

reprinted on March 25, 1999
ICO paper
article by Ad Lagendijk
in book: Current Trends in Optics
edited by J.C. Dainty
(Academic, London, 1994)

CHAPTER 4

Can Light Be Localised?

Ad Lagendijk
FOM-Institute for Atomic and Molecular Physics
Kruislaan 407, 1098 SJ Amsterdam, The Netherlands.
and
Van der Waals-Zeeman Laboratorium, University of Amsterdam
Valckenierstraat 65-67, 1018 XE Amsterdam, The Netherlands.

1 Introduction

The problem with papers that have a question mark in their title is that the question posed in the title is usually not answered in the paper. This paper will not be an exception.

Sometimes scientists working in a new fashionable field of science monopolise the use of well-proven and established words that used to represent important and general concepts. Suddenly these terms denote a much more specific and limited concept. *Chaos* is a good example of such a word, *localisation* another. Localisation of an object, in its sober meaning, signifies that the item is confined to a certain region of space. The ultimate degree of localisation would be confinement to a point. Let us stretch our ideas about localisability a little, and consider a *function* to be localised in a certain region of space if its magnitude is (at least) exponentially decaying in any direction away from this neighbourhood. This exponential localisation is quite simple to obtain when solving the wave equation so dear to the heart of each and every condensed matter physicist: the Schrödinger wave equation. The question whether or not localisation exists in quantum mechanics can be rephrased as: “Does the Schrödinger equation support bound states?” The answer is affirmative. Any book on quantum mechanics is puffed up with them. Consequently quantum objects like electrons and atoms can be localised in the sober sense of localisation. For *classical* (that is non-quantummechanical) wave equations (like Maxwell’s equations, scalar wave equations, etc.) this question is much more subtle. In dielectric systems no bound states are supported by Maxwell’s equations in three dimensions. The guided modes in a glass fibre may look like bound states, but they

are only bound in two of the three dimensions. We need the existence of a negative dielectric constant in part of space to obtain localised solutions, but this practically means that we have to deal with rather lossy (metallic) systems. Absorption is a real killer of the phenomena I will discuss here, and for that reason I will not consider dissipative systems in any detail.

What is the “symbolic” or “frivolous” meaning of localisation in condensed matter physics? In quantum mechanics of the solid state this invariably refers to *localisation induced by disorder*. Consequently this particular form of localisation is a property of some random system. If there is something workers in optics hate and condensed matter physicists love, it is *randomness*. Condensed matter physicists have given up the hope of ever preparing the perfect crystal. So why not turn a disaster into a success story and make the study of disorder the story of your life. Believe it or not it is in these “dirty” materials that some absolutely beautiful properties, sometimes promoted to the status of “paradigms” have emerged. “Localisation of electrons in a disordered system”[1] is one of them. Many of these novel phenomena in random media have been formulated in the language of condensed-matter physics, but the basic concepts can easily be rephrased using optical terminology: the effects are all connected to the importance of *interference* in multiple scattering of waves. This statement immediately brings up the question whether these findings might also apply to other wave equations.

In our group we realized in an early stage the analogy of the role of disorder between the Schrödinger wave equation and the classical wave equations. Our experiments are always performed on strongly scattering, milky white, materials: the nightmare of every optical physicist interested in imaging. One should realize that, maybe apart from the study of surface roughness, exploration of the influence of disorder was a very unpopular activity in optics. The interest in surface scattering stems from the fact that surface roughness of optical components puts limitations on the overall quality of an optical system.

The definition of localisation of electrons or light is not simple.[2] One expects that when a wave propagates through a random medium standard radiative transport and diffusion theory is appropriate. However, depending on a number of parameters including dimensionality, we now know that under strong-scattering conditions this picture may break down. Due to interference in multiple scattering the mean free path gets reduced up to a point where it vanishes and no propagation exists any more. The wave is captured in what we could call a “random cavity”. For electrons this is a well-studied phenomenon with enormous consequences for the electrical conductivity. Notwithstanding the tremendous efforts in that field lasting for at least thirty years now, there are still a large number of principal questions about the nature of localised electrons. What about the likelihood of inducing localisation of light?[3] This field is still young (I estimate its age seven years, at present) and my answer to the question in the title is a clear: “I really do not know.” (I *do* care by the way). Happily the chase for this ultimate goal of trapping light in a (three-dimensional) random cavity keeps on bringing us incredible surprises, partly because the study of light propagation in this multiple-scattering regime has never been so exhaustive.

Localisation of light is intimately connected to localisation of electrons. To make the mapping between the two cases as transparent as possible, I will introduce, in addition to using Lorentz-Heavyside electromagnetic units, the following simplifying notation: the speed of light, Planck’s constant and the electron mass are normalized to one, $c \equiv \hbar \equiv m \equiv 1$. This has a surprising positive effect on the cosmetic appearance of the equations presented in this chapter and makes all energy and (angular) frequency scales having the dimension of inverse length in the classical wave equations and of

inverse length squared in the Schrödinger equation. The frequency component of light will be referred to as E . So the time dependence of a monochromatic classical wave is $e^{\pm iEt}$.

I will review in this article some simple concepts that are relevant for the discussion of localisation of light. Quantities that will be discussed are speeds of propagation, mean free path and the diffusion constant. Another crucial quantity is the density of states. Finally I will point out some important problems encountered in the theory of light localisation.

2 Classical waves versus electrons

Let us look only superficially at a classical wave equation and a quantum-mechanical equation describing scattering. In the case of Schrödinger's equation we suppose that the inhomogeneity is represented by a spatially varying potential $V(\mathbf{r})$ (in units of $2m/\hbar^2$). In scalar wave equation we have a spatially varying dielectric constant $\varepsilon(\mathbf{r})$. After substituting monochromatic solutions in these equations we wind up with

$$(\nabla + E)\Psi = V(\mathbf{r})\Psi \quad \leftarrow \text{Schrödinger equation,} \quad (1)$$

$$(\nabla + E^2)\Psi = [1 - \varepsilon(\mathbf{r})]E^2\Psi \quad \leftarrow \text{scalar wave equation.} \quad (2)$$

The left hand sides of these equations refer to the homogeneous system and the right hand sides constitute the source of the scattering. When comparing both these right hand sides we observe that contrary to scattering of Schrödinger waves from potentials, scattering of scalar waves from dielectric particles is characterized by an “energy-dependent” potential. A comparison of the time-independent Schrödinger equation with potential $V(\mathbf{r})$ and energy E to the scalar-wave equation with the spatially varying dielectric constant $\varepsilon(\mathbf{r})$ at frequency E , suggests to identify a “potential” $V(\mathbf{r}, E) = [1 - \varepsilon(\mathbf{r})]g(E)$. The energy (or frequency) dependence $g(E) = E^2$ is due to the second-order time derivative in the scalar wave equation. The notion of an energy-dependent potential is most apparent in the low-energy limit where the energy dependence E^2 is responsible for the characteristic E^4 Rayleigh scattering, whereas for the Schrödinger equation the low-energy limit is energy independent (s-wave limit). The presence of an energy-dependent potential also signifies different conservation laws as compared to cases where the potential is energy independent.

