

## Huygens’s principle

(Talk given on 22 May 1924)

By HADAMARD

Translated by D. H. Delphenich

---

### I.

As one knows, Huygens’s principle first presented itself at the very birth of the wave theory of light. It was essential to one of the first, and more memorable, stages of development of that theory that one had to explain the fact that light propagated along straight lines, which appeared paradoxical in the wave theory. In order to provide that explanation, one had to analyze the effect that is produced by a signal of very short duration that is emitted perceptibly from a unique point  $O$  and consequently perceived at an arbitrary point  $A$  in space after a time interval  $T$  that equals the quotient of the distance  $OA$  with the velocity  $c$  of light. The brilliant idea of Christian Huygens consisted of considering the state of the medium between  $O$  and  $A$ , not at the initial or final instant, but at an intermediate instant, and substituting the consideration of the latter state for that of the original perturbation entirely.

Fresnel gave a form to the argument that was thus glimpsed by Huygens that was already close to the one that it presents to us today. However, his ideas did not triumph without some difficulties, and one will find an analysis of the spirited polemic that arose between Fresnel and Poisson in Poincaré’s *Leçons sur la Théorie mathématique de la lumière*. It took the emergence of Kirchhoff and his successors, among which, in addition to Beltrami and Maggi, I will mention Duhem and Volterra, to explain the most important of those ideas. A reading of books such as Duhem’s *Hydrodynamique, Élasticité, Acoustique* will suffice to show how delicate the question is.

It is on the topic of that question that I would like to speak today. By retracing the state, indicating how Fresnel’s arguments can be legitimized and the objections that Poisson raised, I hope to also show, not only that several interesting and important problems present themselves, but furthermore, that those problems provoke or clarify a number of interesting questions in various domains of analysis, which would be normal in a science like ours into which everything fits.

Of course, a large part of the difficulties and discussions that were raised by Huygens’s principle came from the fact that it was poorly posed and the fact that the authors that treated it did not attach all of the same significance to that principle that I spoke of.

If I have the temperament of a historian, and I must assume that the majority of those who are listening to me do so well, then in order to educe that significance, I must obviously appeal to the texts by commencing with the oldest of them and reproducing the essential lines of reasoning

in the polemic that took place between Fresnel and Poisson by quoting and commenting upon the passage from Huygens, as Poincaré did in his *Théorie mathématique de la lumière*. However, despite the interest that the historical viewpoint might arouse in us and despite the beautiful victories that have been won on that battlefield by French scholarship, I believe that I see more mathematicians around me than historians of mathematics, properly speaking, and for my own part, I must confess to a certain congenital incapacity in that regard. Reading it is what gives me the most anxiety in my life and I have definitely given up learning it. Under those conditions, instead of entering into a territory that is not my own, I shall be content to recall the argument that was presented in the *Traité de Lumière* <sup>(1)</sup>, and for the rest, I shall refer to the cited book by Poincaré, after which, if you will permit, I will take up the matter from the opposite end: I shall address the argument in its form that is classical today, after which I might possibly study, with much prudence, if not timidity, how that same argument was presented in the school of its original creators.

As I said, there are a certain number of subtle distinctions to be pointed out, and there is one of them that we must say a word about before entering into the debate. Indeed, it is the overpowering question of what one means by a wave, and one knows that this notion is taken to have two very different accepted meanings, according to the situation. The one that Huygens started with is nothing but the one that has been made more precise by Hugoniot since then. To fix ideas, in a medium that we suppose to be initially at rest, we produce a mechanical or electromagnetic disturbance at a certain point. Starting from that moment, space will be divided into two regions that we call 1 and 2 at each instant. The first one can be influenced by only the perturbation, while everything still remains at rest in the region 2. The separation surface between those two regions, or *wave front*, will vary with time, moreover, in such a way that the same molecule will belong to 1 or 2 depending upon the instant when one considers it: *The first region will propagate towards the second one*, i.e., when one makes  $t$  increase, the region 2 will tend to include, little by little, molecules that are progressively enveloped by the region 1. The moment when the wave front reaches a well-defined molecule is characterized by a brief variation in the state of that molecule (for example, a brief variation in its acceleration or one of its higher-order accelerations, if one is dealing with an ordinary motion).

However, it is well-known that physicists do not generally see things that way. In most of their theories, the propagation by waves applies to periodic motions that are sinusoidal. In that case, there will be no division of the medium into regions whose states vary progressively in time; on the contrary, one must suppose that a permanent regime has been established with no brief variations. All phenomena are expressed with the aid of sines or cosines, which are perfectly regular functions whose derivatives of all orders are finite and continuous.

The intuition of physicists has taught them to recognize the kinship between those two types of phenomena, while distinguishing between them. For the analyst, they would seem to be profoundly different to begin with.

Meanwhile, one must concede that a rapprochement is established between a quantity that varies briefly and a quantity that submits to sinusoidal variations that are hardly appreciable, but extremely rapid. Since I have promised to point out to you the various mathematical developments

---

<sup>(1)</sup> Pages 21 and following in the edition of *Maîtres de la pensée scientifique*, Paris, Gauthier-Villars, 1920.

that Huygens's principle gave rise to, I will say a word about how the analytical translation of that rapprochement illuminated the study of what Klein called *approximate mathematics*, which is when one takes into account the fact that nothing is known to us that does not have a certain degree of error to it, and we shall prove that this approximate mathematics must often inspire, not the simplest parts of exact mathematics, but rather the most delicate ones. In any case, such a rapprochement is plainly shown to be at work in the theory that we address: Whether one deals with waves from the viewpoint of Hugoniot or vibratory waves, the calculations are absolutely parallel, and that will often be true of their most minor details and ultimate consequences.

I would not have insisted upon that distinction if it had not factored precisely in the polemic between Fresnel and Poisson. Having to respond to Poisson's objections, Fresnel believed that he could find that response in the vibratory character of the propagating motion and in the behavior of the corresponding interferences, which is classical today. Conforming to what was just said, that fact does not, in reality, have the importance that Fresnel attributed to it in this case, and the true explanation to which he arrived a bit later, as one can read in Poincaré's book, applied just as well to what one calls a "solitary wave," which corresponds to a periodic system of oscillations in Hugoniot's conception of things.

The first question is therefore vacuous: Although there is most assuredly a distinction to be made between the two senses of the work wave, we can generally use them interchangeably today by considering the first one to be, above all, the one to which the name of Hugoniot is attached, moreover.

Having said that, we have pointed out that Huygens's principle involves three successive instants: The first one  $t_0$ , when one is given an initial disturbance, an intermediate instant  $t_1$ , and a final instant  $t_2$ , when one proposes to calculate the effect that is produced. In its form that is currently classical, the argument can be decomposed in the following manner:

*A. In order to deduce the effect that is produced at a later instant  $t_2$  by a known phenomenon at the instant  $t_0$ , one can begin by calculating the effect at an intermediate instant  $t_1$  and then starting from the latter in order to deduce the effect at  $t_2$ .*

*B. If the initial perturbation at the instant  $t_0$  is localized in the neighborhood of a well-defined point  $O$  then its effect at the instant  $t_1$  will be zero everywhere except in the neighborhood of a sphere  $S_1$  with its center at  $O$  and a radius of  $c(t_1 - t_0)$ , in which  $c$  denotes the speed of propagation.*

*C. From the standpoint of its effect at the final instant  $t_2$ , the initial perturbation can be replaced with a system of perturbations that take place at the intermediate instant  $t_1$  and are distributed conveniently over the surface of the sphere  $S_1$ .*

One sees that one then has something in the manner of a syllogism in that, where A is its major premise, B is the minor one, and C is the conclusion.

