{(3)}$ field.)

The $\vec{B}^{(3)}$ field is detectable by the so-called inverse Faraday effect that has been known experimentally since the 1960s [22]. Some experimental setups (e.g., the “magnifying transmitter” of Tesla [16, 17]) claim to use these longitudinal components. They can be considered to consist of an extended resonance circuit where the capacitor plates each have been displaced to the transmitter and receiver site (Fig. 6.7). In an ordinary capacitor (or cavity resonator), a very high-frequency

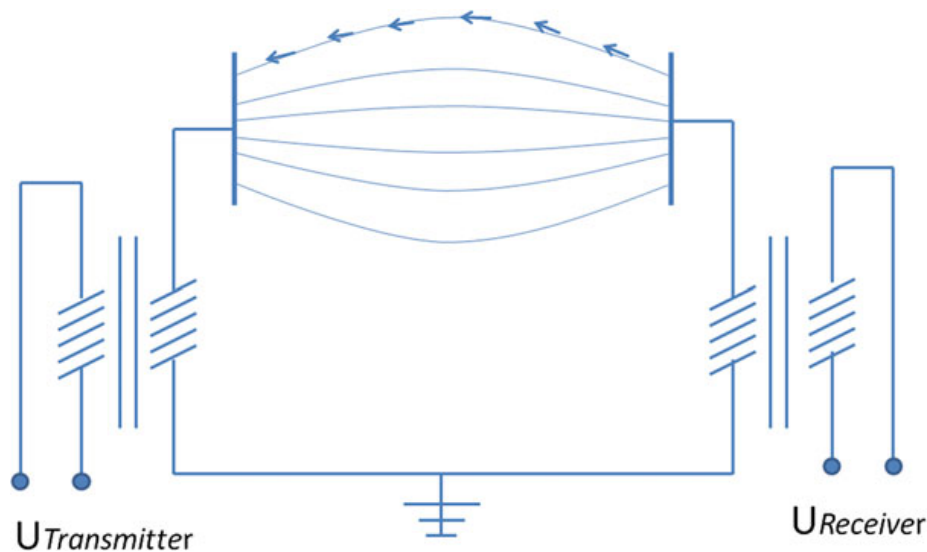


Fig. 6.7 Propagation of longitudinal electric waves according to Tesla

wave (GHz or THz range) leads to significant runtime effects of the signal so that the quasistatic electric field can be assumed to be cut into pulses.

These represent the near-field of an EM wave and may be thought of as longitudinal. For lower frequencies the electric field between the capacitor plates remains quasistatic and therefore longitudinal too. We do not want to go deeper into this subject here. Having given hints for the possible existence of longitudinal electric and magnetic fields, we leave this area of concentration on mechanisms that allow transmission of signals, even without any detectable EM fields, to readers.

Before we move on with more details about the SW, we need to lay groundwork about the types of waves and where the SW falls under that category; thus, we need to have some idea about TWs and LTWs and what their descriptions are. This subject was discussed in a previous section of this chapter quite extensively; however, we describe the subject of longitudinal potential waves in the next section.

6.5 Longitudinal Potential Waves

In the following we develop the theory of EM waves with vanishing field vectors. Such a field state normally is referred to as a *vacuum state* and was described in full relativistic detail by Eckardt and Lindstrom [22]. Vacuum states also play a role in the microscopic interaction with matter. Here we restrict consideration to ordinary electrodynamics to give engineers a chance to fully understand the subject.

With $\vec{E}$ and $\vec{B}$ designating the classical electric and magnetic field vectors, a vacuum state is defined by:

$$\vec{E} = 0 \quad (6.61)$$

$$\vec{B} = 0 \quad (6.62)$$

The only possibility to find EM effects then is by the potentials. These are defined as vector and scalar potentials to constitute the “force” fields $\vec{E}$ and $\vec{B}$:

$$\vec{E} = -\nabla U - \dot{\vec{A}} \quad (6.63)$$

$$\vec{B} = \vec{\nabla} \times \vec{A} \quad (6.64)$$

with electric scalar potential U and magnetic vector potential $\vec{A}$. The dot above the $\vec{A}$ in Eq. 6.63 denotes the time derivative. The vacuum conditions, as stated in Eqs. 6.61 and 6.62, will lead to the following sets of equations:

$$\nabla U = -\dot{\vec{A}} \quad (6.65)$$

$$\vec{\nabla} \times \vec{A} = 0 \quad (6.66)$$

From Eq. 6.66 it follows immediately that the *vector potential* is vortex-free, representing a laminar flow. The gradient of the scalar potential is coupled to the time derivative of the vector potential, so both are not independent of one another. A general solution of these equations was derived by Eckardt and Lindstrom [22]. This is a wave solution where $\vec{A}$ is in the direction of propagation (i.e., this is a LW). Several wave forms are possible, which may even result in a propagation velocity different from the speed of light c . As a simple example, we assume a sine-like behavior of vector potential $\vec{A}$:

$$\vec{A} = \vec{A}_0 \sin(\vec{k} \cdot \vec{x} - \omega t) \quad (6.67)$$

with direction of propagation $\vec{k}$ (wave vector), space coordinate vector $\vec{x}$, and time frequency ω . Then it follows from Eq. 6.66 that

$$\nabla U = \vec{A}_0 \omega \cos(\vec{k} \cdot \vec{x} - \omega t) \quad (6.68)$$

This condition must be met for any potential U . We make the approach as follows:

$$U = U_0 \sin(\vec{k} \cdot \vec{x} - \omega t) \quad (6.69)$$

to find that:

$$\nabla U = k U_0 \cos(\vec{k} \cdot \vec{x} - \omega t) \quad (6.70)$$

which, compared to Eq. 6.68, defines the constant $\vec{A}_0$ to be:

$$\vec{A}_0 = \vec{k} \left(\frac{U_0}{\omega} \right) \quad (6.71)$$

Obviously, the waves of $\vec{A}$ and U have the same phase.