3 Single scattering

Paying tribute to the reductionistic character of physics every physicist is looking for building blocks. Whenever possible he would like to use building blocks that are as simple as possible. In multiple scattering a possible candidate for a simple building block is evidently *single scattering*. The (single) scattering off an object with general shape is already very complicated to characterise. In line with our desire for extreme simplicity it seems a good idea to limit ourselves to scatterers having a spherical shape. One can immediately envisage three regimes of scattering off these spheres: the radius of the sphere r_0 being larger, smaller, or on the order of the wavelength of the wave being scattered. In optics the latter regime is usually referred to as the Mie regime. In this region the scattering exhibits many shape resonances. Inspection of a book on optics in which the Mie solution is presented explicitly immediately demonstrates that even the Mie solution is not “simple” at all. To prove this point I present in Fig. (1)

the cross-section of a Mie scatterer, with radius r_0 , as a function of the dimensionless size parameter $x = Er_0$. The figure is quite busy with resonances.

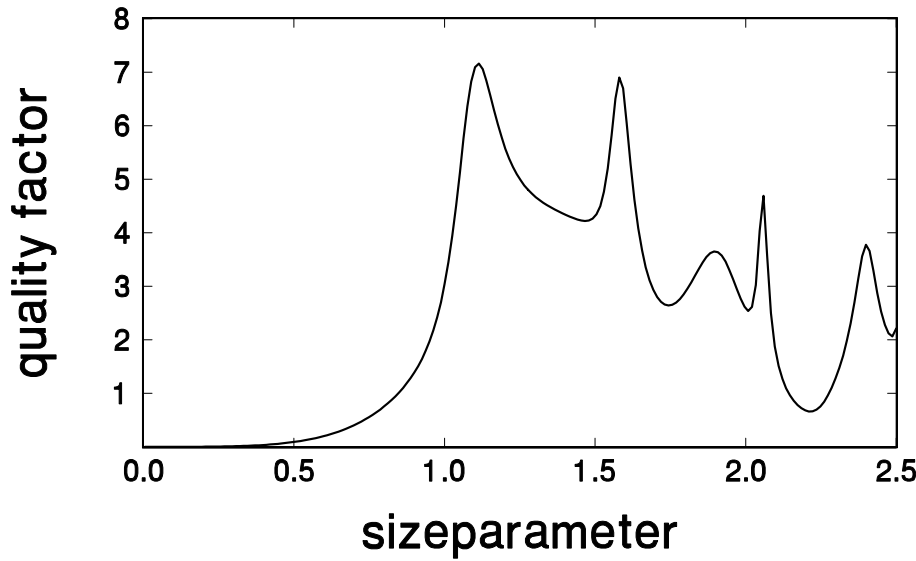


Figure 1: *Dimensionless quality factor of a dielectric sphere, defined as the total cross-section divided by the geometric cross-section, $Q \equiv \sigma/(\pi r_0^2)$, as a function of the scaled frequency $Er_0(\equiv \text{size parameter } x)$. Index of refraction of the sphere $m = 2.7$.*

When we put many of these Mie scatterers together we anticipate that not all the fine details of the properties of one single scatterer will matter in a situation where multiple scattering is dominant. Why not replace a collection of finite-sized spherical scatterers by a collection of even simpler scatterers: point scatterers? This is a popular procedure in condensed matter physics, where one is dealing with the Schrödinger equation. One knows that introducing in this equation point potentials represented by delta functions give problems (divergencies!). In optics these divergencies unfortunately also show up when using point scatterers. I am not going to bother you with details how to solve this mathematical tedious problem. It is possible to introduce regular point scatterers for light but they *must* have an *internal energy* resonance. A classical electron bound by a harmonic force and a two-level system are both good examples of systems for which the interaction with light can be well represented by point scatterers having an internal energy resonance. Scatterers with an internal energy resonance are often used as a model for atoms in atomic physics[4] but they can also be used in multiple scattering to mimic the shape resonances of finite-size scatterer. We recently showed that these point-scattering models can be obtained straightforwardly from the classical wave equations.[5] The scattering of such a point scatterer is characterized by three parameters, a “radius” r_0 , the position of the energy resonance E_0 , and the width of the resonance $1/T_2$. In Fig. (2) the cross-section of such a point scatterer is plotted as a function of the frequency of the incoming wave. Note the extreme simplification with respect to Fig (1).

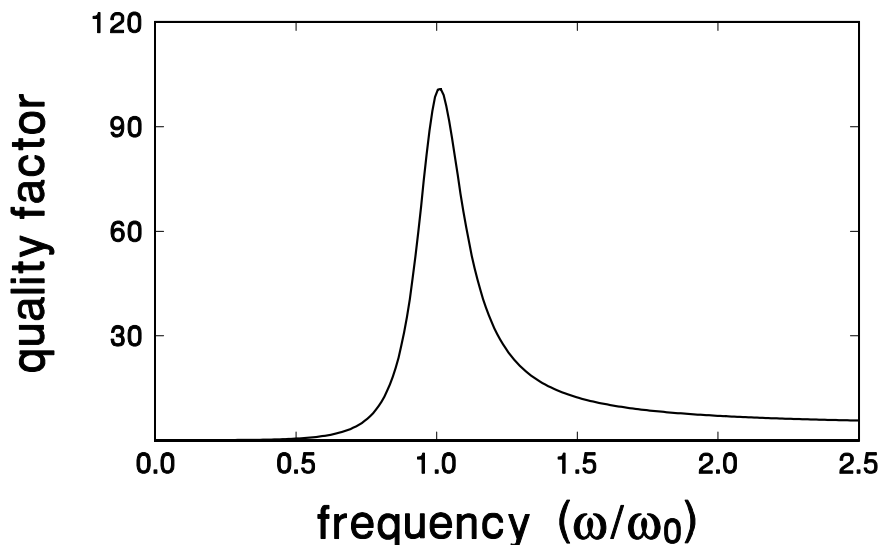


Figure 2: *Dimensionless quality factor of a classical oscillator (point scatterer) having an internal resonance at E_0 as a function of the scaled frequency E/E_0 . The damping $1/(T_2 E_0) = 0.1$. The quality factor is defined as the total cross-section divided by the “geometric cross-section”, $Q \equiv \sigma/(\pi r_0^2)$. Note the extreme simplification with respect to Fig. (1)*