It would seem convenient to examine those three propositions in their natural order. However, I am decidedly condemned to take everything backwards today, because we shall, if you would really like to, begin with the last of them, namely, the conclusion.

In the absence of logic (or what passes for logic at the moment), I must say that I have history on my side this time. Indeed, after Huygens himself, along with Fresnel and Poisson, the historical order of events led us to the work of Kirchhoff, in which some of the points that had preoccupied the first three authors that were cited just now were subjected to detailed calculation for the first time.

Now, as we shall see, it was proposition C that Kirchhoff began with, and he did not call upon the major premise A or the minor one B at any moment. Consequently, nowadays we see that the logical link that unites our three propositions is not as close as it first seemed to us: Indeed, we will be led to consider them to be relatively independent and to apply very different judgements to each of them.

## II.

In order to address the question analytically, one must obviously (and this is what Kirchhoff did) begin with the mathematical law of the phenomenon. Under the conditions that Huygens imposed, that law is expressed by the *spherical wave* equation, i.e., by the partial differential equation:

$$(E) \quad \frac{1}{c^2} \frac{\partial^2 u}{\partial t^2} = \frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} + \frac{\partial^2 u}{\partial z^2} .$$

Under the conditions that Huygens imposed, here is what that must say: We suppose that we are dealing with a perfectly homogeneous and isotropic medium in which, on the other hand, passive resistances (such as viscosity) play no role, and it is ultimately left to itself. That medium will be the ether if one is dealing with light, and air, or any other homogeneous gas, when one is dealing with the propagation of sound, to which everything that we have said today is applicable. Everywhere that the preceding conditions all apply, the phenomenon will depend upon one or several quantities  $u$  that are functions of  $x, y, z, t$  and verify the partial differential equation (E). The simplest case is the one in which the medium fills up all of space, in other words, one deals with the ether where no material bodies exist. For example, we can suppose that this is the case and that the ether is left completely left to itself in all of space *after* the instant  $t_0$  when we emitted a light signal (but not *during* the emission of that signal, since the ether would then cease to be left to itself and its oscillations would be disordered).

Kirchhoff proposed to establish Huygens's principle (and for the moment, that principle is understood to take the form C) by studying a phenomenon that is governed by the spherical wave equation and consequently, in short, by integrating that equation.

That had been done for the first time by Poisson in 1819 in the simplest case that we spoke of a while ago, namely the one in which the medium fills up space completely and is left to itself starting from the instant  $t_0$ , while supposing, on the other hand, that we are given the state of the phenomenon at that initial instant  $t_0$ . The problem is formulated analytically as the study of an unknown function  $u$  that is a solution of equation (E) and is characterized, in addition, by the following conditions that relate to  $t = t_0$  ("defining" conditions):

$$(1) \quad \begin{cases} u(x, y, z, t_0) = g(x, y, z), \\ \frac{\partial u}{\partial t}(x, y, z, t_0) = h(x, y, z), \end{cases}$$

in which  $g$  and  $h$  denote two given functions. That is what one calls a *Cauchy problem*, and the one that Poisson posed is a *well-posed* problem, i.e., the combination of the indefinite equation (E) and the defining conditions (1) determines one and only one function, which Poisson obtained by a particularly simple formula that is all of the articles.

Kirchhoff imposed some other conditions. He assumed the existence of centers of perturbation that were point-like or extended and in which the partial differential equation would cease to be valid. For that reason, and also in order to take into account the possible presence of bodies that were capable of modifying the phenomenon, he no longer supposed that the indefinite equation was verified in all of space, but only in the region  $R$  of that space that is external to one or more closed surfaces  $\sigma$ .

Kirchhoff achieved the integration under those new conditions. Now, the result to which he arrived can be interpreted by saying that the phenomenon that thus plays out in the region  $R$  can be considered to be generated by a system of disturbances that are produced at convenient instants and with convenient amplitudes at conveniently-chosen points on the surfaces  $\sigma$ . One sees that this is proposition C, but in a slightly more general form: It will become proposition C itself when the phenomenon under study is due to a unique perturbation that is produced at the instant  $t_0$  at a point  $O$  and one takes  $\sigma$  to be a sphere of center  $O$ .

In that case, proving Huygens's principle will then amount to only integrating the differential equation for spherical waves.

At that point, a problem obviously presents itself, namely, the problem of doing what was just done for the partial differential equation of spherical waves for other equations and proving Huygens's principle in the form C by integrating them.

Indeed, it is clear that the constitution of the medium will not always be the one that we have supposed up to now and that the partial differential equations must be modified as a consequence: For example, in our terrestrial atmosphere, with its variations due to altitude, the propagation of sound will depend upon a different equation. Other physical phenomena will lead to other equations.

Let us then see what we can say about those new equations insofar as their integration is concerned, and consequently the property C, and even first of all see how the notion of a wave that defines the main topic of our present study will present itself for them.

First consider a rectilinear telegraph cable that is homogeneous in all of its parts. In order to represent the variation of the electric state along such a cable, it is obviously convenient to make time the ordinate, as it is in railroad charts, and consequently trace out a two-dimensional space-time diagram. The cable considered will be represented by a line at the origin of time, namely, the  $x$ -axis, and the same cable, when taken at the instant  $t = 1$  (for example, after 1/1000 second), by a line that is parallel to the first one and is located at a distance of 1 from it, and so on. One knows that any of the quantities upon which the electric state depends (the potential, for example) satisfies the partial differential equation (the "telegraph" equation):

$$(\mathcal{E}) \quad A \frac{\partial^2 u}{\partial t^2} + 2B \frac{\partial u}{\partial t} = C \frac{\partial^2 u}{\partial x^2},$$

in which  $A$ ,  $B$ ,  $C$  denote three positive constants. Having said that, suppose that the cable is originally in the neutral state (where  $u$  is identically zero) up to  $t = 0$ , and one then emits a signal at a well-defined point  $x_0$  on the cable at that instant: The perturbation, thus-produced, will propagate in both directions with a constant velocity  $\omega = \sqrt{C/A}$ , properly speaking, in the form of two waves whose advance will be illustrated on our graph by two half-lines with angular coefficients  $\pm 1/\omega$ . If a point  $(x_1, t_1)$  on our diagram is on one of those two lines then that will signify that the wave that propagates our perturbation arrives at the point  $x = x_1$  at precisely the instant  $t_1$ . If the point  $x_1$  is given, which will be, for example, the position of a receiving post, then the preceding condition will determine a certain instant  $t'_1$  (which corresponds to the ordinate of our half-line), which will be the one when the signal that issues from  $x_0$  at the origin of time will be perceived. We say that the point  $(x_1, t'_1)$  on the diagram is *on the same wave* as  $(x_0, 0)$ . For  $t_1$  less than  $t'_1$ , the point  $(x_1, t_1)$  will be outside the angle of the half-line: One can say that it is *outside the wave* whose center of initial perturbation is  $(x_0, 0)$ , and under those conditions, the signal that issues from the latter center can have no effect at  $(x_1, t_1)$ , i.e., at the instant  $t_1$ , and it can have no effect at the receiving post  $x = x_1$  either. Finally, a third case is obviously the one in which the point  $(x_1, t_1)$  is inside the angle of our two half-lines (so  $t_1 > t'_1$ ). We are then dealing with an instant  $t_1$  that is *later* than the one when the wave that propagated the signal passed the receiving post. In the latter case, we say that the point  $(x_1, t_1)$  on our diagram is *inside of the wave* with  $(x_0, 0)$ .

