

Mapping of a crack edge by ultrasonic methods

A. N. Norris and J. D. Achenbach

The Technological Institute, Northwestern University, Evanston, Illinois 60201

(Received 15 November 1981; accepted for publication 15 February 1982)

Two methods are proposed for the mapping of cracklike flaws in homogeneous, isotropic, elastic media. The methods require as input data the travel times of diffracted ultrasonic signals. The first method maps points on the crack edge by a process of triangulation with the source and receiver as given corner points of the triangle. By the use of travel times for neighboring positions of the source and/or the receiver, the direction of signal propagation, which is the necessary constituent required to complete the triangle, can be computed. The inverse mapping is global in the sense that no *a priori* knowledge of the location of the crack is required. The second method is a local edge mapping which determines sets of planes relative to a known point close to the crack edge. Each plane contains a flash point. The intersection of the envelopes of two sets of planes maps an approximation to the crack edge.

PACS numbers: 43.20.Fn, 43.35.Zc, 43.20.Bi, 43.20.Dk

INTRODUCTION

The field generated by scattering of ultrasonic waves by a flaw contains a substantial amount of information on the flaw's location, its size and shape. The extraction of this information requires the solution to an inverse problem.

Inverse scattering theories may be divided into two general categories, which use data in the time and frequency domains, respectively. If the inhomogeneity is expected to be a smooth convex cavity, then the physical optics inverse scattering theory (see Bojarski¹ and Lewis²) may be a suitable approach. This is a frequency-domain theory based on a Kirchhoff approximation to the solution of the direct problem. The input data is the observed backscattered field. Bojarski showed that the backscattered field is directly related to the characteristic function of the cavity. This function is defined as unity inside the cavity and zero outside. Boerner³ has noted the connection between this method and the use of Radon transforms or projection mapping methods.

A variation of the Bojarski theory in which the function to be mapped is singular on the surface of the scatterer and zero elsewhere was proposed by Cohen and Bleistein.⁴ Such a theory is suitable for inverting crack-scattered data, since the volume of a crack is zero but its surface area is not. Hence the characteristic function has zero support, while the singular surface function is singular over the entire crack face. If the crack is flat, then it is completely defined by the crack edge. This fact was utilized by Achenbach *et al.*,⁵ to develop an inversion scheme which maps the crack edge. Also in the area of crack characterization, Teitel⁶ has shown how the low-frequency scattered field can be used to determine all the relevant parameters for a flat elliptical crack. However, Teitel's low-frequency method requires the measurement of the exact amplitude of the scattered wave. In practice the measured amplitude may, however, be quite different from the theoretical one, due to effects of coupling of the transducer to the material and dissipation within the material.

The methods proposed in the present paper are based on arrival times. They do not require knowledge of the absolute amplitude of the scattered signal. In homogeneous elastic materials the travel time of a signal from one transducer to another can be measured quite accurately. The first received scattered signal satisfies Fermat's Principle of least time. It is well known that such a signal describes a geodesic curve or stationary ray path. The point on the scatterer from which the scattered signal emanates is called the flash point. If the shape and location of the scatterer are known, then the position of the flash point can be computed by the laws of geometrical optics. In general, for every stationary ray path there is an associated flash point on the scatterer. The methods of this paper use travel times to map the locations of the flash points.

For the host material we consider a homogeneous, isotropic, linearly elastic solid. For materials which display significant inhomogeneity and/or anisotropy, and hence significant wave velocity variations, the present methods cannot be expected to yield accurate results. In principle such effects can be taken into account, but the necessary computations become rather difficult. The flaw is assumed to be a crack with a well-defined edge. Except in the domain of specular reflection, the first signals arriving via the crack are produced by diffraction at flash points on the edge of the crack. For the direct problem the positions of the flash points follow from Snell's law of edge diffraction, which is defined within the context of the geometrical theory of diffraction for elastodynamics, as discussed by Achenbach and Gautesen.⁷ In the inverse problem flash points on the crack edge are determined from the arrival times of observed signals.

Two methods are proposed in this paper. In the first method, which is a global triangulation method, the source and the receiver are given corner points of the triangle, and we compute the flash point as the third corner point. The triangle may be completed if the direction of the stationary ray path is known at either the source or the receiver. This

direction can be computed by measuring the spatial gradient of the travel time. The crack edge may be mapped by locating a sufficient number of flash points.

The second method is a local mapping technique. It is local in the sense that a base point near the crack must be known *a priori*. This base point can be determined by the global triangulation method. The travel times corresponding to several different source-receiver pairs form surfaces on whose intersection the flash points must be located. To a first approximation, the surfaces can be replaced by planes, whose intersection is easily computed. By iteration, the solution converges to a section of the crack edge. Numerical tests of the method have been carried out by the use of synthetic data.

It is assumed that the travel times for all relevant stationary ray paths can be measured. This may not be possible for some diffracted signals, whose arrival times may be too close to that of preceding signals. In this case the relevant arrival time may be inferred from measurement of the spacing of peaks in the high-frequency interference spectrum. For more details, the reader is referred to Achenbach and Norris.⁸ Finally, we note that throughout the paper the source and the receiver have been assumed separated, i.e., the measurement method is pitch-catch. The analysis can easily be modified to accommodate pulse-echo data.