4 Multiple scattering and diffusion

The most important concepts in physics is symmetry and corresponding conservation laws. If a quantity is conserved globally, the removal of a local imbalance will be slow (dynamic restoration with an inverse power law in time) as transport over all of space is needed to get rid of the asymmetry. If there is no conservation law operating we have the uninteresting situation that a local disturbance is removed quickly (that is exponential in time). I will give an illustrative example. If you would find a collection of a few thousand soap bubbles in your house and you would like to restore the homogeneous situation (that is clean your house) this can be done very quickly. If you would stumble on a few thousand billiard balls in one of your rooms the spreading of these balls over space will be painfully time consuming: the number of billiard balls is conserved in this example and the number of soap bubbles not. Whenever a quantity is conserved the transport of its density is slow. This is the situation where over long distances the transport can be described with a random walk and a diffusion law. Sometimes the conservation law is not obeyed rigorously and a small amount of local absorption is possible. If the absorption is small it can be dealt with as a perturbation. If the infringement of the conservation law is very serious the whole concept of a random walk and diffusion breaks down.

During the propagation of classical waves through (elastically) scattering media the total energy is conserved and consequently transport of energy density will be slow and can be described in terms of diffusion. A heuristic deduction of the diffusion equation from multiple scattering is simple (A strict derivation is demanding). The basic ingredient in any derivation of the diffusion equation is the *neglect* of all interference effects between multiple-scattered waves. When all interference is neglected one can

invoke random walk arguments to obtain a transport equation. The motion is characterised by a (transport) mean free path ℓ . This length should depend only on the cross-section σ of the individual scattering objects and on their density n . Each of these properties should have an inverse effect: on increasing either of them the mean free path should become shorter. This brings us quite smoothly to $\ell = 1/n\sigma$. To translate the step length into time we need some characteristic velocity of propagation v (which we shall discuss in some detail later).

The diffusion limit corresponds to a situation in transport when both q and ω are small (the “Kubo limit”). The wavevector q refers to the inverse of a typical length scale of the imbalance in the quantity that has to be transported and ω refers to a typical frequency scale of the perturbation. How small these quantities have to be for the diffusion approximation to become valid, can easily be estimated. The relevant length scale is the (transport) mean free path ℓ and the relevant time scale is the mean free time $\tau = \ell/v$. So the diffusion approximation applies when

$$\begin{cases} q\ell \ll 1, \\ \omega\tau \ll 1. \end{cases} \quad (3)$$

I want to stress the assertion that the term “diffusion approximation” in radiative transport is used for two different regimes: the static and the dynamic regime. Most radiative transport theory is concerned with steady-state transport. There is no time dependence then, and this situation applies when ω has been taken to be zero. One might immediately ask where the boundary between steady-state and transient behaviour is located. Any experiment, and certainly any steady-state experiment, is performed by definition on a finite-sized sample. Using a simple outcome of the diffusion equation one calculates the time for transport through the sample to be of order L^2/D , where L is the linear size of the sample (being much larger than the mean free path) and the steady-state diffusion limit applies when $\omega \ll D/L^2$. Apparently we can discriminate between the two diffusion regimes on a quantitative basis:

$$\text{transient diffusion : } \begin{cases} q\ell \ll 1, \\ (\ell/L)^2 \ll \omega\tau \ll 1, \end{cases} \quad (4)$$

$$\text{steady-state diffusion : } \begin{cases} q\ell \ll 1, \\ \omega\tau \ll (\ell/L)^2. \end{cases} \quad (5)$$

In the regime of transient diffusion the diffusion constant D is measured. In the regime of steady-state diffusion the transport mean free path ℓ is the only physical parameter that can be obtained experimentally. Why do I put so much emphasis on the difference between the mean free path and the diffusion coefficient? In a simple Boltzmann picture the diffusion coefficient is given by

$$D = \frac{1}{3}v\ell. \quad (6)$$

This relation between mean free path and diffusion coefficient is very general and can be found, besides the field of radiative transport theory, in many other areas of statistical physics where diffusion is discussed. The velocity v that comes to mind that could be used in formula (6) is for instance the phase velocity v_p or, as time-dependent transport is involved, the group velocity v_g . However, the velocity featuring in the diffusion coefficient is neither of them.[6, 7, 8, 9] The velocity that should be put into equation (6) is the *transport velocity* v_E . This transport or energy velocity is a highly

non-trivial quantity. Only after we became aware of the rich structure of the energy velocity, we realised the importance of the difference between transient and steady-state diffusion. The existence of this difference makes clear that from an experimental point of view D and ℓ are quantities that can be determined *independently*. This makes the velocity v_E an observable quantity.

Localisation of light is connected with the break-down of *steady-state* diffusion in a situation where according to inequalities (4) and (5) steady-state diffusion theory should apply. For infinite systems the situation is more subtle as no steady-limit exists and the diffusion constant is the only relevant parameter.

5 Reduction of velocity

In this section I will describe an important effect (that had been overlooked until recently) that gives rise to a reduction of the *diffusion constant* that is unrelated to localisation effects. Localisation of light is associated with the reduction of the (*transport*) *mean free path* of light due to interference in multiple scattering. Given the simple relation between mean free path and diffusion coefficient most scientists do not bother even to make the distinction: localisation is the vanishing of either the diffusion constant or the mean free path. However, the proportionality factor between D and ℓ , the transport velocity v_E , is a highly significant quantity. We have shown that transport velocity can be smaller than the phase velocity due to the long dwelling of classical waves inside resonant scatterers.[6, 7, 8, 9]

The transport velocity v_E can be defined as the velocity appearing in the diffusion constant. If one would incorrectly assume that all theoretical results obtained with the Schrödinger equation can immediately be used for scalar waves, one would find that $v_E = 1/v_p$ (remember that $c = 1$). In contrast, up to now every worker used $v_E = v_p$, without any justification. This non-rationalised habit was apparently based on intuition. Our theory, which correctly takes into account the difference between the two cases (See again Eqs. (1) and (2)) predicts the transport velocity v_E to be (approximately) equal to the phase velocity v_p *only* for frequencies away from resonances of the scatterers. Near the resonances both the transport velocity and the Boltzmann diffusion constant can become dramatically reduced, the amount of reduction being proportional to the number density of the scatterers.