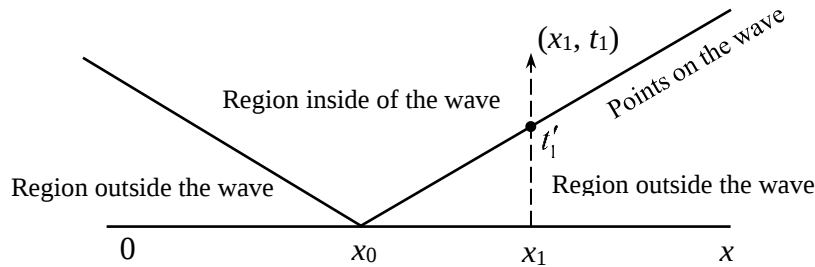


Figure 2.

Now, if the cable is not homogeneous, so its electrical properties vary (in a continuous manner) from one point to another, then the partial differential equation will have a different form: The speed of propagation of the wave will vary with  $x$ , and the lines that appear on the graph (which one calls the *characteristics* in mathematical language) will be curves and no longer lines. Nevertheless, the definitions of the notions of points on the wave, outside the wave, and inside of the wave will extend to those new condition immediately.

Can one state proposition C for such phenomena? The telegraph equation, like the most general equation that corresponds to a heterogeneous cable, belongs to the type that is called the Laplace equation. The general method that is appropriate to the integration of equations of that type was given before even the work of Kirchhoff in a celebrated paper by Riemann. The formulas to which

the Riemann method leads, just like that of Kirchhoff, can be interpreted easily in such a manner as to deduce our property C for all of the corresponding phenomena.

Nonetheless, that interpretation was not given explicitly by either Riemann or even by those (I mean Du Bois Reymond and Darboux) for whom the Riemann method took on its definitive form but who did not have that aspect of the question in mind. The first work to treat an equation other than that of spherical waves from the standpoint of Huygens's principle was the fundamental paper of Volterra (*Acta mathematica*, t. 18) that was dedicated to the *cylindrical wave equation*:

$$(e) \quad \frac{1}{c^2} \frac{\partial^2 u}{\partial t^2} = \frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} + \frac{\partial^2 u}{\partial z^2},$$

which applies to only small motions of a two-dimensional medium, so a medium that reduces to a plane. Here again, we can trace out figures in space-time by considering time to be a third coordinate that is carried perpendicularly to the plane of the given figure. The wave that starts from a unique point  $(x_0, y_0)$  at a well-defined instant  $t_0$  propagates circularly. At an arbitrary later instant  $t_1$ , its front will have the form of a circle that is centered in the original perturbation and has a radius that is proportional to  $(t_1 - t_0)$ . The surface that is generated in space-time by those circles when  $t_1$  varies is a cone of revolution  $o(x_0, y_0, z_0)$  and an axis that is parallel to the  $t$ -axis, or rather the “forward sheet” of that cone (i.e., the sheet that points to positive  $t$ ). The point  $A(x_1, y_1, z_1)$  is *on the wave* with respect to  $O$  if it is located on that conical sheet, i.e., if one has:

$$d^2 = c^2 (t_1 - t_0)^2$$

upon denote the distance between the two points  $(x_0, y_0)$  and  $(x_1, y_1)$  on the plane by  $d$ . On the contrary, it will be *outside of the wave* with respect to  $O$  if one has:

$$d^2 > c^2 (t_1 - t_0)^2 \quad (\text{point exterior to the conical sheet})$$

and *inside of the wave* with respect to it if one has:

$$d^2 < c^2 (t_1 - t_0)^2 \quad (\text{point interior to the conical sheet}).$$

One should note that those conditions can just as well be interpreted by discussing the position of the point  $O$  with respect to that of the point  $A$  that was considered to begin with. The two points will be on the wave, outside the wave, or inside of the wave with respect to each other according to whether  $O$  belongs to, is exterior to, or interior to a conical sheet of revolution that has  $A$  for its summit, respectively, which is a past sheet this time, since it points towards negative  $t$ , which is representative of a wave that issues from  $A$ , but continues to ascend in the course of time. In the theory that we are concerned with, as in that of relativity, two points that are outside of the wave with respect to each other are considered to not exist with respect to each other. There is never any reason to write down a relation in which both of them appear.

It is clear that the notions that were just defined and the reciprocity between them extend to the spherical wave equation that we began with. In order to show that, it is only necessary to read about the incursions into four-dimensional space that the study of the cylindrical wave equation permits us to make.

Finally, the study of anisotropic or heterogeneous media will introduce some other linear partial differential equations with coefficients that are variable, in general, and for which the waves will no longer be spherical or circular, in such a way that their progression will no longer take place along cones with rectilinear generators. That will not interfere with the existence of the preceding distinction between points on the wave, outside of the wave, and inside of the wave, or the reciprocity property that we have confirmed in the two cases that we examined.

Nonetheless, we recall that the notions that we just defined relate to points of space-time: In order to give meaning to them, we must not just give two points in space, but say what instants that were are considering.

Inspired by both the method of Kirchhoff and that of Riemann, Volterra could integrate the cylindrical wave equation, and the formula to which he arrived exhibited the property C for that equation.

Thanks to the methods that he created, the same results have since then been extended considerably to some more general cases by integrating the corresponding partial differential equations. I shall not retrace the path of those integration methods here. I shall say only that that they were inspired by the one that is followed in the study of harmonic functions.

Now, one knows that in the case of three variables, that study is based essentially upon the use of the elementary potential  $1 / r$ , in which  $r$  denotes the distance between two points, and in the case of two variables, upon the use of the elementary logarithmic potential  $\log r$ . In order to integrate the partial differential equations that were spoken of in the foregoing, one must likewise introduce certain auxiliary quantities that are functions of two points (like the elementary potential). For example, the application of the Riemann method demands that one must construct a function of that type for each partial differential equation. Moreover, one can direct one's calculations in such a manner that the quantity thus-introduced is the exact analogue of the elementary potential. That quantity, or *elementary solution*, can become infinite (like the elementary potential): That will happen whenever the two points that it depends upon are *on the wave*. If the number  $m$  of independent variables is odd (as in the case of the cylindrical wave equation) then its form will be completely analogous to  $1 / r$ , namely:

$$v = \frac{V}{\Gamma^p},$$

in which  $\Gamma = 0$  is the condition for the two points to be on the wave and in which  $p = (m - 2) / 2$ . If  $m$  is even then there will generally be one more complication. It can happen that the elementary solution again has the preceding form. However, most often one must add a logarithmic term to the fractional term that was just written down in such a way that the elementary solution will have the form:

$$v = \frac{V}{\Gamma^p} + \mathcal{V} \log \Gamma,$$

in which  $V$  and  $\mathcal{V}$  are always finite, regular functions.