I. GLOBAL TRIANGULATION

If an observed signal is known to emanate from a flash point, then the inverse problem of mapping the scatterer may be viewed as a problem of triangulation, i.e., completing the triangle with corner points at the source, the receiver, and the flash point. The positions of the source and the receiver and the time delay between emission and reception of the diffracted signal are known. In the following it is shown that knowledge of the signal propagation direction at either the source or the receiver is generally sufficient to complete the triangle.

A. Parametric dependence of the flash point

In a homogeneous, isotropic, linearly elastic medium there are two wave speeds c_L and c_T corresponding to longitudinal and transverse waves, respectively. If ρ is the density and λ, μ are the Lamé elastic constants, then

$$c_L = [(\lambda + 2\mu)/\rho]^{1/2}, \quad c_T = (\mu/\rho)^{1/2}. \quad (1)$$

The slownesses s_L and s_T are the inverses of the wave speeds.

At the receiver, four different signals may be observed, corresponding to the different types of the emitted and received waves. Let T_0 be the time delay for a signal of emitted type α , $\alpha = L, T$ and received type β , $\beta = L, T$. If $\mathbf{x}_S$ and $\mathbf{x}_Q$ represent the position vectors of the source and receiver, respectively, then the surface $E_{\alpha\beta}$ on which the flash point lies may be described as

$$E_{\alpha\beta} = \{\mathbf{x}_{\alpha\beta} s_\alpha |\mathbf{x}_{\alpha\beta} - \mathbf{x}_S| + s_\beta |\mathbf{x}_{\alpha\beta} - \mathbf{x}_Q| = T_0\}. \quad (2)$$

For any point $\mathbf{x}_{\alpha\beta}$ on $E_{\alpha\beta}$, we define the unit vectors $\mathbf{p}$ and $\mathbf{q}$ and the distances r_S and r_Q as

$$\mathbf{x}_{\alpha\beta} = \mathbf{x}_S + r_S \mathbf{p} = \mathbf{x}_Q + r_Q \mathbf{q}. \quad (3)$$

Also, we define the vector $\mathbf{X}$ from the source to the receiver as

$$\mathbf{X} = \mathbf{x}_Q - \mathbf{x}_S. \quad (4)$$

Thus the unit vector from the source in the direction of the receiver is $\hat{\mathbf{X}} = \mathbf{X}/X$, where $X = |\mathbf{X}|$. The dependence of the distances r_S and r_Q on $\mathbf{p}$ and $\mathbf{q}$ follows from the definition of the surface $E_{\alpha\beta}$. We have, for example, that

$$r_S = f_{\alpha\beta}(\mathbf{p} \cdot \hat{\mathbf{X}}), \quad (5)$$

where

$$f_{\alpha\beta}(\xi) = \begin{cases} (s_\alpha^{-2} T_0^2 - X^2)/2(s_\alpha^{-1} T_0 - X\xi), & \alpha = \beta, \\ \{s_\beta^2 X\xi - s_\alpha T_0 \pm [(s_\beta^2 X\xi - s_\alpha T_0)^2 + (s_\beta^2 - s_\alpha^2) \times (T_0^2 - s_\alpha^2 X^2)]^{1/2}\} / (s_\beta^2 - s_\alpha^2), & \alpha \neq \beta. \end{cases} \quad (6)$$

The ambiguity with the $\pm$ in the definition of $f_{\alpha\beta}(\xi)$, $\alpha \neq \beta$, is taken care of as follows: If $T_0 > s_T X$, then only the plus sign is taken. However, if $\alpha = L$ and $T_0 < s_T X$, the distance r_S becomes a double valued function of the unit vector $\mathbf{p}$. In other words the source is not contained within E_{LT} if $T_0 < s_T X$. When this happens, either sign is admissible for $f_{LT}(\xi)$.

The distance r_Q from the receiver to the flash point is

$$r_Q = f_{\beta\alpha}(-\mathbf{q} \cdot \hat{\mathbf{X}}). \quad (7)$$

When $\alpha = \beta$, the surface $E_{\alpha\beta}$ is simply a spheroid with foci at the source and receiver. The major axis length is $s_\alpha^{-1} T_0$ and the eccentricity is $s_\alpha X / T_0$. Some examples of the surface $E_{\alpha\beta}$, $\alpha \neq \beta$, are presented in Fig. 1.

We note that $\mathbf{p}$ and $\mathbf{q}$ are not independent but are related by the identity

$$r_S \mathbf{p} - r_Q \mathbf{q} = \mathbf{X}. \quad (8)$$

The explicit dependence is

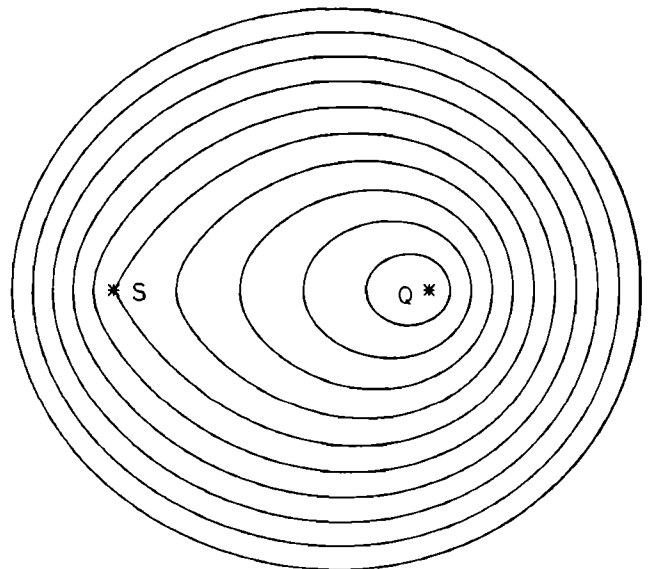


FIG. 1. Sections of the surfaces E_{LT} in a plane containing the source S and receiver Q . Dimensionless travel times are $T_0 c_L / X = 1.2(0.2)3$. Poisson's ratio is $1/3$.