Next, we consider the energy density of such a combined wave. This is given in general by:

$$w = \frac{1}{2} \epsilon_0 \vec{E}^2 + \frac{1}{2\mu_0} \vec{B}^2 \quad (6.72)$$

From Eqs. 6.65 and 6.66 it can be seen that the magnetic field disappears identically, but the electric field is a vanishing sum of two terms that are different from zero.

These two terms evoke an energy density of space where the wave propagates. This cannot be obtained out of the force fields (these are zero) but must be computed from the constituting potentials. As discussed in a paper by Eckardt and Lindstrom [20], we must write:

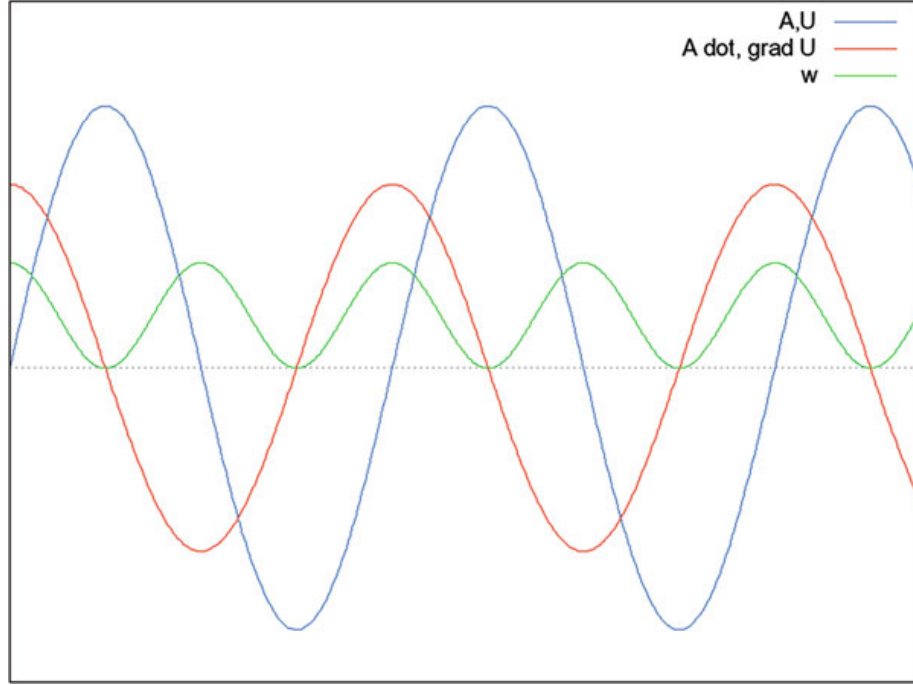


Fig. 6.8 Phases of potentials $\vec{A}$ and U and energy density w

$$w = \frac{1}{2} \epsilon_0 \left(\dot{\vec{A}}^2 + (\nabla U)^2 \right) \quad (6.73)$$

With Eq. 6.67 and Eq. 6.69, it follows that:

$$w = \epsilon_0 k^2 U_0^2 \cos^2(\vec{k} \cdot \vec{x} - \omega t) \quad (6.74)$$

This is an oscillating function, meaning that the energy density varies over space and time in phase with the propagation of the wave. All quantities are depicted in Fig. 6.8. Energy density is maximal where the potentials cross the zero axis. There is a phase shift of 90° between both plots that can be observed in the figure.

There is an analogy between longitudinal potential waves and acoustic waves. It is well known that acoustic waves in air or solids are mainly longitudinal too. The elongation of molecules is in the direction of wave propagation, as shown in Fig. 6.9. This is a variation in velocity. Therefore, the magnetic vector potential can be compared with a velocity field. The differences in elongation evoke a local pressure difference. Where the molecules are pressed together, the pressure is enhanced and vice versa. From conservation of momentum, the force $\vec{F}$ in a compressible fluid is given by:

$$\vec{F} = \dot{\vec{u}} + \frac{\nabla p}{\rho} \quad (6.75)$$

In this equation the term $\vec{u}$ is the velocity field, p is the pressure, and ρ is the density of the medium.

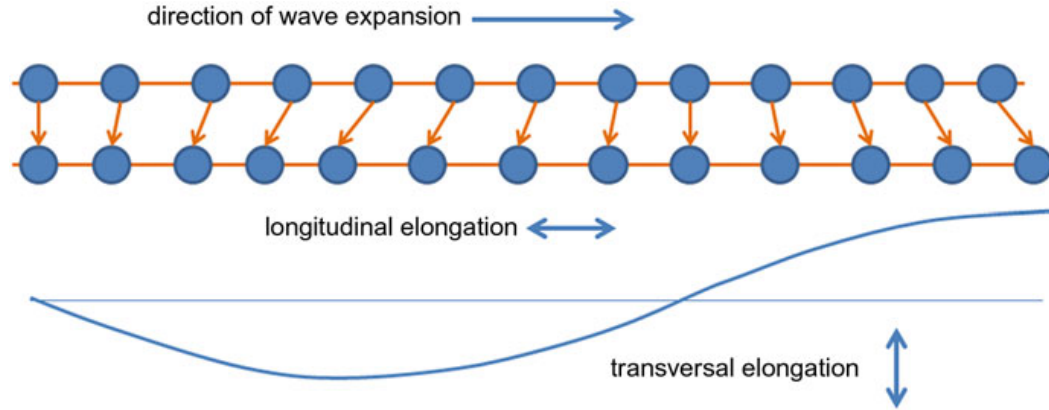


Fig. 6.9 Schematic representation of longitudinal and transversal waves

This is a full analysis of Eq. 6.63. In particular, we can see that in the EM case spacetime must be “compressible”; otherwise, there is no gradient of the scalar potential. As a consequence, space itself must be compressible, leading us to the principles of general relativity.