By a convenient use of the elementary solution thus-defined, one will arrive at an extension of any second-order linear partial differential equation of the type (viz., hyperbolic normal type) that one is likely to encounter in physical applications and which includes, in particular, the equations of spherical and cylindrical waves and the calculations of Riemann, Kirchhoff, and Volterra. As a consequence, Huygens's principle, in its form C, is also a general property that belongs to all of the partial differential equations with the character of possessing an elementary solution, and one has the same means for explicitly constructing the distribution of the active centers on the surfaces  $\sigma$  that replace an arbitrarily-given perturbation that is created inside of those surfaces.

We are then committed to the form C of the principle. I have been obliged to insist a bit further that up to now this has not been strictly necessary in the method that was followed in order to prepare us for what we shall now say about the other two propositions A and B.

### III.

Proposition A must be considered to be immediately obvious (and it is the only one of the three that is). It is not distinct from our very principle of scientific determinism. Indeed, that principle expresses the idea that if one knows that state of the world at a well-defined instant  $t_0$  then one must be able to deduce the state of the world at an arbitrary later instant  $t_0 + h$ , where  $h$  is no particular positive time.

If one knows the state that relates to  $t_0$  then one can also deduce the one that relates to the instant  $t_0 + h + k$ . However, it must also be possible to calculate that same state that was just spoken of (when  $k$  is positive) with the aid of the state at the instant  $t_0 + h$ , which is itself supposed to be calculatable when one starts from the state at  $t_0$ . *The two methods of calculation must lead to the same result* with no contradiction: That is precisely what constitutes our proposition A.

A is therefore a sort of truism, but that is no reason for that proposition to not attract the attention of the mathematician. Indeed, it is clear that the calculation by which one passes from the state at  $t_0$  to the state at  $t_0 + h$  constitutes a transformation that depends upon the parameter  $h$ . The proposition A expresses the idea that the set of all those transformations when  $h$  takes on all possible positive values (and one can even add the negative values, at least if one assumes that there is reversibility) constitutes a *group*. The transformation of the parameter  $h + k$  coincides with the product of the two transformations with the parameters  $h$  and  $k$ , respectively.

That fact was discovered for the first time in 1895 by Picard for the case in which the phenomenon is regulated by simply an ordinary differential equation. Somewhat later, in 1903, it was discovered by Le Roux for the general case in which a partial differential equation appears. In the first case, one will obtain an ordinary Lie group. In the second one, the group that is obtained will have a new character because it is composed of *functional* operations.

We can now profit from the fact that thanks to what we know about the integration of partial differential equations, we are informed about the nature of the operations in question. From the analytical standpoint, calculating the ultimate state of a motion when we know the initial state at a well-defined instant  $t_0$  is a Cauchy problem and is precisely the one that Poisson posed for the spherical wave equation (up to the choice of partial differential equation). We saw that such a

problem can be solved by quadratures once we have constructed what we called the elementary solution. The “major premise of Huygens” that I call our proposition A then expresses a property of that elementary solution.

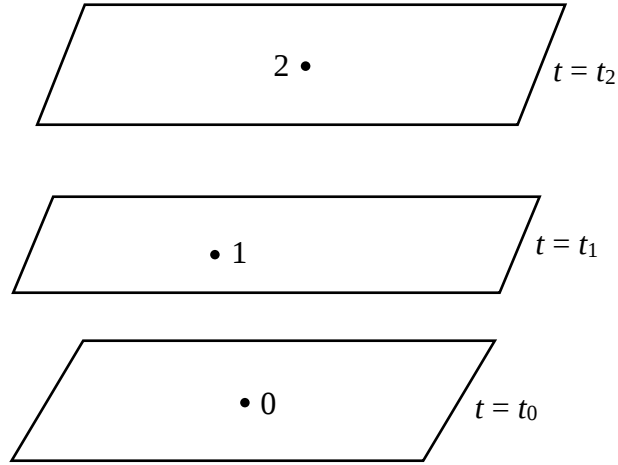


Figure 2.

If one is dealing with the passage from the state at  $t_0$  to the state at  $t_2 = t_0 + h + k$  then one must find the values of that solutions for two points 0 and 2, one of which is taken to be in the plane  $t = t_0$ , while the other is taken from the plane  $t = t_2$  (Fig. 2), and I shall call that value  $v_{02}$ . Similarly, the two partial operations with parameters  $h$  and  $k$ , respectively, introduce, in one case, the elementary solution  $v_{01}$  that corresponds to the point 0 and a point 1 that is taken from the intermediate plane  $t = t_1 = t_0 + h$ . In the second case, the elementary solution  $v_{12}$  is constructed from the same point 1 and the point 2. The “major premise of Huygens” translates into a relation between those three expressions  $v_{01}$ ,  $v_{12}$ ,  $v_{02}$ . It is a true *addition theorem* that the quantity  $v$  must satisfy. However, it is an *integral addition theorem* because if the points 0 and 2 are given then in order to obtain an expression for  $v_{02}$ , one must vary the point 1 in its plane and integrate over the domain of all possible positions for that point. Any elementary solution to a linear second-order partial differential equation will admit an integral addition theorem of that type.

That is not the only situation in which such integral addition theorems present themselves in analysis. An important chapter in the theory of integral equations that Volterra created, namely, the theory of the composition of kernels, is likewise a source of theorems of that type. It would be interesting to compare and contrast the two paths for calculation that open up in that way. By one or the other, one will arrive at some remarkable identities that are sometimes difficult to establish directly.

One of the applications of that integral addition theorem consists of the analytic continuation of the elementary solution. Indeed, it is often the case that the methods that serve to construct it are valid only within a certain radius, in other words, when the two points upon which it depends are not too far apart, so, for example, when the parameter  $h$  that was recently in question does not exceed a well-defined limit  $H$ . Now, the integral addition theorem permits one to free oneself from such a restriction, since the combination of the operations that correspond to two values of the

parameters, both of which are found between 0 and  $H$ , one can produce the result for a value of  $h$  that is found between  $H$  and  $2H$ . Hence one can likewise pass on to values that are even larger than the latter, and so on.

The possibility of that continuation is a particularly remarkable fact when (as is the case for certain forms of the problem) the characteristic surface along which  $v$  must become infinite admits singularities in its own right, which is what happens for reflected waves that present caustics. Even in the regions where that is true, the preceding method will continue to apply and will permit one to define the required solution.

I have not developed that part of my talk any further, first of all, because that exposition would be very lengthy, such as it is, and then because I have already treated it in the article that the Société published on the occasion of the present date. I will then be content to refer you to a point at which I will soon say a word about it and pass on to the fourth part, which is the study of proposition B.