$$\mathbf{q} = g_{\alpha\beta}(\mathbf{p}, \hat{\mathbf{X}}) [f_{\alpha\beta}(\mathbf{p}, \hat{\mathbf{X}}) \mathbf{p} \cdot \hat{\mathbf{X}} - X]^{-1} [X - f_{\alpha\beta}(\mathbf{p}, \hat{\mathbf{X}}) \mathbf{p}]. \quad (9)$$

The unit vector $\mathbf{p}$ follows as a function of $\mathbf{q}$ by replacing $\mathbf{p} \leftrightarrow -\mathbf{q}$ and $\alpha \leftrightarrow \beta$. Here $g_{\alpha\beta}(\xi)$ is the root of

$$f_{\alpha\beta}[g_{\alpha\beta}(\xi)]g_{\alpha\beta}(\xi) + f_{\alpha\beta}(\xi)\xi - X = 0. \quad (10)$$

This equation may be solved to give

$$g_{\alpha\beta}(\xi) = \frac{1}{2}(h - X)(s_\alpha^2 X h - T_0^2)^{-1} \times [s_\beta T_0 + s_\alpha [T_0^2 + (s_\beta^2 - s_\alpha^2)Xh]^{1/2}], \quad (11)$$

where

$$h = 2f_{\alpha\beta}(\xi)\xi - X. \quad (12)$$

When $\alpha = \beta$, the expression for $g_{\alpha\beta}(\xi)$ simplifies somewhat, and the resulting form of Eq. (9) is

$$\mathbf{q} = (T_0^2 + s_\alpha^2 X^2 - 2s_\alpha T_0 \mathbf{p} \cdot \mathbf{X})^{-1} \times [(T_0^2 - s_\alpha^2 X^2)\mathbf{p} - 2s_\alpha(T_0 - s_\alpha \mathbf{p} \cdot \mathbf{X})\mathbf{X}]. \quad (13)$$

B. Determination of the ray direction

In order to complete the triangulation scheme, we require knowledge of either of the ray directions $\mathbf{p}$ or $\mathbf{q}$. Consider the unit vector $\mathbf{p}$, which is the direction of the ray leaving the source. In the Appendix it is shown how the spatial gradient of the travel time for a stationary ray is related to the ray direction at the source. This relation is a direction consequence of a generalized form of Fermat's Principle that includes stationary as well as minimum ray paths. From Eq. (A13) in the Appendix, we have for the homogeneous isotropic elastic material that

$$\mathbf{p} = -c_\alpha \nabla_s T_0, \quad (14)$$

where $\nabla_s T_0$ is the gradient of the travel time with respect to the source position $\mathbf{x}_s$. Similarly

$$\mathbf{q} = -c_\beta \nabla_Q T_0, \quad (15)$$

where $\nabla_Q T_0$ is the gradient of T_0 with respect to the receiver position $\mathbf{x}_Q$.

Experimentally, one could estimate the vector $\nabla_s T_0$ ($\nabla_Q T_0$) by shifting the source (receiver) successively in three linearly independent directions and measuring the time delay for each new position. For example, let the source be moved successively to the three new positions $\mathbf{x}_s + h\mathbf{e}_j$, $j = 1, 2, 3$, where h is some small distance and $\mathbf{e}_j$, $j = 1, 2, 3$, are three linearly independent unit vectors. Denoting by T_{0j} , $j = 1, 2, 3$, the measured value of the travel time for each new source position, we have by finite differences that

$$c_\alpha \nabla_s T_0 = \gamma, \quad (16)$$

where

$$\gamma_j = c_\alpha (T_{0j} - T_0)/h - \Delta\gamma_j, \quad j = 1, 2, 3 \quad (17)$$

and the computational error $|\Delta\gamma|$ is of order h times the Laplacian $c_\alpha \nabla_s^2 T_0$. By choosing different finite difference procedures, the computational error in the gradient of T_0 can be made of order h^2 or less. Later we analyze the effect of errors, both computational and experimental, on the inversion result.

In practice, one is interested in locating a defect inside an elastic body. The source and receiver are transducers,

which are positioned on the surface of the body and must remain directly coupled to the surface. Therefore the source (receiver) position has only two degrees of freedom. Consider the source at $\mathbf{x}_s$. Locally it may be moved in the tangent plane to the surface. Let $\mathbf{e}_1$ and $\mathbf{e}_2$ be two orthogonal unit vectors in this plane. Then by shifting the source successively in the directions $\mathbf{e}_1$ and $\mathbf{e}_2$, and forming finite differences as in Eq. (17), we may compute approximately the quantities $c_\alpha \mathbf{e}_j \cdot \nabla_s T_0$, $j = 1, 2$. Let us assume for the moment that the errors are zero. Then, from Eq. (14) we have the components of the unit vector $\mathbf{p}$ in the $\mathbf{e}_1$ and $\mathbf{e}_2$ directions. It remains to determine the component of $\mathbf{p}$ into the body. Since $|\mathbf{p}| = 1$, it follows that this component is equal to $[1 - (\mathbf{p} \cdot \mathbf{e}_1)^2 - (\mathbf{p} \cdot \mathbf{e}_2)^2]^{1/2}$. Similarly, we could determine the vector $\mathbf{q}$ by shifting the receiver in two orthogonal directions tangential to the surface and using Eq. (15).