After our original paper [6] had appeared quite some workers were very sceptical about our results, despite the fact that a reduction of the transport velocity was already well-known in atomic physics.[4]. These workers continued to believe that results obtained for the Schrödinger equation could immediately be used for classical waves. Our results contradicted this general conviction. The difference between the two cases is already evident when comparing the wave equations (1) and (2). Apart from a lot of critical discussions during conferences, critical correspondence, and extremely critical referee reports on our papers, actually only one paper was published that explicitly addressed the same question and found, contrary to us, that always $v_E = v_p$. [10] We have found the error in their reasoning and have pointed it out.[9] These authors now agree that their original criticism, as formulated in Ref. [10] was unjustified. The whole controversy rests on the fact that workers are not used to treat multiple scattering in which an energy-dependent potential is concerned. The good thing about the criticism was that it forced us to rederive our results in several independent ways. By being pushed to do this we discovered some very interesting relations between dwell times[11], Wigner phase delay times, and absorption coefficients. We have now five, essentially different but consistent, expressions for the energy velocity.[6, 7, 8, 9, 12] I am convinced now

that the community has accepted our new results as being correct.

Let us give the simplest, but basically rigorous derivation (essentially reproduced from Ref. [8] and inspired by Loudon [4]) of v_E by looking *directly* at the conserved quantity. This can be most conveniently done for scalar waves. A velocity is by definition a flux divided by a density. For scalar classical waves the energy density before averaging is $W = \frac{1}{2}\varepsilon(\mathbf{r})|\partial_t\psi|^2 + \frac{1}{2}|\nabla\psi|^2$, with corresponding ‘‘Poynting vector’’ $\mathbf{S} = -\text{Re}(\partial_t\psi)(\nabla\psi)^*$. After averaging over realizations and cycles, both terms in W represent an equal amount of energy. For monochromatic modes $\psi \sim \psi_E \exp(-iEt)$ and using the slowly-varying amplitude approximation (with introduction of wave vector $\mathbf{k}_E$) we arrive at

$$\langle W_E \rangle = E^2 \left\{ \langle |\psi_E|^2 \rangle + \langle [\varepsilon(\mathbf{r}) - 1] |\psi_E|^2 \rangle \right\}, \quad (7)$$

$$\langle \mathbf{S}_E \rangle = E \mathbf{k}_E \langle |\psi_E|^2 \rangle = \hat{\mathbf{k}} \frac{E^2}{v_p} \langle |\psi_E|^2 \rangle. \quad (8)$$

The brackets denote ensemble averaging in the thermodynamic limit, but can (on the basis of ergodicity) be replaced by space averaging. Equations (7) and (8) look so similar to well-known equations that can be found in many textbooks on classical electrodynamics that their form could already have been guessed without going through the ‘‘derivation’’ based on the rigorous formula’s for the energy density and current density given above. The second term in $\langle W_E \rangle$ is non-zero only inside the scatterers and is therefore at least first order in density. Since it is *not* present for Schrödinger scattering (where the particle number density $W_E \propto |\psi_E|^2$), this term gives concrete form to the different conserved quantity mentioned earlier. The current $\mathbf{S}_E$ is essentially identical to the current in the Schrödinger equation and does not give rise to a difference between the two cases. We used $k_E \approx E/v_p$, where v_p is the phase velocity. By normalizing ψ_E we equate the first term in Eq. (7) within braces to unity, and the second to $\Delta(E)$. Hence $\langle W_E \rangle = E^2 [1 + \Delta(E)]$. The quantity $E^2\Delta(E)$ represents the ‘‘potential energy’’ inside the scatterers. The transport velocity v_E becomes the ratio of the flux over the density,

$$v_E \equiv \frac{\langle S_E \rangle}{\langle W_E \rangle} = \frac{1}{v_p} \frac{1}{1 + \Delta(E)}. \quad (9)$$

The second equality makes use of Eqs. (7) and (8). The quantity $\Delta(E)$ can be calculated into lowest order of the number density n of the scatterers by inserting for ψ_E the normalized eigenfunction ψ_E^+ of one single scatterer S . We get

$$\Delta(E) = n \int_S d\mathbf{r} [\varepsilon(\mathbf{r}) - 1] |\psi_E^+(\mathbf{r})|^2 + \mathcal{O}(n^2). \quad (10)$$

It is known from literature [11, 12] that the integral of the square of the eigenfunction over the *scatterer region*, which appear in Eq. (10) is one of the possible definitions of the dwell time. The dwell time was originally developed for one-dimensional tunnelling problems only, but we have extended it into three dimensions.[12] The dwell time can become quite large in a resonant scatterer. We have calculated the energy velocity for a variety of particles, among which a collection of point scatterers with an internal energy resonance, and an assembly of finite-size scatterers in scalar and vector wave equations. In Fig. (3) I have plotted typical results for a collection of point scatterers, and in Fig (4) I have plotted the transport velocity of finite-size Mie scatterers for parameters relevant

for many of our experiments. We see that in all cases, either internal energy resonances or external shape resonances, the velocity can be dramatically smaller than the vacuum speed or medium speed of light. The retardation effect I have discussed so far is already present on the level of the Boltzmann approximation and is not an interference effect and therefore *not* connected to localisation. In these figures the magnitudes of phase and group velocities have also been displayed and they will be discussed later. The reader might already get intrigued by the large differences between either of these velocities and the energy velocity.

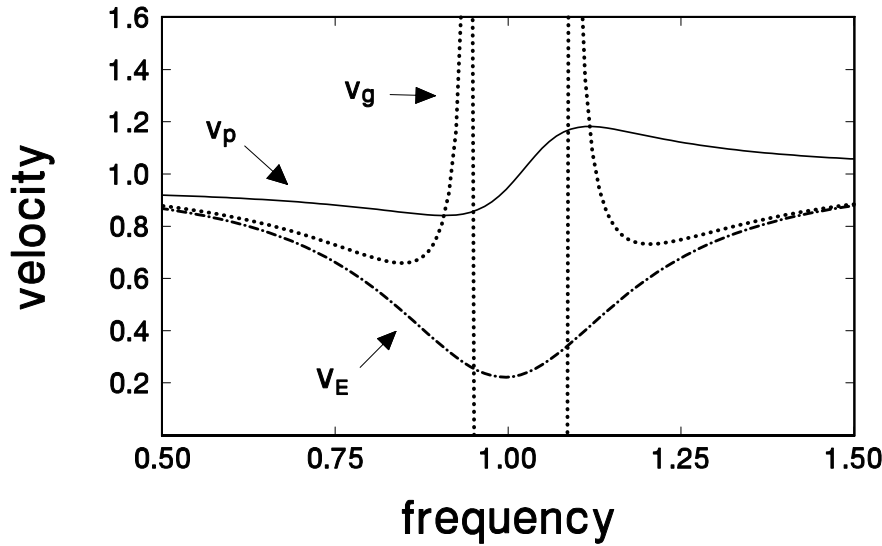


Figure 3: Phase velocity v_p (solid line), group velocity v_g (dotted), and transport velocity v_E (dash-dotted) of scalar waves propagating through a collection of point scatterers having an internal resonance with frequency E_0 at a scaled density $2\pi nr_0^3 = 0.0028$, as a function of the scaled frequency E/E_0 . The damping $1/(T_2 E_0) = 0.1$. Note the large reduction of v_E from one.