6.6 Transmitters and Receivers for Longitudinal Waves

A sender for longitudinal potential waves must be a device that avoids producing $\vec{E}$ and $\vec{B}$ fields but sends out oscillating potentials waves. We discuss two propositions about how this can be achieved technically. In the first case, we use two ordinary transmitter antennas (with directional characteristics) with a distance of half a wavelength (or an odd number of half waves). This means that ordinary EM waves cancel out, assuming that the near-field is not disturbed significantly. Because the radiated energy cannot disappear, it must propagate in space and is transmitted in the form of potential waves. This is depicted in Fig. 6.10.

A more common example is a bifilar flat coil (e.g., from the patent of Tesla [23])—see second drawing in Fig. 6.10. The currents in opposite directions effect an annihilation of the magnetic field component, while an electric part may remain because of the static field of the wires, as shown in the Fig. 6.11.

Construction of a receiver is not so straightforward. In principle no magnetic field can be retrieved directly from $\vec{A}$ because of Eq. 6.66. The only way is to obtain an electrical signal by separating both contributing parts in Eq. 6.63 so that the equality is outweighed and an effective electric field remains that can be detected by conventional devices [22]. A very simple method would be to place two plates of a capacitor at a distance of half a wavelength (or odd multiples of it). Then the voltage in space should influence the charge carriers in the plates, leading to the same effect as if a voltage had been applied between the plates. The real voltage in the plates or the compensating current can be measured (Fig. 6.12). The *tension of space* operates directly on the charge carriers while no electric field is induced. The $\vec{A}$ part

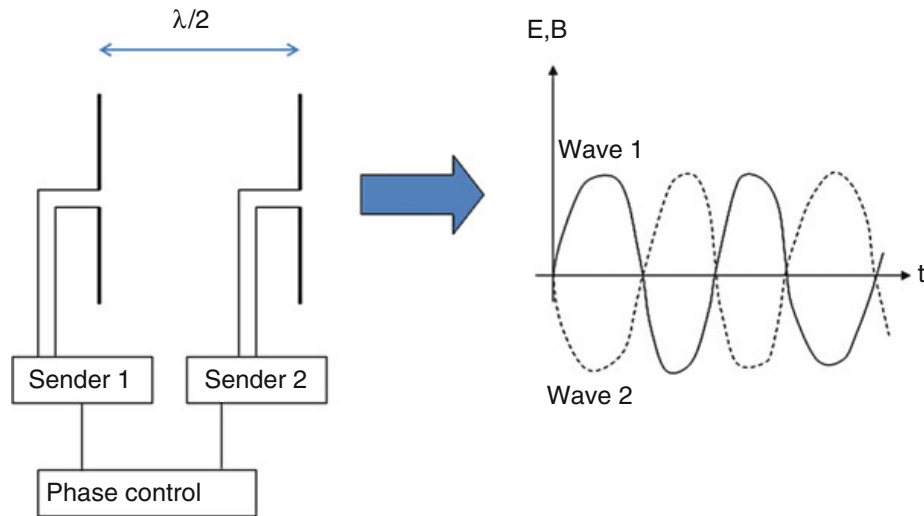
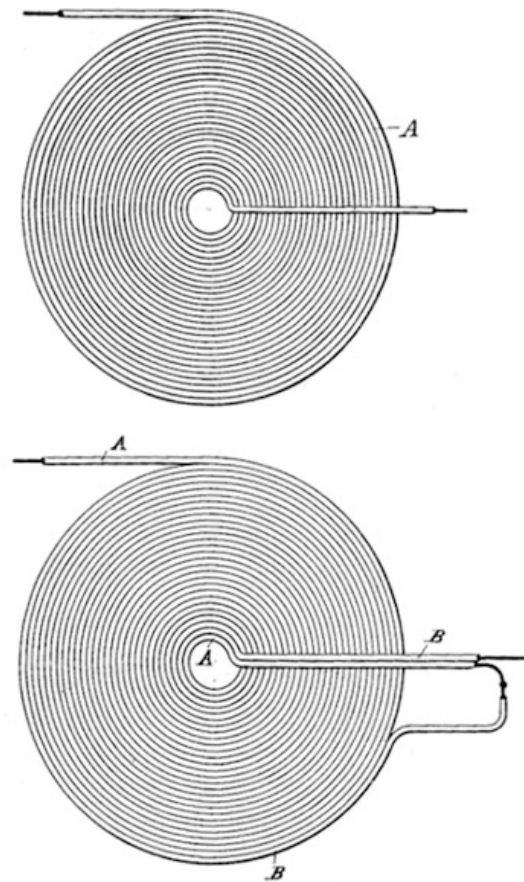


Fig. 6.10 Suggestion for a transmitter of longitudinal potential waves

Fig. 6.11 Tesla coils according to the patent [23]



is not contributing because the direction of the plates is perpendicular to it (i.e., no significant current can be induced).