Now suppose that we shift both source and receiver, but only in one direction for each. Let $\mathbf{e}_s$ and $\mathbf{e}_Q$ be unit tangent vectors to the surface at the source and receiver, respectively. Assuming no errors, computational or experimental, we can obtain the quantities γ and η , where

$$\gamma = c_\alpha \mathbf{e}_s \cdot \nabla_s T_0, \quad (18)$$

$$\eta = c_\beta \mathbf{e}_Q \cdot \nabla_Q T_0. \quad (19)$$

Our problem now reduces to finding $\mathbf{p}$ from the three equations

$$\mathbf{p} \cdot \mathbf{e}_s + \gamma = 0, \quad (20)$$

$$\mathbf{q} \cdot \mathbf{e}_Q + \eta = 0, \quad (21)$$

and

$$|\mathbf{p}| = 1. \quad (22)$$

From Eq. (9), we have $\mathbf{q} = \mathbf{q}(\mathbf{p})$. This relation, in combination with Eq. (10), can be used to rewrite Eq. (21) as the equation for a surface in $\mathbf{p}$ space. This surface is

$$f_{\alpha\beta}(\mathbf{p}, \hat{\mathbf{X}}) \mathbf{p} \cdot \mathbf{e}_Q - \mathbf{X} \cdot \mathbf{e}_Q + \eta f_{\alpha\beta}[g_{\alpha\beta}(\mathbf{p}, \hat{\mathbf{X}})] = 0, \quad (23)$$

where $f_{\alpha\beta}$ and $g_{\alpha\beta}$ were defined in Eqs. (6) and (11). This surface will intersect the plane defined by Eq. (20) along some curve and this curve will in general intersect the unit sphere of Eq. (22) at two points, $\mathbf{p}^+$ and $\mathbf{p}^-$. When $\alpha = \beta$, the surface (23) simplifies to the plane

$$\mathbf{p} \cdot \mathbf{e}'_Q + \eta' = 0, \quad (24)$$

where

$$\begin{aligned} \mathbf{e}'_Q &= \mathbf{e}_Q + 2(R^2 - X^2)^{-1}(\mathbf{e}_Q \cdot \mathbf{X} - R\eta)\mathbf{X} \\ &= (R^2 - X^2)\mathbf{e}_Q \cdot \nabla_Q [X/(R^2 - X^2)], \end{aligned} \quad (25)$$

$$\begin{aligned} \eta' &= \eta + 2(R^2 - X^2)^{-1}(\eta X^2 - R\mathbf{e}_Q \cdot \mathbf{X}) \\ &= (R^2 - X^2)\mathbf{e}_Q \cdot \nabla_Q [-R/(R^2 - X^2)], \end{aligned} \quad (26)$$

and $R = c_\alpha T_0$. Thus, for $\alpha = \beta$, the two points $\mathbf{p}^+$ and $\mathbf{p}^-$ may be found in closed form as the points of intersection of the two planes (20) and (24) with the unit sphere. The necessary condition for finding the points of intersection is that the planes (20) and (24) are not parallel, i.e., that

$$|\mathbf{e}_s \wedge \mathbf{e}'_Q| > 0. \quad (27)$$

When $\alpha \neq \beta$, the points of intersection of the three surfaces (20), (22), and (23) must be found by solving an algebraic equation.

The dilemma of choosing the vector $\mathbf{p}$ from the pair $\mathbf{p}^+$ and $\mathbf{p}^-$ can usually be resolved quite simply. For example, if $\alpha = \beta$ and $\mathbf{e}_s, \mathbf{e}_Q$, and $\mathbf{X}$ are all coplanar, then the only difference in $\mathbf{p}^+$ and $\mathbf{p}^-$ is in their components normal to this plane, i.e., only one of them is directed into the solid. However, if the source and receiver are on opposite sides of a slab with parallel sides, such that $\mathbf{e}_s$ and $\mathbf{e}_Q$ are parallel to one another but perpendicular to $\mathbf{X}$, then extra information is required to choose between $\mathbf{p}^+$ and $\mathbf{p}^-$.

In summary, we have three alternative methods of determining the unit vector $\mathbf{p}$ (or equivalently $\mathbf{q}$). The first method (I) consists of shifting the source transducer in two directions tangential to the body surface, and forming directional derivatives of the travel time to the fixed receiver. In the second method (II), the roles of the source and receiver are reversed, i.e., the latter is shifted while the former is not. In the third method (III), both source and receiver are shifted, but each in only one direction. If the received signal is mode converted, and in addition $T_0 > s_T X$, then only one of methods I and II will produce a unique flash point. The correct method is the one in which the endpoint corresponding to the faster wave speed is kept fixed; for example, if $\alpha = L$ and $\beta = T$, then it is method II.

C. Error analysis

In practice the inversion result will be incorrect due to experimental and computational errors. The former are the inevitable result of inaccurate measurements of such quantities as the source position, the delay time of the first arriving signal, etc. It is assumed that errors due to inhomogeneity of the host material and variations of the wave speed along the ray path are negligible. Often this may not be the case, which may lead to errors on the order of magnitude of the size of the scatterer. We define the computational errors as those incurred in using the finite difference approximation to the gradients. For the moment we consider all errors together and find the resultant error in the flash point position.