#### IV.

Proposition B, which I can call “Huygens's minor premise,” is the one that was justifiably the most controversial of the three, and to the extent that we shall examine it, we will see, little by little, how false the idea was of making it concur with the proof of the conclusion C, although it initially seemed natural.

If a perturbation is localized to a unique point A at the origin of time (I intend that to mean inside of an infinitely-small sphere with that point for its center) then it will be constantly true, as we have seen, that the wave that propagates in a medium that is governed by equation (E), which was written previously, will have a sphere  $S_1$  of radius  $ct_1$  at an arbitrary instant  $t_1$  and that there is no effect outside of the aforementioned sphere at that instant, in other words, that no effect exists for the points that are *outside of the wave* with its center at the original perturbation. However, proposition B asserts that it no longer exists *inside* the sphere, i.e., for the points that are *inside the wave* with that perturbation. If that is the case then one will say that the waves that are governed by equation (E) do not *diffuse*.

It is curious that Huygens did not seem to have believed so firmly in that proposition that one assumes in the polemic that was followed. When one refers to its text <sup>(1)</sup>, one will see that in order to evaluate the effect that is produced at the final instant  $t_2$  by the spherical wave that emanates from a unique A at an instant  $t_0 = 0$ , he replaced it, not with the new wave that issued at the instant  $t_1$  from the points of the corresponding sphere  $S_1$ , but with a system of spherical waves that issued from any sort of points inside of  $S_1$ .

It is true that for him those new spherical waves were limited to the surfaces of the corresponding spheres. As for the final effect, he considered it to be non-existent, or at least negligible, anywhere but the surface of the sphere  $S_2$  that constituted the wave front at the instant  $t_2$ :

“Each of those waves [he said] can only be infinitely weak in comparison to the wave [which we call the wave front] to the composition of which all of the others

---

<sup>(1)</sup> Particularly *loc. cit.*, pp. 22.

contribute by way of the part of their surface that is furthest from the center A,” and later on “...The waves... do not compete at the same instant to collectively compose a wave that terminates the motion, which is precisely the circle  $CE$  that is their common tangent,” and even later “...The parts of the particular waves that extend outside of the space  $ACE$  [i.e., the part that lies inside of the sphere in the figure that was considered] are too weak for them to produce light.”

That argument will obviously leave the reader confused, and we do not know very much about whether the portions of the wave that Huygens considered to be negligible are indeed such. Meanwhile, Poincaré was content to say that this argument does not stand up to a rigorous analysis, which seems quite severe to me, and I ask that one should appeal to his judgement. We should be more enlightened today than we were ten years ago so as to find it a bit too radical, being given the phenomena that have become familiar to us in the meantime. The more militarily inclined of us have indeed acquired some notions from the study of artillery that are undoubtedly quite *passé*: One knows that the same projectile whose motion through the air has generally translated into what we hear as a whistling sound will give rise to a bang for an observer at a well-defined instant, i.e., to the sound of an explosion that is analogous to the one that accompanies the sudden departure and is even more violent. The whistling and the bang, although different to our sense, thus have one and the same origin, namely, the sound waves that are born at the contact between the shell and the atmosphere that it cleaves, except that in the first case, they will reach us in isolation, while along their envelope, they will add together and surge forth in a way that does great damage to our eardrums. The bang is to the whistling exactly as the final motion that Huygens retained is to the one that he proposed to neglect, and we see that such an accumulation of waves is capable of giving the impression of a new physical phenomenon, whereas in reality there is only the difference between less and more. That is certainly due to the fact that this difference is infinitely large. In the case of Huygens, as in the case of the bang, which is completely identical to it, the vibratory energy density at the level of the envelope is infinitely large in comparison to what that density is below the envelope.

We shall leave behind the text of Huygens and move on to Fresnel and Poisson, whose discussions had precisely the proposition that we spoke of as their subject. Poisson's objection was quite simple. Assume, if one so desires, that there exists no effect inside of the wave front at the instant  $t_1$  and apply Huygens's argument to the final instant  $t_2 = t_1 + k$ . Assume furthermore, as Huygens did, that the spherical waves that issue from the various points of the sphere  $S_1$  at the instant  $t_1$ , which is the wave front at that instant, have no effect anywhere that they do not accumulate, i.e., everywhere along their envelope. However, that envelope is not composed of only the sphere  $S_2$  of radius  $ct_2 = c(t_1 + k)$ . It also includes a second sphere  $S'_2$  that is concentric to the first one and has a radius of  $c(t_1 - k)$ . All of the reasoning that served to prove the existence of a perceptible effect on one of those spheres seems to apply to the other. There are certainly some points (viz., those of the sphere  $S'_2$ ) that are interior to  $S_2$  and at which some effect of the original perturbation must nonetheless remain at the instant  $t_2$ .

In order to immediately see that this objection, as compelling as it might seem, is not well-founded, we hastily say that it is not at all special to proposition B, which we are presently

examining, and that the same difficulty will arise even if we suppose that this proposition is not true, i.e., if we assume that there can be diffusion.

Let us refer to the considerations that were presented just now in regard to Huygens's major premise as it applies to either the spherical wave equation or any other partial differential equation of an analogous type. I will illustrate matters as if we were dealing with the cylindrical wave equation, although the argument works better when we concentrate on equations with an even number of independent variables, for reasons that I will not insist upon at the moment. We again start from the perturbation that is localized at a point  $O$  at the instant  $t_0$  and whose effect propagates by a wave that progressively expands about that point. Let  $S_1$  be the wave front the instant  $t = t_1$ . There will or will not be any diffusion effects in the region of space that is interior to  $S_1$ , but certainly (and this emerges from the aforementioned calculation that led to the conclusion C) there will be another special effect along the surface  $S_1$  itself. We must then (conforming to principle A) study the repercussions of all of that at the final instant  $t_2 = t_1 + k$ . Now, the waves that are produced at the various points of the surface  $S_1$  at the instant  $t_1$  will have an envelope that is composed of two sheets that will give two distinct wave fronts at the instant  $t_2$  in the general case, as in the one that Poisson had in mind, namely, an external front that is nothing but the original wave from A and an internal front that is analogous or identical to Poisson's sphere  $S'_2$ . Exactly as in Poisson's argument, the calculation will predict, in the first place, a special term that corresponds to that internal front. It is quite obvious that it will appear only there, so the special term will drop out after all of the calculations have been done, since the instant  $t_0$  and  $t_2$  are given, so the intermediate instant  $t = t_1$  can be chosen arbitrarily in such a way that the internal front in question will have all sorts of positions inside of  $S_1$  at will. However, the analytical verification of that fact will be a difficulty of exactly the same order whether one does or does not have diffusion, i.e., whether proposition B is or is not true.

The polemic between Fresnel and Poisson then centered around the issue of whether one would know if the effects at the internal front cancel and how that would be possible. Fresnel sought the explanation for that fact in the undulatory, sinusoidal, character of the perturbation. As we have said, that is not the proper question.

Moreover, you, (like myself) might find it superfluous to study whether the compensation in question is conceivable, since the real question is knowing whether it actually takes place and whether the proof of the second fact would follow obviously from that of the first.