Another possibility of a receiver is to use a screened box (*Faraday cage*). If the mechanism described for the capacitor plates is valid, the electrical voltage part of

Fig. 6.12 Suggestion for a receiver of longitudinal potential waves (capacitor)

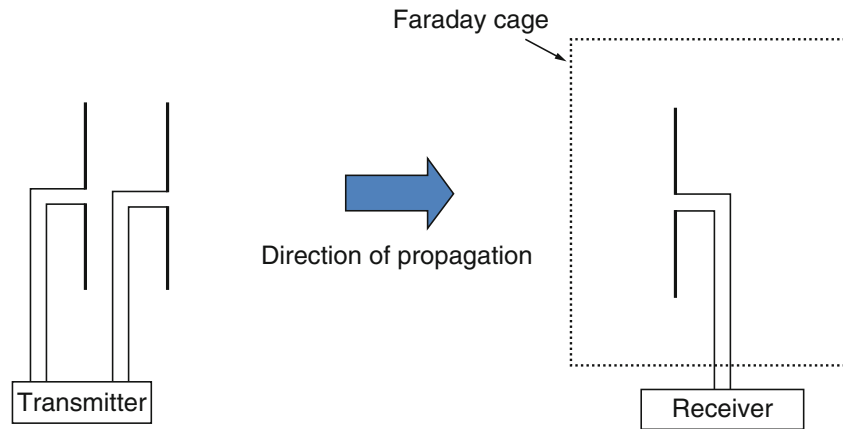
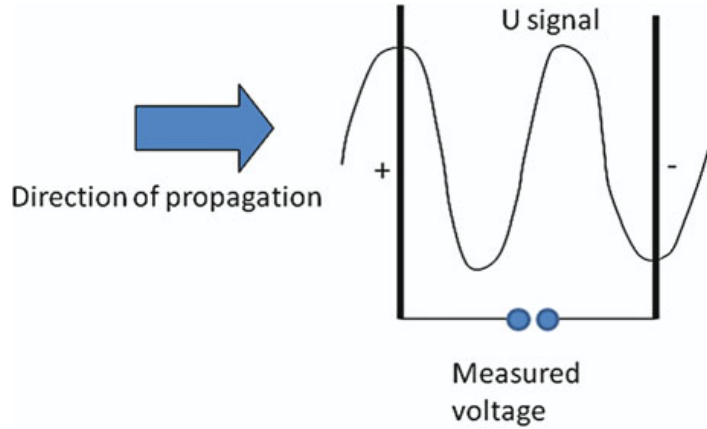


Fig. 6.13 Suggestion for a receiver of longitudinal potential waves (Faraday cage)

the wave creates charge effects that are compensated immediately because of the high conductivity of the material. As is well known, the interior of a Faraday cage is free of electric fields. The potential is constant because it is constant on the box's surface. Therefore, only the magnetic part of the wave propagates in the interior where it can be detected by a conventional receiver (Fig. 6.13).

Another method of detection is using vector potential effects in crystalline solids. As is well known from solid-state physics, the vector potential produces excitations within the quantum mechanical electronic structure, provided the frequency is near to the optical range. Crystal batteries work in this way. They can be engineered through chemical vapor deposition of carbon. In the process you get strong light-weight crystalline shapes that can handle lots of heat and stress by high currents. For detecting LWs, the excitation of the electronic system must be measured (e.g., by photoemission or other energetic processes in the crystal).

All these are suggestions for experiments with LWs. Additional experiments can be performed for testing the relationship between wave vector k and frequency ω to check whether this type of wave propagates with the ordinary velocity of light c :

$$c = \frac{\omega}{k} \quad (6.76)$$

where k is defined from the *wavelength* λ by the following relation:

$$k = \frac{2\pi}{\lambda} \quad (6.77)$$

As pointed out in the paper by Eckhardt and Lindstrom [22], the speed of propagation depends on the form of the waves.

This even can be a non-linear step-function. The experimental setup shown earlier in Fig. 6.11 can be used directly for finding the $\omega(\vec{k})$ relationship because the wavelength and frequency are measured at the same time. There are rumors that Eric P. Dollard [24] found a propagation speed of LWs of $(\pi/2) \cdot c$, which is 1.5 times the speed of light; however, no reliable experiments on this have been reported in the literature.

The ideas worked out in this section may not be the only way LWs can be explained and technically handled. As mentioned in the introduction, electrodynamics derived from a unified field theory (Evens et al. [19]) predicts effects of polarization in all space and time dimensions and may lead to a discovery of even richer and more interesting effects.

6.6.1 *Scalar Communication System*

The basic scalar communication system indicates that the communications antenna does not make any sense according to normal EM theory. The goal of a scalar antenna is to create powerful repulsion and/or attraction between two magnetic fields to create large scalar bubbles/voids. This is done by using an antenna with two opposing EM coils that effectively cancel out as much of each other's magnetic field as possible. An ideal scalar antenna will emit no EM field (or as little as possible) because all power is being focused into the repulsion–attraction between the two opposing magnetic fields. Normal EM theory suggests that because such a device emits no measurable EM field, it is useless and will only heat up.

A scalar signal reception antenna similarly excludes normal EM waves and only measures changes in magnetic field attraction and repulsion. This typically will be a two-coil powered antenna that sets up a static opposing or attracting magnetic fields between the coils, and the coils are counterwound so that any normal radio frequency (RF) signal will be picked up by both coils and effectively canceled out.

It has been suggested that scalar fields do not follow the same rules as EM waves and can penetrate through materials that would normally slow or absorb them. If true, a simple proving method is to design a scalar signal emitter and a scalar signal receiver and encase each inside separate shielded and grounded metal boxes, known as Faraday cages. These boxes will absorb all normal EM energy and will prevent any regular non-scalar signal transmissions from passing from one box to the other.

Some people have suggested that organic life may make use of scalar energies in ways that we do not yet understand. Therefore, caution is recommended when experimenting with this fringe technology. Nevertheless, keep in mind that if scalar fields do exist, we are likely already deeply immersed in an unseen field of scalar noise all the time, generated anywhere two magnetic fields oppose or attract. Common scalar field noise sources include AC electrical cords, powerlines carrying high currents, and electric motors that operate on the principle of powerful spinning regions of repulsion and attraction.