We consider the inversion using method I for arbitrary type of the emitted and received waves. Without loss of generality we may assume that the source is shifted in two orthogonal directions, specified by the unit vectors $\mathbf{e}_1$ and $\mathbf{e}_2$. Define the dimensionless numbers γ_1 and γ_2 to be the actual values of the directional derivatives of the travel time, i.e.,

$$\gamma_j = c_\alpha \mathbf{e}_j \cdot \nabla_s T_0, \quad j = 1, 2. \quad (28)$$

Let $\gamma_j + \Delta\gamma_j$ be the corresponding values computed from the experimental data. Thus $\Delta\gamma_j$ incorporates both experimental and computational errors. We assume that the major source of error in the flash point position is attributable to the errors $\Delta\gamma_j$. If $\Delta\gamma$ is small, where

$$\Delta\gamma = [(\Delta\gamma_1)^2 + (\Delta\gamma_2)^2]^{1/2}, \quad (29)$$

then a linear error analysis produces bounds on the error $|\Delta\mathbf{x}_{\alpha\beta}|$ in the position of the flash point. These bounds are as follows:

$$r_s \Delta\gamma \leq |\Delta\mathbf{x}_{\alpha\beta}| < M r_s \Delta\gamma, \quad (30)$$

where

$$M^2 = [1 + (s_\beta X / s_\alpha r_Q)^2] / (1 - \gamma^2) \quad (31)$$

and

$$\gamma^2 = \gamma_1^2 + \gamma_2^2. \quad (32)$$

In order to derive this result, we consider the error induced in the computed value of the unit vector $\mathbf{p}$, which points in the direction of the flash point. We have that $\mathbf{p} = \mathbf{p}^+$ or $\mathbf{p}^-$ where

$$\mathbf{p}^\pm = -\gamma_1 \mathbf{e}_1 - \gamma_2 \mathbf{e}_2 \pm (1 - \gamma^2)^{1/2} \mathbf{e}_1 \wedge \mathbf{e}_2. \quad (33)$$

The error in the calculated value of $\mathbf{p}$ is $\Delta\mathbf{p} = \Delta\mathbf{p}^+$ or $\Delta\mathbf{p}^-$, where

$$\begin{aligned} \Delta\mathbf{p}^\pm = & -\Delta\gamma_1 \mathbf{e}_1 - \Delta\gamma_2 \mathbf{e}_2 \\ & \mp (\gamma_1 \Delta\gamma_1 + \gamma_2 \Delta\gamma_2) (1 - \gamma^2)^{-1/2} \mathbf{e}_1 \wedge \mathbf{e}_2. \end{aligned} \quad (34)$$

Thus the absolute magnitude of $\Delta\mathbf{p}$ is

$$\Delta p = [(1 - \gamma^2)(\Delta\gamma_1)^2 + (1 - \gamma_1^2)(\Delta\gamma_2)^2]^{1/2} (1 - \gamma^2)^{-1/2}, \quad (35)$$

from which it follows that

$$\Delta\gamma \leq \Delta p \leq (1 - \gamma^2)^{-1/2} \Delta\gamma. \quad (36)$$

The error in $\mathbf{p}$ induces a resultant error in the estimated flash point position, which follows from Eq. (3) as

$$\Delta\mathbf{x}_{\alpha\beta} = r_s \Delta\mathbf{p} + (\Delta\mathbf{p} \cdot \nabla_p r_s) \mathbf{p}. \quad (37)$$

The gradient of r_s with respect to $\mathbf{p}$ may be calculated from the quadratic equation

$$r_s^2 (s_\alpha^2 - s_\beta^2) - 2r_s (s_\alpha T_0 - s_\beta^2 \mathbf{p} \cdot \mathbf{X}) + T_0^2 - s_\beta^2 X^2 = 0. \quad (38)$$

Differentiation of this equation with respect to $\mathbf{p}$ produces $\nabla_p r_s$, which when inserted in Eq. (37) gives

$$\begin{aligned} \Delta\mathbf{x}_{\alpha\beta} = & [(s_\alpha r_Q + s_\beta r_s) \Delta\mathbf{p} \\ & + s_\beta \mathbf{X} \wedge (\mathbf{p} \wedge \Delta\mathbf{p})] r_s / (s_\alpha r_Q + s_\beta r_s - s_\beta \mathbf{p} \cdot \mathbf{X}). \end{aligned} \quad (39)$$

Since $\mathbf{p}$ is a unit vector, it follows that $\mathbf{p} \cdot \Delta\mathbf{p} = 0$ to first order, i.e., $\mathbf{p}$ and $\Delta\mathbf{p}$ are perpendicular. Hence by Eq. (39)

$$|\Delta\mathbf{x}_{\alpha\beta}| = r_s [(\Delta p)^2 + (s_\beta \mathbf{X} \cdot \Delta\mathbf{p})^2 / (s_\alpha r_Q + s_\beta r_s - s_\beta \mathbf{p} \cdot \mathbf{X})^2]^{1/2}.$$

Now,

$$s_\alpha r_Q + s_\beta r_s - s_\beta \mathbf{p} \cdot \mathbf{X} = r_Q (s_\alpha + s_\beta \mathbf{p} \cdot \mathbf{q}) > s_\alpha r_Q, \quad (40)$$

because by definition, $\mathbf{p} \cdot \mathbf{q}$ must be positive. Therefore we obtain the following bounds for the flash point error:

$$r_s \Delta p \leq |\Delta\mathbf{x}_{\alpha\beta}| < r_s [1 + (s_\beta X / s_\alpha r_Q)^2]^{1/2} \Delta p. \quad (41)$$

Combining this with Eq. (36) we arrive at the result (30).