Now the punchline to the story (*le piquant de l'histoire*) is that this proof was given by Poisson himself. One asks how no allusion was made in that controversy (which took place in 1823) to the formula that had been established in 1819 that resolved the issue entirely. It was the one that provided the general integral of the spherical wave equation when one was given the values of the unknown  $u$  and its derivative  $\partial u / \partial t$  as a function of  $x, y, z$  for  $t = 0$ . In order to obtain the numerical value that is taken by  $u$  at the point A ( $x_1, y_1, z_1, t_1$ ) in space-time under those conditions, i.e., the point ( $x_1, y_1, z_1$ ) and the instant  $t_1$ , one knows that one must describe a sphere of radius  $ct_1$  and its center at the point ( $x_1, y_1, z_1$ ) at the instant  $t_1$ . The desired value of  $u_A$  is then obtained with the aid of the double integral that extends over the surface of that sphere and that one can describe explicitly with the aid of just the givens of the problem (i.e., with the aid of what happens for  $t = 0$ ).

The answer to the question that we have posed then becomes obvious. If the medium is originally at rest and the perturbation takes place at the instant  $O$  (or rather at an infinitesimally close prior instant  $t = -\varepsilon$ ) solely in the immediate neighborhood of the point  $O$ , while  $u$  and  $\partial u / \partial t$  will generally be identically zero for  $t = 0$ . It will be non-zero only inside of a very small sphere whose center is that point. In order for the double integrals that enter into Poisson's formula to have elements that are non-zero, and consequently for  $u_A$  to be non-zero, it is necessary that this small sphere must meet Poisson's sphere  $\Sigma$ , i.e., that (up to an infinitesimal) the points  $O$  and  $A$  are *on the wave*. If they are *inside the wave*, i.e., if the small sphere that we just spoke of is interior to the Poisson sphere, then  $u$  will be zero at  $A$ , just as if the two points were outside the wave. In other words, conversely, if we consider the wave that issues from the point  $O (x_1, y_1, z_1)$  at the instant  $0$  and that  $u$  is identically zero at the point  $(x_1, y_1, z_1)$  up to the instant  $t'_1$  when that wave reaches the point in question then it can be non-zero for  $t = t'_1$ , since the points (which are space-time points)  $(x_0, y_0, z_0, 0)$  and  $(x_1, y_1, z_1, t'_1)$  will be on the wave. However, it will become zero again and remain zero at the later instant, since one will then have world-points that are on the wave with the point  $O$  that is the site of the initial perturbation.

That is precisely the stated proposition B that was the source of the dispute between Poisson and Fresnel.

It hardly needs to be said that the same fact will read analogously in the formula that Kirchhoff obtained.

However, now that we have focused upon the question the relates to spherical waves, it would not be appropriate for us to limit our curiosity in that way. We must demand that things happen in the same way when the partial differential equation of the problem is nothing but our original equation (E).

It is clear that an initial response will emerge from an examination of the equations that were treated first, namely, those of Riemann and those of Volterra.

The answer is negative. This time, the question is not being presented from solely the theoretical viewpoint: It was raised in a manner that was most precise and most necessary for a practical application, namely, telegraph transmissions. That residual effect inside of the wave front that disappears completely in the cases of Poisson and Kirchhoff is then produced, and the operating faults that it occasions along the lines were initially a great source of frustration to the telegraph engineers. Poincaré and Picard each provided an interpretation in their own ways. That of Picard is particularly brilliant. As classical as it has become today, we shall briefly recall it.

If an electric cable is assumed to be perfectly homogeneous in all of its parts then we have said that the potential  $u$ , when considered as a function of the abscissa  $x$  and time  $t$ , must satisfy the "telegraph" equation:

$$A \frac{\partial^2 u}{\partial t^2} + 2B \frac{\partial u}{\partial t} = C \frac{\partial^2 u}{\partial x^2} .$$

That equation is suggested by the Riemann method and the Riemann function could be easily constructed by Picard. Then suppose that a perturbation is given at  $t = 0$  that is localized to a very small segment  $(0, \alpha)$  of the cable, i.e., that one demands that  $u$  and  $\partial u / \partial t$  must be zero for  $t = 0$  outside of the segment in question and that one is given their values inside of that segment more

or less arbitrarily. One then calculates, with the aid of those givens, the value of  $u$  for a given (positive) instant  $t = t_1$  and a given abscissa  $x = x_1$ , which, to fix ideas, we likewise suppose to be positive and even noticeably greater than  $\alpha$ , since that would amount to studying a signal that is received at a great distance. If, with Picard, we apply the Riemann method, purely and simply, then we will be led to draw the two characteristics:

$$x + \omega t = \text{const.}, \quad x - \omega t = \text{const.}$$

through the point  $A(x_1, t_1)$  by letting  $\omega$  denote the speed of propagation of the electrical oscillations under the conditions of the present problem. If  $M$  and  $N$  are the two points where those two characteristics cut the  $x$ -axis respectively then the Riemann method will provide the desired value  $u_A$  as a sum of three terms, one of which comes from  $u_M$  and the second of which comes from  $u_N$ , while the third one is a definite integral over the segment  $MN$ . The first of those terms is certainly zero, since the point  $M$  is obviously external to just the segment  $(0, \alpha)$  of the  $x$ -axis where the givens are non-zero.

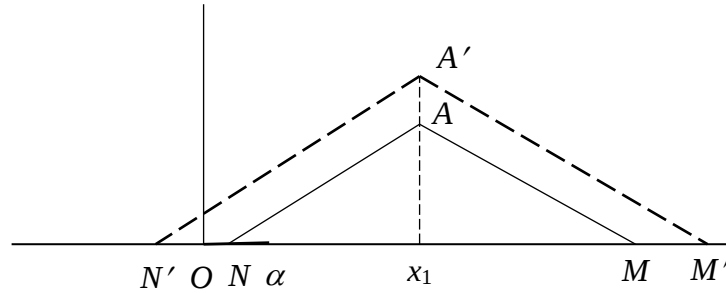


Figure 3.

On the contrary, the second one can be non-zero. If we consider a well-defined point of the cable, i.e., we suppose that  $x_1$  is given once and for all, then we give positive and increasing values to  $t_1$  such that second term will become non-zero at a well-defined moment  $t'_1$ , or more precisely (if  $t'_1$  has the significance that was indicated before)  $t'_1 - \alpha / \omega$ , which is the one when the corresponding point will be *on the wave* with the given perturbation. It is nothing but the reception of the telegraph signal, and the moment of that reception is indeed that of its emission, shifted by the time  $x_1 / \omega$  that is necessary for its propagation.

However, when one starts from that moment and one then continues to increase  $t_1$ , the third term, which is the integral term, will also begin to contribute. Soon (after a very small time  $\alpha / \omega$ ), that third term be the only one left, since the point  $(x_1, t_1)$  will then cease to be on the wave with the original signal, i.e., the reception of that signal will cease. Now at that moment, the third term (i.e., the diffusion term), which no longer involves points on the wave with the point  $A'$ , but the points *inside that wave*, will remain non-zero. That constitutes a *harmful effect* that perturbs the following transmissions to the telegraphers.