The presence of the fascinating delay effect has also a certain disadvantage. When, for instance in search for localisation, one is looking for an experimental verification of some reduction of the mean free path by measuring the diffusion constant, one should be prepared for the fact that a large reduction of the diffusion constant can also be brought about by the velocity v_E , an effect which is argued not to be connected to localisation at all. That is why we have emphasized so much the difference between steady-state and transient diffusion. In steady-state measurements all dynamic delay effects are absent and the mean free path is probed.

6 Multiple backscattering

Localisation of light is connected to the importance of interference of multiple scattering of waves. This effect has not been observed unequivocally yet in three dimensions. Another interference effect in multiple scattering that has been observed, and that by now should be well-known, is considered to be a precursor of localisation of light. It

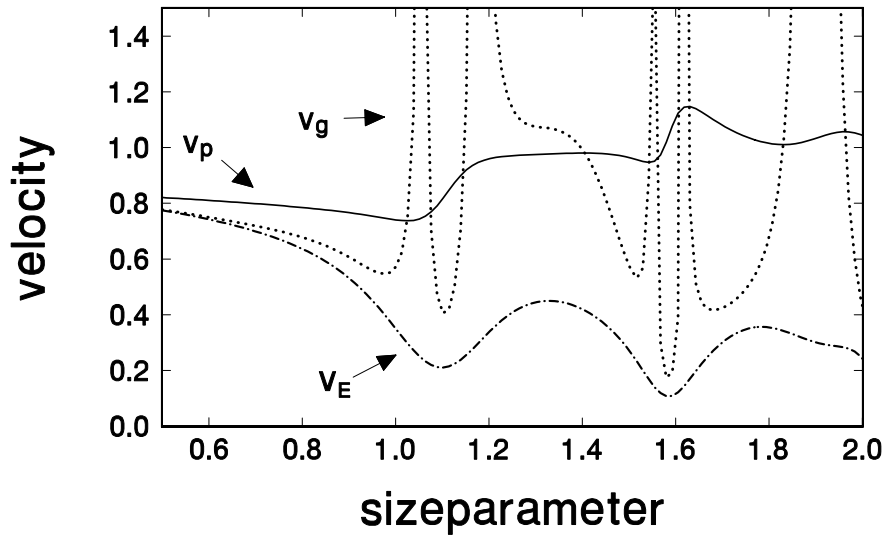


Figure 4: Phase velocity v_p (solid), group velocity v_g (dotted), and transport velocity v_E (dash-dotted) of vector (light) waves propagating through a collection of finite size scatterers at a volume fraction of $f = 0.2$ and index of refraction of the scatterers $m = 2.7$, as a function of the scaled frequency Er_0 ($\equiv$ size parameter x). Note the large reduction of v_E from one.

has acquired several different names: “weak localisation”, “coherent backscattering” or “enhanced backscattering”. Long before localisation is believed to set in (that is long before the transport mean free path vanishes), one can observe this interference effect. The idea is very simple and appealing. Suppose a wave propagates in a random medium from A to B. One has to sum all possible paths and square the result to obtain the final probability for arriving at B. (See Fig. (5a)). The outcome would consist of the sum of squares (incoherent contribution) and the rest (interference terms). Given the random character of the phases of the interference terms one could assume that all interference terms vanish on the average. On large length scales ($q\ell \ll 1$) this brings us back to the conventional theory of diffusion. We have explained that this is absolutely wrong when one is in the localisation regime, but even in the regime of weak multiple scattering it is possible to pinpoint a serious flaw in the assumption of vanishing interference terms. Imagine one wants to propagate a wave from A to A. In Fig. (5b) two possible paths are depicted: a randomly chosen path and the same path but traversed in opposite sense (time-reversed path). These two routes always interfere constructively and should be added “coherently”. As a consequence the *returning* probability of this wave is *twice* as large as would be expected on the basis of a theory which would only include the incoherent terms: $1^2 + 1^2 = 2$ and $(1 + 1)^2 = 4$. This constructive interference will always exist, no matter how long and intricate the path, provided that no inelastic scattering occurs along the road. This constructive interference is already present in lower orders of multiple scattering in contrast to (strong) localisation which needs very large orders of scattering. Would it be possible to observe this constructive interference effect occurring in already in lower-order multiple scattering directly? The answer is definitely affirmative for light scattering. Indeed, the constructive interference also

applies to events not starting in the medium but beginning outside the sample. This type of paths is typically monitored in a conventional scattering experiment. Paths returning to their starting point correspond to backscattering in such an experimental set-up. This interference effect of lower-order multiple scattering events, known as enhanced backscattering, has clearly been observed.[13, 14]

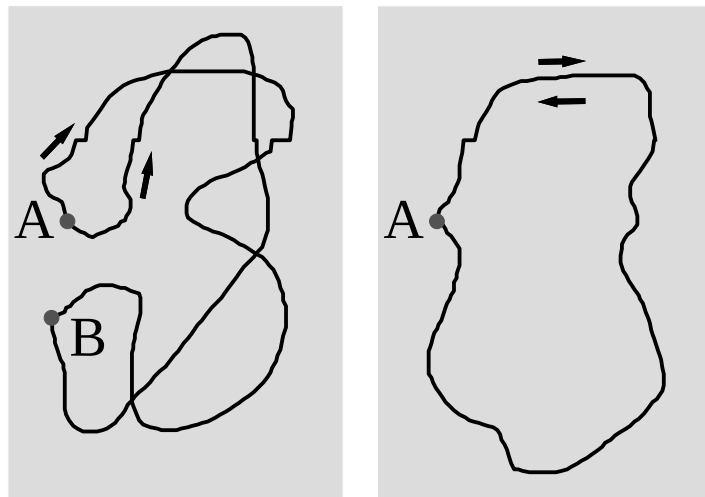


Figure 5: Left part (a): Two randomly chosen paths in a disordered medium from A to B. Right part (b): A randomly chosen path in a disordered medium from A to A and its time-reversed counterpart. Arrows indicate direction of propagation.

Let me clarify the close connection between weak localisation (which actually is no localisation at all!) and strong localisation. Enhanced backscattering pertains to an enhanced probability of returning towards (the neighbourhood) of the starting point. This indicates that a random walker gets less far than was anticipated on the basis of a picture in which interference was neglected: the mean squared displacement has decreased. In other words the mean free path of the random walker gets reduced due to interference. If owing to the presence of a large degree of disorder the reduction due to interference becomes complete (reduction of 100%) the random walker will be localised.