For example, if a difference scheme like the one described in Eq. (17) is used to compute γ_1 and γ_2 , then the computational error is

$$\Delta\gamma = O(h c_\alpha \nabla_s^2 T_0). \quad (42)$$

The Laplacian of the travel time follows quite simply from Eq. (14) as

$$c_\alpha \nabla^2 T_0 = -\nabla_s \cdot \mathbf{p}. \quad (43)$$

The divergence of $\mathbf{p}$ with respect to the source position may be obtained as follows: Consider the reverse signal which is emitted from the point $\mathbf{x}_Q$ as a wave of type β , scatters from the flash point $\mathbf{x}_{\alpha\beta}$, and is observed at the point $\mathbf{x}_s$ as a wave of type α . The travel time for this signal is T_0 , and its dependence on the location of $\mathbf{x}_s$ is the same as that of the $S \rightarrow Q$

signal. The reverse scattered signal propagates as a wave front which passes $\mathbf{x}_s$ at time T_0 after emission at $\mathbf{x}_Q$. Let ρ_{s1} and ρ_{s2} be the principal radii of curvature of the wave front surface as it passes $\mathbf{x}_s$, defined so as to be negative if the center of curvature is in the direction of propagation of the wave front. Thus the vector $\mathbf{p}$ is the unit normal to the wave front surface, pointed in the direction of positive curvature. By resolving the gradient operator ∇_s into three components along $\mathbf{p}$ and along the principal lines of curvature we deduce that

$$-\nabla_s \cdot \mathbf{p} = \rho_{s1}^{-1} + \rho_{s2}^{-1}. \quad (44)$$

From Eqs. (30), (42)–(44), we have for the first-order difference scheme, that

$$|\Delta \mathbf{x}_{\alpha\beta}| = O(Mhr_s/\rho_s), \quad (45)$$

where $\rho_s^2 = \text{Min}(\rho_{s1}^2, \rho_{s2}^2)$.

Consider the special case in which the flash point lies on a smooth edge. This occurs, for example, if the scattering object is a flat crack, with a smooth edge. The observed signal is diffracted from the edge and the flash point may be predicted by the geometrical theory of diffraction (GTD). A description of GTD for problems of elastodynamic diffraction by cracks is given by Achenbach and Gautesen.⁷ As far as the radii of curvature of the diffracted signal are concerned, it is immaterial whether the edge is that of a crack or of a wedge-like inclusion. Whatever the local geometry of the edge at the flash point, the principal radii ρ_{s1} and ρ_{s2} are

$$\rho_{s1} = r_s, \quad \rho_{s2} = r_s + \bar{\rho}_\beta^\alpha, \quad (46)$$

where $\bar{\rho}_\beta^\alpha$ is the signed distance from the flash point to the caustic. The general form of $\bar{\rho}_\beta^\alpha$ for an arbitrary incident wave front is given by Gautesen *et al.*⁹ In our case, the incident wave front is spherical, with center at $\mathbf{x}_Q$. We find that

$$\bar{\rho}_\beta^\alpha = c_\beta \sin^2 \phi_\alpha [\mathbf{n} \cdot (c_\beta \mathbf{p} + c_\alpha \mathbf{q})/a + c_\alpha \sin^2 \phi_\beta / r_Q]^{-1}. \quad (47)$$

The angles ϕ_α and ϕ_β are the angles between the direction vectors $\mathbf{p}$ and $\mathbf{q}$ and the tangent to the edge, respectively. The unit vector $\mathbf{n}$ is directed from the flash point towards the center of curvature of the edge, and a is the corresponding radius of curvature. Snell's law for diffraction is

$$c_\alpha \sin \phi_\beta = c_\beta \sin \phi_\alpha. \quad (48)$$

If $\mathbf{n} \cdot [\mathbf{p} + (c_\alpha/c_\beta)\mathbf{q}] = o(1)$, the diffracted ray is near either the shadow or reflection boundary. We assume, as is generally the case, that the signal is in the zone of pure diffraction. In this case, the above quantity is of order unity. If, in addition, the radius of curvature of the edge is small as compared with the distances r_s and r_Q , then by Eq. (47),

$$\bar{\rho}_\beta^\alpha = O(a), \quad (49)$$

and by Eq. (46),

$$\rho_{s2} = r_s + o(r_s). \quad (50)$$

Therefore, from Eqs. (45), (46), and (50), we have

$$|\Delta \mathbf{x}_{\alpha\beta}| = O(Mh). \quad (51)$$

This states that the flash point error is proportional to the transducer shift h if the simplest difference scheme is adopted. The same result could also be shown for methods II and III of above.

The minimum number of shifts required in order to proceed with the inversion is two, whether the method used is I, II, or III. This corresponds to taking three measurements of the travel time. The only possible difference scheme is the one described in Eq. (17), and the resultant error in the flash point position is of order h . However, if one more shift is performed, then a difference scheme can be used which ensures that the error is of order $O(h^2)$. Thus one more measurement increases the computational accuracy dramatically. To see this, we consider an example using method III and synthetic data.