6.7 Scalar Waves Experiments

It can be shown that SWs normally remain unnoticed and are very interesting in practical use for information and energy technology for reasons of their special attributes. The mathematical and physical derivations are supported by practical experiments. Such demonstrations show the following:

1. Wireless transmission of electrical energy
2. Reactions of the receiver to the transmitter
3. Free energy with an overunity effect of about 3
4. Transmission of SWs with 1.5 times the speed of light
5. Inefficiency of a Faraday cage to shield SWs

6.7.1 Tesla Radiation

Shown here are five extraordinary science experiments that are incompatible with textbook physics. Short courses, that were given by Meyl [25] show the transmission of longitudinal electric waves.

This is a historical experiment because 100 years ago the famous experimental physicist Nikola Tesla measured the same wave properties as this author. From him stems a patent concerning the wireless transmission of energy (1900) [26]. Because he also had to find out that very much more energy arrives at the receiver than the transmitter takes up, he spoke of a *magnifying transmitter*.

By the effect back on the transmitter Tesla sees that he has found the resonance of the Earth that lies according to his measurement at 12 Hz. Because the Schumann resonance of a wave, which goes with the speed of light lies at 7.8 Hz, however, Tesla came to the conclusion that his wave was 1.5 times the speed of light c [27]. As founder of diathermy Tesla, already had pointed to the biological effectiveness and to the possible use in medicine. The diathermy of today has nothing to do with Tesla radiation; it uses the incorrect wave and, as a consequence, hardly has any medical importance.

The discovery of the Tesla radiation has been denied and is not mentioned in textbooks anymore. For that there are two reasons:

1. No high school ever has rebuilt a “magnifying transmitter,” The technology simply was too costly and too expensive. In that way the results have not been reproduced, as it is imperative to acknowledge. This author has solved this problem using modern electronics by replacing the spark a gap generator with a function generator and it operates with high-tension with 2–4 V low-tension. Meyl [25] sells the experiment as a demonstration-set so that it is reproduced as often as possible. It fits in a case and has been sold more than 100 times. Some universities already could confirm the effects. The measured degrees of effectiveness lie between 140 and 1000%.
2. The other reason why this important discovery could fall into oblivion can be seen in the missing suitable field description. Maxwell’s equations in any case only describe TWs, for which the field pointers oscillate perpendicular to the direction of propagation.

The vectorial part of the wave equation derived from Maxwell’s equations are presented here:

$$\left\{ \begin{array}{l} \vec{\nabla} \times \vec{E} = -\frac{\partial \vec{B}}{\partial t} \\ \vec{\nabla} \times \vec{H} = \vec{J} + \frac{\partial \vec{D}}{\partial t} \\ \vec{B} = \mu \vec{H} \\ \vec{D} = \epsilon \vec{E} \\ \vec{J} = 0 \end{array} \right. \Rightarrow \text{In Linear Media} \quad (6.78)$$

and

$$\vec{\nabla} \times (\vec{\nabla} \times \vec{E}) = -\mu \frac{\partial (\vec{\nabla} \times \vec{H})}{\partial t} = -\mu \epsilon \left(\frac{\partial^2 \vec{E}}{\partial t^2} \right) \quad (6.79)$$

Then, from the result of Eqs. 6.78 and 6.79, we obtain the wave equation:

$$\left\{ \begin{array}{l} \nabla^2 \vec{E} = \vec{\nabla} (\vec{\nabla} \cdot \vec{E}) - \vec{\nabla} \times (\vec{\nabla} \times \vec{E}) = \frac{1}{c^2} \frac{\partial^2 \vec{E}}{\partial t^2} \\ \mu \epsilon = \frac{1}{c^2} \end{array} \right. \quad (6.80)$$

See Chap. 4 of this book for more details on derivation of wave equations from Maxwell’s equations. Note that in all these calculations, the following symbols apply:

$\vec{E}$ = electric field or electric force

$\vec{H}$ = auxiliary field or magnetic field

$\vec{D}$ = electric displacement ($\vec{D} = \epsilon \vec{E}$ in linear medium)

$\vec{B}$ = magnetic intensity or magnetic induction

$\vec{J}$ = current density

Now breaking down the first equation of in the sets of Eq. 6.80 will be as follows:

$$\underbrace{\nabla^2 \vec{E}}_{\substack{\text{Laplace} \\ \text{operator over } \vec{E}}} = \underbrace{\vec{\nabla}(\vec{\nabla} \cdot \vec{E}) - \vec{\nabla} \times (\vec{\nabla} \times \vec{E})}_{\substack{\text{if } \vec{\nabla} \cdot \vec{E} = 0 \text{ then we have Transversal Wave} \\ \text{if } \vec{\nabla} \times \vec{E} = 0 \text{ then we have Longitudinal Wave}}} = \underbrace{\frac{1}{c^2} \frac{\partial^2 \vec{E}}{\partial t^2}}_{c \text{ is speed of light}} \quad (6.81)$$

Note that in the equation that if $\vec{\nabla} \cdot \vec{E} \neq 0$, then we have a situation that provides the SW conditions, while the following relationships apply as well:

$$\vec{E} = -\vec{\nabla} \phi : \begin{cases} (1) \cancel{\vec{\nabla}(\vec{\nabla} \cdot \vec{E})} = \cancel{\vec{\nabla}} \left[\frac{1}{C^2} \frac{\partial^2 \phi}{\partial t^2} \right] \\ (2) \quad \vec{\nabla} \cdot \vec{E} = -\vec{\nabla} \cdot \vec{\nabla} \phi \end{cases} \quad (6.82)$$

$$\vec{\nabla} \cdot \vec{D} = \rho : \begin{cases} (3) \quad \vec{\nabla} \cdot \vec{E} = \frac{\rho}{\epsilon} \end{cases}$$

From this equation we also can conclude that the *plasma wave* is:

$$\nabla^2 \phi = \frac{1}{c^2} \cdot \left(\frac{\partial^2 \phi}{\partial t^2} \right) - \frac{\rho}{\epsilon} \quad (6.83)$$

The results found in Eqs. 6.81 and Eq. 6.82 are the scalar part of the wave equation describing longitudinal electric waves, which end up with a deviation of plasma waves, as can be seen in Eq. 6.83. In these equations symbol ϕ represents a *scalar field*, as described in Chap. 4.