7 “Amplitude” Properties

The purpose of this section is to show how the density of states of a cavity changes when resonant scatterers are present in the cavity. I will use for scatterers again the generic model of point scatterers. As is well-known in optics when light propagates through scattering media it is important to discriminate between the (spatially exponentially attenuated) “coherent” beam and the “diffuse” intensity. In the extreme case of very weak scattering there is almost no scattering and thus no diffuse intensity; in the other extreme case of strong scattering the mean free path is very short and there is almost no coherent beam. One sometimes says that the coherent beam is an “amplitude” property and the diffuse beam an “intensity” property. Of course the intensity in the coherent beam is an intensity and not an amplitude but it consists of the cross product of a scattered beam and the homogenous incoming beam and that is why it is

called an “amplitude” property. In contrast the diffuse intensity contains the scattered beam *squared*, and for that reason is referred to as an “intensity” property. Quantities like the index of refraction, extinction coefficient, group velocity and phase velocity are all “amplitude” properties. Localisation is concerned with the transport of scattered waves and is therefore an “intensity” property. However, it is clear that amplitude properties (I will drop the quotation marks from now on), which feature in Huygens’ principle, and will seriously influence transport properties. Another way of indicating the intimate relation between amplitude and intensity is pointing to the existence of conversation laws, like the conservation of energy, that contains both amplitude and intensity properties. So a calculation of amplitude properties is an essential ingredient of an intensity calculation. If some amplitude property suffers from a dramatic experience, the intensity properties will also feel a backlash. In particular if a calculation of the amplitude tells us there is almost no density of states at a certain frequency, transport at that frequency will seriously be affected by this small density of states.

Let me review some amplitude properties of a collection of simple point scatterers placed in a homogeneous background. (See Loudon [4]) In some contexts this is known as polariton theory. In almost any book on optics one can find figures that visualise how the index of refraction in the visible of air is changed when a collection of atoms, usually taken to be sodium, is added to the air. These pictures look very much like the solid line in Fig. (6a). This particular figure has been calculated by using some typical parameters for a collection of point scatterers.

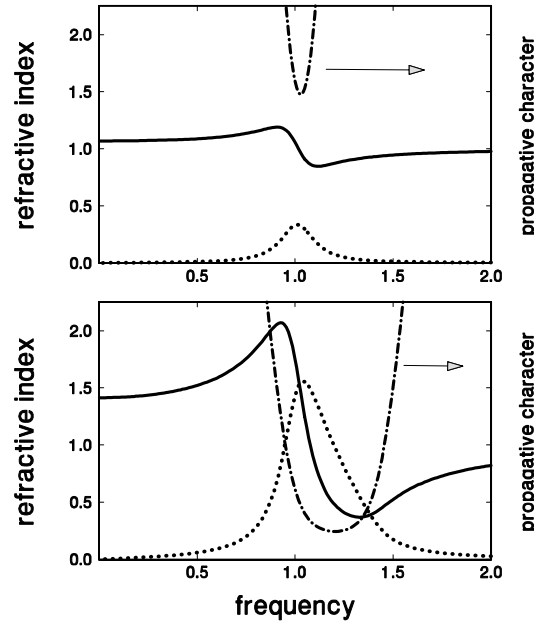


Figure 6: The real part η (solid line), imaginary part κ (dotted) of the index of refraction and the product $k'\ell_{ext}$ (dash-dotted) of a collection of point scatterers having an internal resonance at frequency E_0 as a function of the scaled frequency E/E_0 for two scaled densities (a): $2\pi nr_0^3 = 0.0028$ and (b): $2\pi nr_0^3 = 0.02$. The damping $1/(T_2 E_0) = 0.1$. Note that in the high-density case there is a large region where the imaginary part is larger than the real part (evanescent wave).

The behaviour of the average amplitude is described by a complex valued wavevector k and by a complex index of refraction m :

$$k \equiv k' + ik'' \equiv Em \equiv E(\eta + i\kappa), \quad (11)$$

Extinction, existing of absorption and scattering, is connected with the imaginary part of the wavevector k'' or the imaginary part of the index of refraction κ . These quantities describe how the amplitude is exponentially attenuated. With this exponential decay a characteristic length, the extinction mean free path ℓ_{ext} can be associated. We find the following relation

$$\ell_{ext} \equiv 2/k''. \quad (12)$$

The factor of two originates from the fact that extinction mean free path is defined as the characteristic decay length of the coherent *intensity* and not amplitude. Extinction can be brought about by two decay channels: elastic scattering and absorption. Both processes have their own cross-sections and mean free paths associated with them: $(\ell_{ext})^{-1} = (\ell_{scat})^{-1} + (\ell_{abs})^{-1}$. When scattering is predominantly elastic $\ell_{scat} \approx \ell_{ext}$, but when the extinction is mainly brought about by absorption $\ell_{scat} \gg \ell_{ext}$. In the low-density Boltzmann regime, apart from a geometrical factor of order one, the scattering mean free path ℓ_{scat} is equal to the transport mean free path ℓ .

For our discussion we will need the phase velocity. In vacuum the dispersion relation is $E = k'$. For our medium filled with point scatterers this becomes $E = k'/\eta$. A sufficiently accurate definition of the phase velocity is $v_p \equiv E/k' = 1/\eta$. It is this result that was plotted in Fig. (3). The group velocity, relevant for pulse propagation of the coherent beam, is defined as the derivative of the dispersion relation $v_g = \frac{\partial E}{\partial k'}$. In Fig. (3) the group velocity has also been plotted. The divergences in the group velocity are connected to “S-curves” in the dispersion relation. These divergences become even more apparent in Fig. (4) where the group velocity has been plotted for a collection of finite-size spheres. The divergences in the group velocity are nothings to worry about. The group velocity describes the propagation of a pulse in the coherent beam, but due to strong scattering almost all intensity is in the scattered beam and not in the exponentially attenuated coherent beam.

An important property of our system of point scatterers is the density states per unit volume. Whether or not the density of states is an amplitude or an intensity property I will not discuss here since this is a subtle point. I will only count the states that obey the dispersion relation and call them *propagating* states. A solution of the wave equation can be characterised by its frequency E and its momentum variable $\mathbf{p}$. In vacuum there is no dispersion, $E^2 = p^2$, and all states are “on-shell”. The distribution function that counts the number of states with energy E and momentum $\mathbf{p}$ is defined to be spectral density $S(E, p)$. Summation over all momenta of this quantity will give the *density of states* $\rho(E)$. To comply with conventional treatments in Green’s functions we normalise $S(E, p)$ so that integration over momenta gives 2π times the density of states. So for vacuum we find

$$S(E, p) = 2\pi E \delta(E^2 - p^2). \quad (13)$$

This gives indeed as the density of the states per unit volume

$$\rho(E) = \frac{E^2}{2\pi^2}, \quad (14)$$

showing the well-known quadratic frequency dependence of the density of states. In the presence of scatterers the spectral density of propagating states, that are the states obeying the relation $E = v_p p$ is

$$S(E, p) = 2\pi E \delta(E^2/v_p^2 - p^2). \quad (15)$$

This gives the following density of propagating states

$$\rho(E) = (1/v_p) \frac{E^2}{2\pi^2}. \quad (16)$$

This result is most important result of this whole section. We see that an increase in phase velocity causes a depression in the density of states. In Fig. (7) I have plotted the density of propagating states for an empty cavity and for a cavity filled with point scatterers for two different densities. From the figure one infers that a resonance induces both an increase and a decrease of the density of states compared to the empty cavity. The decrease always occurs at higher frequencies than the resonance frequency and is connected to the fact that the scatterers (the oscillators) are lagging in phase the driving electric field.