D. Inversion of synthetic data

We present an example of inversion of synthetic crack-scattering data. The travel times are calculated exactly and the only errors come from the size of the transducer shifts and the finite difference scheme used. Only the first received diffracted signals are considered. Both the incident and diffracted waves are longitudinal, i.e., $\alpha = \beta = L$. Method III of above is used, and two difference schemes are compared for accuracy.

Consider a flat crack with a smooth circular edge. Let the radius of the crack be unity, and define a rectangular coordinate system (x, y, z) with origin at the center of the circle and the z axis normal to the crack. Two transducers, our source and receiver, are located at the points $\mathbf{x}_s = (5, 10, 7)$ and $\mathbf{x}_Q = (10, 5, 7)$, respectively. The exact position of the flash point for the first received signal is at $(1/\sqrt{2}, 1/\sqrt{2}, 0)$, and the travel time may be calculated easily. The source is then shifted in the direction of the vector $\mathbf{U} = (1, 1, -2)/\sqrt{6}$ by an amount h . The travel time for the first diffracted signal is again calculated. We note that the flash point position will be slightly different. Now the receiver is shifted in the same direction by the same amount and the travel time is calculated. These three travel times provide sufficient information with which to form gradients according to Eq. (17), and then use method III to invert to find the flash point. For $h = 0.1, 0.5, 1$, and 2 , the results of the inversion and its accuracy are given in Table I. We observe the expected linear error growth with h , in agreement with Eq. (51). Also, we note that the calculations do not involve Poisson's ratio explicitly, since the travel time and the longitudinal wave speed occur only in the combination $c_L T_0$, which is equal to the ray path length. The other data in Table I result from shifting the source back to its original position and calculating the travel time once more. This additional information allows us to take the actual source and receiver positions at $\mathbf{x}_s + (h/2)\mathbf{U}$ and $\mathbf{x}_Q + (h/2)\mathbf{U}$, respectively. The gradients are now computed by a centered difference scheme, and are correct to order h^2 . The error in the position of the computed flash point is seen from Table I to be also of order h^2 . In fact it is very small, even for shifts as large as the dimension of the crack.

In summary, two transducer shifts provide the minimum information necessary for inversion, and the error in the calculated position is linear. However, one more shift produces much greater computational accuracy. Two shifts imply three travel time measurements, while three shifts imply four such measurements.

II. LOCAL CRACK-EDGE MAPPING

The methods discussed above are useful for identifying single points on a scatterer. In order to determine the size and shape of the scatterer, we would have to repeat the inversion procedure a number of times. Each inversion produces one point, the flash point, so that eventually we would hope to have enough points to characterize the scatterer. In the present section we discuss an alternate method of mapping the scatterer. This method is again based on transit times from source to receiver via a flash point. However, instead of using single transit times to produce single flash points, we now use a set of transit times to simultaneously produce a set of flash points. Essentially, the method is an iterative procedure based on Newton's method of approximation. To start with, we assume that a point $\mathbf{x}_0$ near the scatterer is known. It may have been determined by the global triangulation method. Then, the surface $E_{\alpha\beta}$ on which the flash points lie may be approximated by their tangent planes near $\mathbf{x}_0$. For a crack the intersections of the tangent planes give a polygon which approximates the curve of flash points. The procedure may then be repeated by rechoosing $\mathbf{x}_0$, so that a convergent iteration ensues. In what follows, we assume that the scatterer is a flat crack with a smooth edge, although it may be seen that the method would also work for convex voids or inclusions.

A. Tangent plane approximation

We have, as before, a source S at $\mathbf{x}_S$ and a receiver Q at $\mathbf{x}_Q$. The travel time T_0 for the diffracted signals of type (α, β) is assumed known. Then the flash point $\mathbf{x}_{\alpha\beta}$ lies on the surface $E_{\alpha\beta}$ defined by Eq. (2). The unit vectors $\mathbf{p}$ and $\mathbf{q}$ and the distances r_S and r_Q are defined in Eq. (3). Define the new distance R as

$$R = \kappa_\alpha r_S + \kappa_\beta r_Q, \quad (52)$$

where $\kappa_\alpha = c_L/c_\alpha$, $\alpha = L, T$. With respect to an arbitrary point $\mathbf{x}_0$, we define the unit vectors $\mathbf{p}_0$ and $\mathbf{q}_0$, and the distances r_{S0} , r_{Q0} , and R_0 as follows:

$$\mathbf{x}_0 = \mathbf{x}_S + r_{S0}\mathbf{p}_0 = \mathbf{x}_Q + r_{Q0}\mathbf{q}_0, \quad (53)$$

$$R_0 = \kappa_\alpha r_{S0} + \kappa_\beta r_{Q0}. \quad (54)$$

Now, for all $\mathbf{x} \in E_{\alpha\beta}$,

$$\mathbf{x} - \mathbf{x}_0 = r_S\mathbf{p} - r_{S0}\mathbf{p}_0. \quad (55)$$

Taking the dot product of this with $\mathbf{p}_0$ we obtain

$$(\mathbf{x} - \mathbf{x}_0) \cdot \mathbf{p}_0 = r_S - r_{S0} + r_S(\mathbf{p} \cdot \mathbf{p}_0 - 1). \quad (56)$$