If we derive the field vector from a scalar potential ϕ , then this approach immediately leads to an inhomogeneous wave equation, which is called a *plasma wave*. Solutions are known, such as the electron plasma waves, that are *longitudinal oscillations* of the electron density—*Langmuir waves*.

6.7.2 Vortex Model

The Tesla experiment and this author's historical rebuild, however, show more. Such LWs obviously exist even without plasma in the air and even in vacuum. The questions thus are:

- I. What does divergence $\vec{E}$ describe in this case?
- II. How is the impulse passed on so that a longitudinal standing wave can form?
- III. How should a shock wave come about if there are no particles that can push each other?

We have solved these questions by extending Maxwell's field theory for vortices of the electric field. These so-called potential vortices are able to form structure and propagate in space for reason of their particle nature as a longitudinal shock wave. The model concept is based on the ring vortex model of Hermann von Helmholtz, which Lord Kelvin made popular. In Volume 3 of the Meyl book, *Potential Vortex* [1], the mathematical and physical derivation is described.

In spite of the field theoretical set of difficulties every physicist at first will seek a conventional explanation. We will try three approaches as follows: (1) resonant circuit interpretation, (2) Id interpretation, and (3) vortex interpretation. The details of these approaches are given in the following subsections.

6.7.2.1 Resonant Circuit Interpretation

Tesla presented his experiment to, among others, Lord Kelvin, and 100 years ago Tesla spoke about a vortex transmission. In the opinion of Kelvin, however, vortex transmission by no means concerns a wave but rather radiation. Kelvin recognized clearly that every radio-technical interpretation had to fail because alone the course of the field lines is a completely different one.

It presents itself assuming a resonant circuit, consisting of a capacitor and an inductance (Fig. 6.14). If both electrodes of the capacitor are pulled apart, then between both stretches an electric field. The field lines start at one sphere, the transmitter, and they bundle up again at the receiver. In this manner a higher degree of effectiveness and a very tight coupling can be expected. In this way, without doubt, some but not all, of the effects can be explained.

The inductance is split up in two air transformers, which are wound in a completely identical fashion. If a field in sinusoidal tension voltage is transformed up in the transmitter, then it again is transformed down at the receiver. The output voltage should be smaller or, at most, equal to the input voltage, but it is substantially higher!

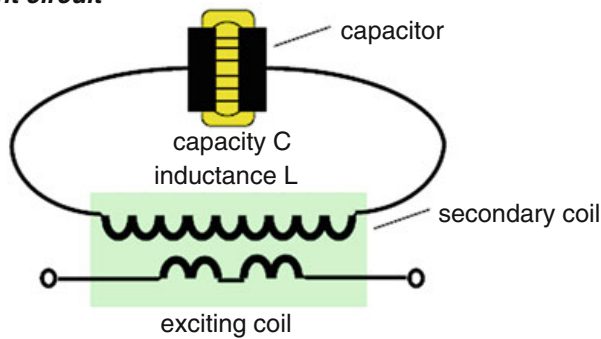
An alternative wiring diagram can be drawn and calculated, but in no case does the measurable result that light-emitting diodes at the receiver glow brightly ($U > 2 \text{ V}$), whereas at the same time the corresponding light-emitting diodes at the transmitter go out ($U < 2 \text{ V}$)! To check this result, both coils are exchanged.

The measured degree of effectiveness lies, despite the exchange, at 1000%. If the law of conservation of energy is not to be violated, then only one interpretation is left: The open capacitor withdraws field energy from its environment. Without consideration of this circumstance, the error deviation of every conventional model calculation lies at more than 90%. In this case, one should do without the calculation.

1. closed resonant circuit

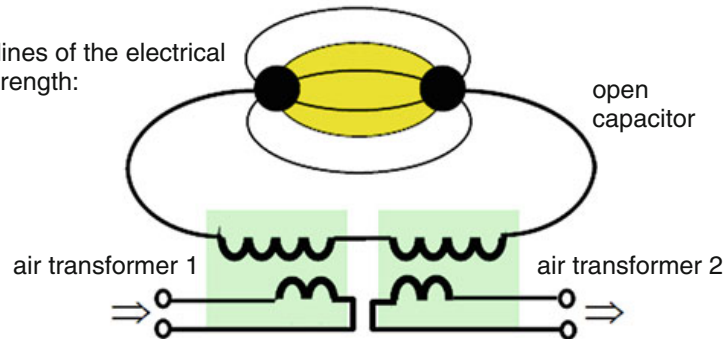
resonance
frequency:

$$f = \frac{1}{2\pi\sqrt{LC}}$$



2. separating the resonant circuit

Field lines of the electrical
field strength:



3. resonant circuit with open capacitor

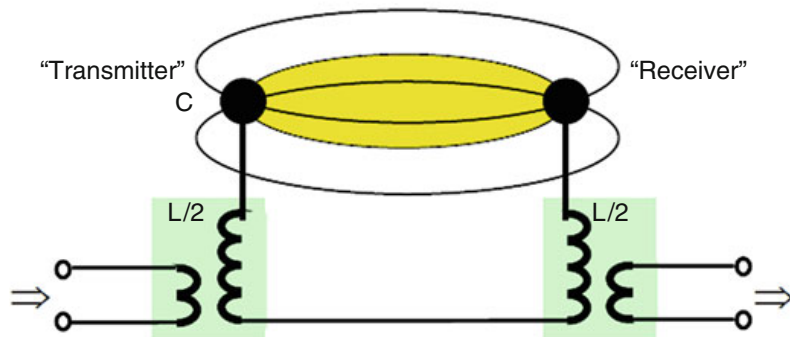


Fig. 6.14 Interpretation as an open resonant circuit

The calculation concerns oscillating fields because the spherical electrodes are changing in polarity with a frequency of approximately 7 MHz. They are operated in resonance. The condition for resonance reads as: identical frequency and opposite phase. The transmitter obviously modulates the field in its environment, while the receiver collects everything that fulfills the condition for resonance. Also, in the open question regarding the transmission velocity of the signal, the resonant circuit interpretation fails. But a HF-technician still has another explanation on the tip of his or her tongue.