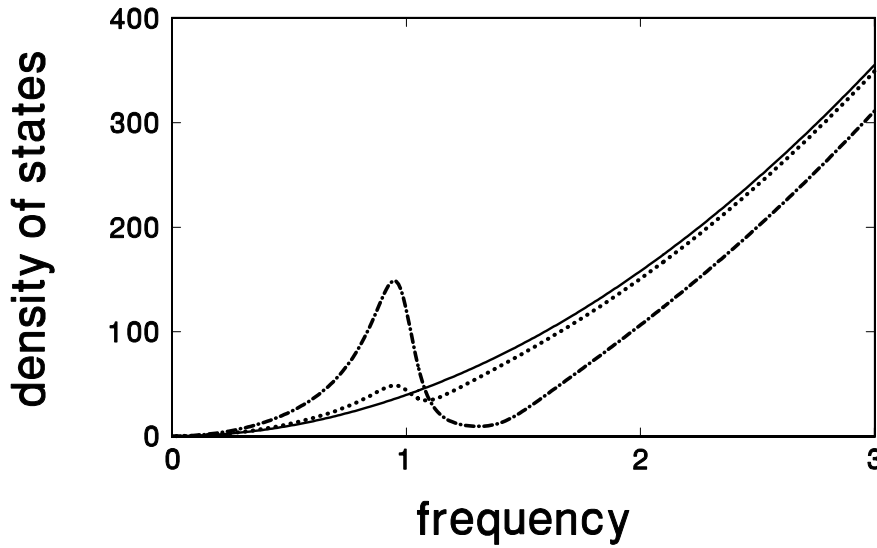


Figure 7: *The density of propagating states of a collection of point scatterers having an internal resonance at frequency E_0 as a function of the scaled frequency E/E_0 for three scaled densities $2\pi n r_0^3 = 0$ (empty cavity, solid line), $2\pi n r_0^3 = 0.0028$ (dotted), and $2\pi n r_0^3 = 0.02$ (dash-dotted). The damping $1/(T_2 E_0) = 0.1$. Note the depression of the density of states at frequencies slightly above the resonance for a filled cavity.*

8 Towards Localisation

To describe localisation of light it is extremely instructive to compare localisation of light with the occurrence of evanescent waves. I will indicate why. We call an

amplitude evanescent when the imaginary part of the wave vector k'' becomes larger than the real part k' : the amplitude is overdamped. An amplitude is overdamped when

$$k'\ell_{ext} < 0.5. \quad (17)$$

This criterion does not signify any sudden spectacular change. It only means that propagation of the amplitude becomes more and more difficult and criterion (17) is a useful beacon to indicate this gradual change. It is clear that the narrower the resonance or the larger the concentration of scatterers the more states will become evanescent. In Fig. (6) the product $k'\ell_{ext} = 0.5k'/k'' = 0.5\kappa/\eta$ has been plotted, together with both the real and imaginary part of the index of refraction η and κ separately for a small and a large density of point scatterers. We notice that the product is smallest for frequencies slightly larger than the resonance frequency. We also see that for a large density of scatterers many propagating states become evanescent, as indicated by criterion (17).

Localisation of light is a transport property that says that if in a strongly elastically scattering medium a certain criterion is met there will be no transport of the intensity. As the most important length scales for interference in elastic scattering are the wavelength and the scattering mean free path we expect that this localisation criterion can be related to the product $k'\ell_{scat}$. One believes that the localisation transition in the intensity occurs when the product $k'\ell_{scat}$ becomes less than some critical value α_c ,

$$k'\ell_{scat} < \alpha_c \approx 1 \quad . \quad (18)$$

The actual value of α_c differs from situation to situation but is always of order one. I want to stress that condition (18) marks a “phase transition” in transport whereas condition (17) only indicates a gradual change of the character of the amplitude. The condition for evanescent waves gives of course a lot of insight on the phenomenon of localisation. Apparently near localisation many states are close to being evanescent and presumably to induce localisation it helps to have a small density of propagating states. This explains why many workers in this field try to suppress the density of states as much as possible, for instance by using three-dimensional periodic arrays of dielectric materials (“dielectric crystals”). Dielectric crystals might have a region in frequency space where the density of states is small or even vanishing (gap). The amplitude will have great problems in propagating through these structures and hopefully the intensity will get localised easier.

The reader might wonder why localisation has not been observed yet. We are again confronted with the difference between electrons and scalar waves as product (18) can be made as small as possible for electrons by going to small frequencies (giving rise to small wave vectors). This is feasible because for electrons the scattering potential is energy independent (See equation (1)). For classical waves small frequencies (Rayleigh limit) will lead to a weak scattering potential (See equation (2)) and to corresponding long mean free paths. Large frequencies bring us to the geometric optics limit where scattering is inefficient and localisation is unlikely. Apparently if localisation exists at all for a certain sample, it will be confined to a window in the frequency domain. It turns out that criterion (18) is extremely difficult to satisfy in practice for any optical system. The reader might wonder why this is so difficult as criterion (17) can easily be satisfied, for instance near stopbands. The catch is that the extinction mean free path is usually much shorter than the scattering mean free path and for many systems criterion (17) can be satisfied without even approaching criterion (18). To find localisation we have to minimise the product $k'\ell_{scat}$ for which we need elastic scattering to dominate over

absorption. A naive view would be to minimise the mean free path but Fig. (6) clearly demonstrates that the influence of the wavevector (and so density of states) is also very important. For the behaviour of the product $k'\ell_{ext}$ the behaviour of the wavevector is overwhelming. That is why all the action occurs at frequencies *larger* than the resonance frequency.