The experimentally-confirmed presence of that harmful effect is therefore quite natural: One can even calculate its value as a function of the givens and consequently study the mode of emission of the signal that is most appropriate for diminishing its importance.

An examination of the results that were obtained by Volterra will lead to some analogous conclusions. Having to integrate the cylindrical wave equation in a three-dimensional (and no longer four) region of space-time is entirely similar to the problem that Kirchhoff considered and Volterra also arrived at (up to singularities, for which there is no reason to discuss at the moment), which is that of representing the phenomena by the action of fictitious centers that are distributed over curves that are exterior to the one that is given. However, this time, those centers will not only act when they are on the wave with the point  $A$ , they will also act when they are inside that wave. There again, as in the previous example, a perturbation that is localized in the immediate neighborhood of a point  $O$  in space-time (I once more mean only three-dimensional space-time) will make its action at an arbitrary point felt, not only at the precise moment when the corresponding wave arrives, but indefinitely after that moment. The wave leaves a residual effect, and one of my colleagues has said that it is indeed fortunate that the space in which we live has three dimensions, because if it had only two, and consequently the propagation of sound in it were governed by the cylindrical wave equation, then we would have to perceive all of the sounds that have been produced since the creation of the world (but very muffled, to be sure).

With that, we are then led to the true judgement that we must give to Huygens's minor premise. Far from being able to serve as the foundation for proposition C, it is hardly true by the same right as the latter. It is neither immediately obvious, like A, nor generally valid, like C, for all of the second-order linear partial differential equations, and consequently, for all corresponding phenomena. It is a particular property of the spherical wave equation.

However, as a result of that confirmation, we see that we can pose a whole series of problems. In the first place (and this is even obvious), we must ask what the partial differential equations would be that satisfy Huygens's minor premise, i.e., the ones for which the propagation of waves will take place without diffusion or residual effect.

If all that one wishes to do is find the necessary and sufficient conditions for the absence of diffusion then the answer will be very simple.

One must first exclude all of the equations with an odd number of independent variables (such as the cylindrical wave equation); all of them, without exception, will give rise to diffusion. The same thing will be true when the number  $n$  of independent variables is equal to 2.

Now take the number  $m$  to be even and equal to at least 4. However, it will not follow that there is no diffusion. For example, consider *the damped spherical wave equation*:

$$\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} + \frac{\partial^2 u}{\partial z^2} + K u = 0 .$$

We propose to find the solution  $u$  such that for  $t = 0$ ,  $u$  will be equal to a given function  $g$  of  $x$ ,  $y$ ,  $z$ , and  $\partial u / \partial t$  will be equal to another given function  $h$  of the same variables. One performs the same construction that Poisson did for equation (E), i.e., if  $(x_1, y_1, z_1, t_1)$  is the point of space-time where one wishes to calculate the numerical value of  $u$  then one describes the sphere of radius  $ct_1$  with its center at the point  $(x_1, y_1, z_1)$ . The sphere  $\Sigma$ , thus-described, will again serve as the domain of integration for the quadratures that bear upon the givens  $g$  and  $h$ . However, this time, those quadratures will be of two types: One of them will be double integrals over the surface of the

sphere, as in the first case. The other will be *triple* integrals that extend over the entire *interior* of that sphere, i.e., ones that are *inside the wave* with  $(x_1, y_1, z_1, t_1)$ . That is due to the fact that for the new equation that we consider at the moment, the elementary solution will contain a term in  $\log \Gamma$ , if  $\Gamma = 0$  denotes the equation of the sphere, rather than the one in equation (E) that reduces to  $1/\Gamma$ .

For an entirely-arbitrary perturbation, *the necessary and sufficient condition for the absence of diffusion is that if  $m$  is even then the elementary solution must contain no logarithmic term.*

That answers the question up to a point. However, it is one thing to indicate a necessary condition, as we just did, and something else entirely to construct all of the equations that satisfy that condition. We know one of them – namely, equation (E) – and the same fact will be true for the analogous equations in 6, 8, ... variables. It will similarly be obvious for all of the ones that we can deduce by some obvious transformations, such as changing the independent variables and multiplying the unknown or the left-hand side of the equation by a factor. Such equations must not be considered to be essentially distinct from the first ones.

Do there exist other cases than those for which our minor premise B is verified (which are certainly all exceptional cases)? That question has not yet been answered. In fact, it is the foregoing that presents the difficulty, moreover, since if a well-defined equation answers the question then the analysis leads to its discovery must, at the same time, provide all of the ones that are derived from it any of the aforementioned transformations.

As such, that question is coupled with the study of point-like transformations of the partial differential equations and the corresponding invariants, i.e., the work of Cotton and also Levi-Civita. At the same time, that will involve, and probably shine a new light upon, a great classical theory of geometry, namely, that of geodesic lines.

However, some other viewpoints demand our further attention.

Return to the Riemann case, namely, that of  $m = 2$ . The fact that there is always diffusion translates into the existence of a *residual integral* that is generally non-zero in the region that is *inside of the wave* and no longer on the wave with the initial perturbation. For the simplest case, namely, that of the vibrating string, or what amounts to the same thing, the equation:

$$\frac{\partial^2 u}{\partial x \partial y} = 0,$$

that residual integral will reduce to a constant (which can obviously be zero, and whose value depends upon the form of the generating perturbation). More generally, the classical studies of Darboux (*Leçons sur la Théorie des Surfaces*, t. II) lead one to ask whether there do not exist equations that belong to the indicated category whose residual integral depends upon only a finite number of arbitrary constants, or for which that integral will necessarily verify a linear partial differential equation that is distinct from the proposed one.

The answer is simple: In order for that to be true, *it is necessary and sufficient that the proposed equation should be integrable by the Laplace method*, i.e., that the Laplace sequence should terminate in at least one direction. That is a consequence of a theorem by which Goursat had

learned to characterize the Laplace equation. Like that theorem, the present consequence characterizes those equations by an intrinsic property of their solutions and not by a well-defined method that we can only imagine in order to obtain them.

Except for the equations that are integrable by the Laplace method, the residual integral does not present itself in a form that is also specialized. Is it as general as the general integral itself? In other words, can an arbitrary solution of the equation be considered to be a residual integral?

Certainly not. First of all, observe that in practice the two elements of the general integral, namely, the residual integral and the one that corresponds, on the contrary, to the passage of the wave front, are initially distinguished by their orders of magnitude. Do not forget that this residual integral is precisely the one that Huygens proposed to neglect as infinitely weak, and we are once more led to see in that an idea that is not unreasonable. If we recall, for example, the study of the telegraph transmission then we will see that the harmful effect, which is measured by a definite integral that extends over the segment  $(0, \alpha)$ , which is the seat of the initial perturbation, has the same order as  $\alpha$ , whereas the signal, properly speaking, will not contain that quantity as a factor. One can then render the harmful effect as small as one desires in comparison to the useful effect if one can truncate the duration of the emission without limit. A similar fact presents itself in all of the other examples. To fix ideas, take that of the damped spherical wave equation and suppose that the initial perturbation is localized in a small sphere of radius  $\alpha$ . If that sphere is traversed by the Poisson sphere then we will have a wave front term that will be a double integral, and as such, will have order  $\alpha^2$ . On the contrary if the small sphere is entirely inside the wave, i.e., the interior of the Poisson sphere, then we will have exclusively the residual integral, whose quality of being a triple integral will have order  $\alpha^3$ , and consequently, if  $\alpha$  is very small then it will be much weaker than the previous one. That is basically the idea that Huygens expressed, but in a different form.