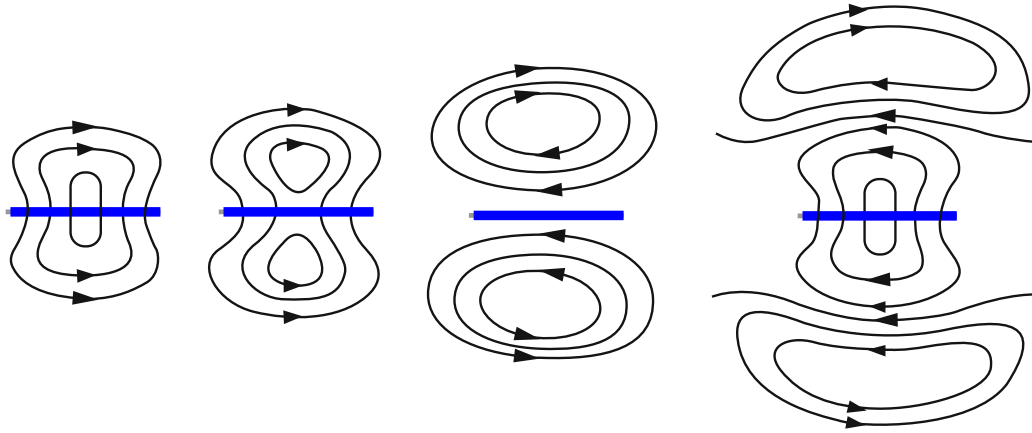


Fig. 6.15 The coming off of the electric field lines of the dipole

6.7.2.2 Near-Field Interpretation

At the antenna of a transmitter in the near-field (a fraction of the wavelength) only SWs (*potential vortex*) exist. They decompose into EM waves in the far-field and further. The near-field is not described by Maxwell's equations and the theory only is postulated. It is possible to pick up only SWs from radio transmissions. Receivers that pickup EM waves actually are converting those waves into potential vortices that are conceived as "standing waves."

This presents itself assuming a resonant circuit consisting of a capacitor and an inductance (Fig. 6.15). If both electrodes of the capacitor are pulled apart, then between both stretches an electric field. The field lines start at one sphere, the transmitter, and they bundle up again at the receiver. In this manner, a higher degree of effectiveness and a very tight coupling can be expected. In this way, without doubt, some but not all, of the effects can be explained.

In the near-field of an antenna effects are measured, on the one hand, as inexplicable because they evade the normally used field theory; on the other hand, can be shown as very close to SW effects. Everyone knows a practical application—for example, at the entrance of department stores, where the customer has to go through or between SW detectors.

In Meyl's experiment [25] the transmitter is situated in the mysterious near zone. Also, Tesla always worked in the near zone. But one who asks for the reasons will discover that the near-field effect is nothing else but the SW part of the wave equation. Meyl's explanation goes as follows: "The charge carriers which oscillate with high frequency in an antenna rod form longitudinal standing waves. As a result, also, the fields in the near zone of a *Hertzian dipole* are *longitudinal scalar wave fields*. The picture shows clearly how vortices are forming and how they come off the dipole."

Like for charge carriers in the antenna rod the phase angle between current and tension voltage amounts to 90° , and occurs in the near-field; likewise, electric and magnetic fields phase shifted for 90° . In the far-field, however, the phase angle is zero. In this author's interpretation the vortices are breaking up, they decay, and transverse radio waves are formed.

6.7.2.3 Vortex Interpretation

The vortex decay, however, depends on the velocity of propagation. Calculated at the speed of light the vortices already have decayed within half the wavelength. The faster the velocity, the more stable they get to remain stable above 1.6 times the velocity. These very fast vortices contract in the dimensions. They now can tunnel. Therefore, speed faster than light occurs at the tunnel effect. Therefore, no *Faraday cage* can shield fast vortices.

Because these field vortices with a particle nature following the high-frequency oscillation permanently change their polarity from positive to negative and back, they do not have a charge on the average over time. As a result, they almost are unhindered penetrate solids. Particles with this property are called *neutrinos* in physics. The field energy that is collected in this experiment stems from the neutrino radiation that surrounds us. Because the source of this radiation, all the same if the origin is artificial or natural, is far away from the receiver, every attempt of a near-field interpretation goes wrong. After all, does the transmitter installed in the near-field zone supply less than 10% of the received power. The 90%, however, which is of concern here, cannot stem from the near-field zone!

6.7.3 Meyl's Experiment

In Meyl's experimental set up he also takes a few other steps in order to conduct his research that is reported here [25]:

At the function generator he adjusts frequency and amplitude of the sinusoidal signal, with which the transmitter is operated. At the frequency regulator I turn so long, until the light-emitting diodes at the receiver glow brightly, whereas those at the transmitter go out. Now an energy transmission takes place.

If the amplitude is reduced so far, until it is guaranteed that no surplus energy is radiated, then in addition a gain of energy takes place by energy amplification.

If he takes down the receiver by pulling out the earthing, then the lighting up of the LED's signals, the mentioned effect back on the transmitter. The transmitter thus feels, if its signal is received.