Obvious parameters to optimise in order to reach localisation are the cross-section of the individual scatterers and the density of scatterers. Increasing the cross-section of the individual scatterers immediately brings us to the use of scatterers on resonance. One should realise that there is a “unitary” upper limit to the cross-section and it is of order wavelength squared, $\sigma_{max} = 4\pi/E^2$. Increasing the density of scatterers brings us in difficulties as well. Up to what point do our estimates of mean free paths and wavevectors still hold in the high-density limit? This is easily checked. We will assume that our formula’s keep on to hold and will look for an inconsistency. We assume that in an elastically-scattering medium $\ell_{scat} = 1/n\sigma$ is still valid at large density of the scatterers. For an order of estimate we assume that $k' = E$. Let us consider two scatterers taking part in a scattering event. They are a distance ℓ_{scat} apart. The optical “size” of the scatterers is of order $\sqrt{\sigma}$. We will draw a sphere around scatterer one with radius ℓ_{scat} . When scatterer one is receiving an incoming wave it will scatter it independently of the other scatterer only if the cross section of the second scatterer is a small fraction of the total area of the sphere. That is if the inequality

$$4\pi\ell_{scat}^2 \gg \sigma, \quad (19)$$

is satisfied. We call this the independent-scattering limit, and this is the basis of all heuristic arguments. In this limit all our formula’s are reliable. If our scatterers are on resonance $\sigma = 4\pi/E^2$ the inequality can be rewritten as

$$k'\ell_{scat} \gg 1. \quad (20)$$

To our great surprise we find that the condition for localisation of light, as expressed by inequality (18), is in conflict with our request for independent scattering of resonating scatterers, as expressed by inequality (20). Leaving the regime of independent scattering is a tremendous complication for a theoretical description as several length scales become important and no small parameters can be identified. This is also a challenge as this regime is more or less virgin territory. An important but unanswered question is: is localisation going to survive in this regime? Dependent-scattering corrections can either help or hinder the occurrence of localisation. Many workers in the field suggested that one should use very high-Q scatterers (like atoms or Mie spheres with a very large contrast in index of refraction) and pack a lot of them together. This is, however, exactly the situation where dependent-scattering corrections are going to be extremely important. Well, will they help or won’t they? We have given a partial answer to this question.[15] It seems that the dependent-scattering corrections *help* for frequencies slightly larger than a resonance frequency and that they *hinder* for frequencies slightly smaller than a resonance frequency. This confirms our general findings that below and beyond a resonance are physically two separate regimes.

The dependent-scattering corrections are to a large degree determined by the scattering mean free path. These corrections are much less important when extinction is largely dominated by absorption, a situation typical for many polaritons. In Fig. (6b) there is a regime where inequality (17) is satisfied. If this figure would refer to a medium in which extinction is dominated by elastic scattering and not by absorption I would have to introduce large dependent-scattering corrections (which would be very difficult). So that part of the figure can only be trusted for polaritons where absorption dominates.

9 Remaining questions

Rather than answering the question in the title I am going to throw up a small number of additional questions that have to be addressed in the future.

Many interesting phenomena occur at either sides of resonances. Developing theories for dependent scattering is crucial in this regime. Recent studies show that there is an intimate connection between dependent scattering and induced dipole-dipole coupling in electromagnetic theory. The study of this should be pursued.

Many people, either experimentalists or workers doing numerical simulations, operate on only one or a relatively small number of realisations of a finite random system. An important issue is the averaging over realisations. People in optics know how difficult it sometimes is to get rid of speckle. Are all the relevant intensity quantities self-averaging? That is to say: is the variance of a quantity vanishing for an infinite system? For an infinite system there is no difference between one and many realisations. It is quite possible that some properties do not self average. Are transport properties self averaging? Is averaging necessary for localisation? Even the definition and measurement of transport properties of one realisation is an open question. What usually can be measured quite easily in practice is the mode structure of one realisation, but the connection to transport is not clear at all. This question is particularly important for the detection of localisation in dielectric crystals.

The field is lively and full of unresolved questions. Very often scientists end their papers by enumerating the great variety of applications that will become accessible due to the advances made by fundamental research in their field. This lip service might be necessary for survival, in many cases it is overstated. Notwithstanding this comment I will make a remark of this type. Multiple scattering of light and other classical waves is a common phenomenon in our daily lives. The pursue of the field of localisation of light has led to an enormous increase in the understanding of this phenomenon with potentially important spin-off. This traffic is a two-way process as fundamentally oriented scientists get acquainted with more applied aspects of their work in a natural way.

I am grateful to Meint P. van Albada for discussions and to Bart A. van Tiggelen for discussions and a critical reading of this manuscript. This work is part of the research program of the “Stichting voor Fundamenteel Onderzoek der Materie” (FOM) and was made possible by financial support of the “Nederlandse Organisatie voor Wetenschappelijk Onderzoek” (NWO).

References

- [1] T.V. Ramakrishnan, in: *Chance and Matter (Les Houches, 1986)* eds. J. Souletie, J. Vannimenus and R. Stora, (North-Holland, Amsterdam, 1986).
- [2] B. Souillard, in: *Chance and Matter (Les Houches, 1986)* eds. J. Souletie, J. Vannimenus and R. Stora, (North-Holland, Amsterdam, 1986).
- [3] *Scattering and Localisation of Classical Waves in Random Media*, edited by Ping Sheng (World Scientific, Singapore, 1990).
- [4] R. Loudon, *The Quantum Theory of Light*, (Oxford, Oxford, 1973) first edition, p. 61.
- [5] Th.M. Nieuwenhuizen, A. Lagendijk, and B.A. van Tiggelen, Phys. Lett. A. **169**, 191 (1992).
- [6] M.P. van Albada, B.A. van Tiggelen, A. Lagendijk, and A. Tip, Phys. Rev. Lett. **66**, 3132 (1991).
- [7] B.A. van Tiggelen, A. Lagendijk, M.P. van Albada and A. Tip, Phys. Rev. B **45**, 12233 (1992).
- [8] B.A. van Tiggelen and A. Lagendijk, to be published in Europhysics Letters.
- [9] B.A. van Tiggelen and A. Lagendijk, Comment to Ref [10] to be published in Phys. Rev. Lett.
- [10] Yu.N. Barabanenkov and V.D. Ozrin, *Phys. Rev. Lett.* **69**, 1364 (1992).
- [11] M. Büttiker in *Electronic Properties of Multilayers and Low-Dimensional Semiconductor Structures*, edited by J.M. Chamberlain et al. (Plenum, New York, 1990).
- [12] B.A. van Tiggelen, A. Tip and A. Lagendijk, J. Phys. A **26**, 1731 (1993).
- [13] M.P. van Albada and A. Lagendijk, Phys. Rev. Lett. **55**, 2692 (1985).
- [14] P.E. Wolf and G. Maret, Phys. Rev. Lett. **55**, 2696 (1985).
- [15] B.A. van Tiggelen, A. Lagendijk, A. Tip, and G.F. Reiter, Eur. Phys. Lett. **15**, 535-540 (1991).