Independently of that first criterion (which will not come into play if one does not know the order of magnitude of the initial perturbation or that of the wave front term), we know, at the very least, an essential primary character that necessarily belongs to the residual integral, but not the general integral. We know that equations of the elliptic type with analytic coefficients (the potential equation, for example) have integrals that are all analytic. However, that is not the case for the hyperbolic equations (i.e., ones that are compatible with the existence of waves), which always admit some non-analytic integrals.

Now, as opposed to the general integral, *the residual integral is always analytic* (of course, when that is true for the coefficients of the equation). The classical arguments that are given in regard to that in potential theory apply without modification, as far as that is concerned.

It remains for us to know the conditions under which an analytic solution to the equations can be considered to be a residual integral. Always for  $m = 2$ , an examination of the classical formulas that led to the Riemann formula show that the question reduces to the study of an integral equation of the first kind. Moreover, it will oblige us to consider such equation in a somewhat different light from the ones that have been envisioned most frequently up to now.

The theory of integral equations of the first kind is, in a sense, elucidated by some work that is known to all, namely, that of Picard, which was the work that the theory gave birth to. By a fortuitous use of the knowledge that we have acquired about the integral equations of the second

kind and the theory of fundamental functions that it is connected with, we will get a system of necessary and sufficient conditions for the solubility of the equation.

Now, despite the fecundity of the viewpoint that has been adopted that was established by the results obtained, one can easily convince oneself that there is good reason for one to not limit oneself to it, and that the close connection that is established between the two types of integral equations (as advantageous as it might be, for the moment, that it permits us to utilize the important body of knowledge in an unexplored domain that was acquired in a neighboring domain), it is not without its inconveniences, moreover. The two domains are profoundly different. On a first inspection of the Fredholm equation, one will perceive that it is essential to make the principal variable  $x$  (which is one that appears outside of the  $\int$  sign) traverse exactly the same interval that the second variable  $y$  (viz., the integration variable) must traverse under the  $\int$  sign, since those two quantities enter into it as arguments of the same function.

Nothing similar is true for an equation of the first type (at least if the kernel has no singularities and the limits are fixed). The two intervals, namely, the integration interval and the one in which one considers the principal variable  $x$ , have no relationship to each other. One can suppose, at will, that they are identical or distinct. In order to do that, it will suffice to change the variable, which one can do arbitrarily with  $x$ , as well as with  $y$ . In the general theory, there will be good reasons to consider those intervals to be different.

It should attract our attention all the more that, by contrast, the nature of the problem will change profoundly with the relation that exists between the two intervals in question. If the solution to the problem is possible and determinate – for example for a certain choice of integration interval  $(\alpha, \beta)$  and the interval  $(a, b)$  that is traversed by  $x$  – then it must become impossible when one diminishes the first integral or augments the second one, and become indeterminate if one does the opposite, in such a way that for given values of  $\alpha, \beta, a$  can correspond to a perfectly-determinate value of  $b$  for which the problem is well-posed.

Some authors (above all Bateman, to my knowledge) have already considered equations of the first type from that viewpoint. That is the one that imposes itself upon us in the application that we are addressing today. In the latter, the two intervals that we just spoke of are, in principle, different and have no relationship to each other.

Now that we have asked what the common properties of all the residual integrals of such an equation might be and what the initial perturbation that gave rise to them might be, we return to the original question. We can study how we can choose the initial perturbation, and if we can choose it in such a way that the residual integral will be annulled when that it not true for an arbitrary initial perturbation.

That is once more a question of integral equations of the first kind, but this time it is the homogeneous equation that we must solve. This new problem calls for the same observations as before, moreover.

One remarks that it is the one that is closest to practice. It is the one that suggests the telegraphic application that was recently of interest. The most classical discovery that was obtained along that path, namely, the sinusoidal mode of emission that was invented by Lord Kelvin for the submarine telegraph, obviously has no relation to the explanation that Fresnel attempted to give for the absence of the internal front and diffusion in spherical waves, and it is based upon the undulatory

character of the initial perturbation. Moreover, that relation (for which there are undoubtedly good reasons to examine the most closely) shows us once more that there is always something to be retained from the intuitions of a Huygens or a Fresnel, even when we are not led to accept such things.

All of those problems are likewise posed for  $m$  greater than 2. One notes that the equations whose residual integrals depend upon a finite number of arbitrary constants, or even satisfy a partial differential equation that is distinct from the proposed one, constitute a special category that is analogous to the equations in two variables for which the Laplace sequence terminates, even if that situation does not give rise to a particular method of integration.

We would have come to an end with that (if not with all of the essential aspects of the question then at least with all of the ones that presented themselves on first glance) if the case of solid obstacles did not open a new chapter of research that we must now say a word about. In the book *Hydrodynamique, Élasticité, Acoustique* by Duhem, one finds an example that seems to contradict our minor premise of Huygens that nevertheless bears upon the spherical wave equation. Duhem considered a pulsating sphere whose fixed center was placed, for example at the coordinate origin and whose radius submitted to infinitely-small oscillations during a certain time interval  $\theta$ : He then studied the small motions that were generated by those oscillation in the homogeneous atmospheric medium in which the sphere was supposed to be immersed. Now, if the sphere definitively returns to a state of rest at the end of the time interval  $\theta$  (and its radius returns to its initial value, as well) then that will not generally be true for the medium of the environment, even after the spherical waves that were created in the course of the pulsations have ceased to pass through it.

That apparent contradiction will dissipate when one takes into account the fact that if solid obstacles are present, and consequently regions in which the partial differential equation ceases to be verified, then the problem in analysis into which one correctly translates the physical question of determining the motion will no longer be the same as it was just now, which corresponds to the fact that one must include not only the direct waves, but also the reflected waves. Moreover, this new problem raises some analytic considerations that are very analogous to the preceding ones. One will have to introduce a new quantity that is analogous to the one that we called the elementary solution (but will become infinite on the reflected wave and not on the directed wave), and according to whether that quantity does or does not contain a logarithmic term, there will or will not be diffusion of the reflected wave, respectively. The problems will then be the same as they were for the direct wave, but with the influence of a new element, namely, the form of the reflecting wave.

But let us not lengthen the list of all of those problems, which is already quite extensive. If a great number of them have been posed then we now see that this is due to the very special character of Huygens's minor premise, while the major premise A is one of those "guiding principles of knowledge," according to the terminology of the philosophers, beyond which we do not know how to think and reason, and the conclusion C is a property that is common to all classes of phenomena that we have dealt with today, while the minor premise B, which is much more esoteric,

differentiates those phenomena from each other. It might or might not be verified according to which of the two that one considers in particular.

That is why it is still mysterious to this day and remains part of the science of tomorrow.

---