The self-resonance of the Tesla coils, according to the frequency counter, lies at 7 MHz. Now the frequency is running down and see there, at approx. 4.7 MHz the receiver again glows, but less bright, easily shieldable and without discernible effect back on the transmitter. Now we unambiguously are dealing with the transmission of the Hertzian part and that goes with the speed of light. Because the wavelength was not changed, does the proportion of the frequencies determine the proportion of the velocities of propagation? The SW according to that goes with $(7/4.7 =)$ 1.5 times the speed of light.

If Meyl puts the transmitter into the aluminum case and closes the door, then nothing should arrive at the receiver. Expert laboratories for EM compatibility in this

case indeed cannot detect anything and, although in spite of that, the receiver lamps glow! By turning the receiver coil, it can be verified that an electric and not a magnetic coupling is present although the Faraday cage should shield electric fields. The SW obviously overcomes the cage with a speed faster than light by *tunneling*. We can summarize what we have discussed so far in respect to the SW in next subsection as follows.

6.7.4 Summary

German professor Konstantin Meyl developed a new unified field and particle theory based on the work of Tesla. Meyl's theory describes quantum and classical physics, mass, gravitation, the constant speed of light, neutrinos, waves, and particles—all explained by vortices. The subatomic particle characteristics are calculated accurately by this model. Well-known equations also are derived by the unified equation. He provides tools replicating one of Tesla's experiments, which demonstrates the existence of SWs. Scalar waves are simply energy vortices in the form of particles. Here is a summary of an interview with Konstantin Meyl on his theory and technologies.

The *unified field theory* describes the electromagnetic, eddy current, potential vortex, and special distributions. This combines an extended wave equation with a Poisson equation. Maxwell's equations can be derived as a special case where Gauss's Law for magnetism is not equal to zero. This means that magnetic charges do exist in Meyl's theory [25]. That electric and magnetic fields always are generated by motion is the fundamental idea that this equation is derived from. The unipolar generator and transformer have conflicting theories under standard theories. Meyl splits them into the equations of transformation of the electric and magnetic fields separately, which describes unipolar induction and the equation of convection, relatively.

Meyl says that the field is always first, which generates particles by decay or conversion. Classical physics does not recognize energy particles (i.e., potential vortices), so they were not included in the theory. Quantum physics effectively tried to explain everything with vortices, which is why it is incomplete. The derivation of Schrodinger's equation from the extended Maxwell's equations means they are vortices. For example, photons are light as particle vortices and EM light is in wave form, which depends on the detection method that can change the form of light.

Gravitation is from the speed of light difference caused by proximity that, proportional to field strength, decreases the distance of everything for the field strength. This causes the spin of the Earth or another mass to move quicker farther away from the greatest other field's influence, thus orbit the Sun or larger mass. The closest parts of the bodies have smaller distances because of larger total fields and thus slower speeds of light. These fields are generated by closed field lines of vortices and largely are matter. Matter does not move as energy because the speed

of light is zero in the field of the vortex because of infinite field strength within the closed field. The more mass in proximity to something, the greater the field strength and the shorter the distances, which causes larger groups of subatomic particles to individually have smaller sizes.

The total field energy in the Universe is exactly zero, but particle and energy forms of vortices divide the energy inside and outside the vortex boundary. When particles are destroyed, no energy is released. No energy was produced when large amounts of matter was destroyed at MIT with accelerated sodium atoms. This is what Tesla predicted but contradicts Einstein's $E = MC^2$. Einstein's equation is correct as long as the number of subatomic particles is only divided; energy comes from mass defect, not from destruction.

There are various kinds of waves. Electromagnetic waves are fields, scalar electric or eddy currents or a magnetic vortex, which Tesla started with, and magnetic scalar or the potential vortex, which Meyl focuses on and is used in nature. The EM is fixed at the speed of light at that specific closed field strength. Scalar vortices can be any speed. Neutrinos travel at $1.6c$ or higher and do not decay to EM. Tesla-type SWs are between c and $1.6c$ and decay at distances proportional to their speed (used in the traditional radio near-field). Under the speed c , the scalar vortex acts as an electron.

Black holes may produce and emit neutrinos by condensing and transforming matter into massive fast particles with apparently no mass or charge because of their very high frequency of fluctuation. Neutrinos oscillate in mass and charge. When neutrinos hit matter and have a precise charge or mass, they produce one of three effects: a gain in mass, a production of EM, or emission of slower neutrinos.

Resonance requires the same frequency, same modulation, and opposite phase angle. Once (scalar) resonance is reached, a direct connection is created from the transmitter to the receiver. Signal and power will pass through a *Faraday cage*.

References

1. http://www.k-meyl.de/xt_shop/index.php?cat=c3_Books-in-English.html
2. B. Zohuri, *Plasma Physics and Controlled Thermonuclear Reactions Driven Fusion Energy* (Springer, Cham, 2016)
3. B. Zohuri, *Directed Energy Weapons: Physics of High Energy Lasers (HEL)*, 1st edn. (Springer Publishing Company, Cham, 2016)
4. L.M. Hively, G.C. Giakos, Toward a more complete electrodynamic theory. *Int. J. Signals Imaging Syst. Engr.* **5**, 3–10 (2012)
5. D.A. Woodside, Three vector and scalar field identities and uniqueness theorems in Euclidian and Minkowski space. *Am. J. Phys.* **77** (May 2009)
6. M.W. Evans, J.P. Vigiér, *The Enigmatic Photon, Volume 1: The Field $\vec{B}^{(3)}$* (Springer, Dordrecht, 1994).; Softcover Reprint of the Originated. Edition
7. M.W. Evans, *Phys. B* **182**, 237 (1992).; **183**, 103 (1993)
8. M.W. Evans, in *Waves and Particles in Light and Matter*, ed. by A. Garuccio, A. Van der Merwe (Eds), (Plenum, New York, 1994)
9. M.W. Evans, *The Photon's Magnetic Field* (World Scientific, Singapore, 1992)