Doing the same thing in terms of r_0 , $\mathbf{q}$ etc., we deduce that

$$(\mathbf{x} - \mathbf{x}_0) \cdot (\kappa_\alpha \mathbf{p}_0 + \kappa_\beta \mathbf{q}_0) - (R - R_0) = -\lambda, \quad (57)$$

where the distance λ is

$$\lambda = \kappa_\alpha r_S(1 - \mathbf{p} \cdot \mathbf{p}_0) + \kappa_\beta r_Q(1 - \mathbf{q} \cdot \mathbf{q}_0) > 0. \quad (58)$$

For any $\mathbf{x} \in E_{\alpha\beta}$, define the distance d as

$$d = |\mathbf{x} - \mathbf{x}_0|. \quad (59)$$

From the cosine rule for the triangle $\mathbf{x}_S$, $\mathbf{x}_0$, and $\mathbf{x}$, we have

$$d^2 = r_S^2 + r_{S0}^2 - 2r_S r_{S0} \mathbf{p} \cdot \mathbf{p}_0. \quad (60)$$

Hence,

$$1 - \mathbf{p} \cdot \mathbf{p}_0 = [d^2 - (r_S - r_{S0})^2] / (2r_S r_{S0}). \quad (61)$$

Similarly we can calculate the quantity $1 - \mathbf{q} \cdot \mathbf{q}_0$, and combine the results to get

$$\begin{aligned} \lambda &= \frac{\kappa_\alpha}{2r_{S0}} [d^2 - (r_S - r_{S0})^2] + \frac{\kappa_\beta}{2r_{Q0}} [d^2 - (r_Q - r_{Q0})^2] \\ &< \frac{1}{2} \left(\frac{\kappa_\alpha}{r_{S0}} + \frac{\kappa_\beta}{r_{Q0}} \right) d^2. \end{aligned} \quad (62)$$

Consider a segment of the surface $E_{\alpha\beta}$, say $F_{\alpha\beta}$, and define the positive number δ_F as

$$\delta_F = [(\kappa_\alpha r_{Q0} + \kappa_\beta r_{Q0}) / (2r_{S0} r_{Q0} R_0)] \sup_{\{\mathbf{x} \in F_{\alpha\beta}\}} (d^2). \quad (63)$$

For a given source and receiver, the number δ_F depends on the arbitrary point $\mathbf{x}_0$ and the segment $F_{\alpha\beta}$. From Eqs. (58), (62), and (63) it follows that the segment $F_{\alpha\beta}$ lies between the two planes:

$$(\mathbf{x} - \mathbf{x}_0) \cdot (\kappa_\alpha \mathbf{p}_0 + \kappa_\beta \mathbf{q}_0) - (R - R_0) = 0 \quad (64)$$

and

$$(\mathbf{x} - \mathbf{x}_0) \cdot (\kappa_\alpha \mathbf{p}_0 + \kappa_\beta \mathbf{q}_0) - (R - R_0) = -R_0 \delta_F. \quad (65)$$

Obviously, in the limit as $\delta_F \rightarrow 0$, the segment $F_{\alpha\beta}$ becomes a segment of the first plane. However, δ_F can equal zero only if $F_{\alpha\beta} = \{\mathbf{x}_0\}$. If $F_{\alpha\beta}$ is taken as some part of $E_{\alpha\beta}$ which includes the flash point, then $\delta_F > 0$. Also, if the dimensions of $F_{\alpha\beta}$ are small compared to the distances r_S and r_Q , and in addition, $\mathbf{x}_0$ is near the flash point, then δ_F will be much smaller than unity. Therefore we may approximate the surface $E_{\alpha\beta}$ near the flash point by the plane (64).

B. Example: Elliptical crack

For the remainder of this section, we consider the inversion of crack-scattering data. In order to demonstrate the local edge mapping, we consider the following situation: A flat crack with a smooth convex edge is located within a body. There are n source positions S_i , $i = 1, n$, and m observer positions Q_j , $j = 1, m$. For each source observer pair (S_j, Q_j) there are four pairs of associated travel times, corresponding to the four types of diffraction processes. We note

TABLE I. Computed flash point and errors. The error is the distance from the exact flash point (0.707, 0.707, 0).

h	Computed point	Two shifts		Three shifts	
		Error	Computed point	Error	
0.1	(0.728, 0.728, -0.405)	0.050	(0.707, 0.707, 0.000)	0.000	
0.5	(0.810, 0.810, -0.204)	0.250	(0.709, 0.702, -0.002)	0.005	
1.0	(0.909, 0.909, -0.410)	0.500	(0.715, 0.688, -0.006)	0.022	
2.0	(1.01, 1.01, -0.832)	1.000	(0.738, 0.629, -0.024)	0.088	

that the sources and receivers may be identical, corresponding to pulse-echo measurements. Suppose that an initial point $\mathbf{x}_0$ has been determined by the method of triangulation, or some other *a priori* procedure. Our starting point for the local edge mapping is to assume that the flash points lie on planes specified by Eq. (64). Then all of the planes considered pass very close to the crack edge; in fact they envelope a curve which is an approximation to a segment of the crack edge. The inversion thus reduces to finding the congruence of all planes. If we had an infinite number of sources and receivers located arbitrarily on a sphere of very large radius with the crack at the center, then the congruence would be a smooth closed curve. Since our data is finite, we will obtain a polygon of points which approximate a segment of the crack edge.

If the geometry of the problem is two dimensional, then the crack is completely specified by the two flash points at either end of the crack. The planes become lines and the edge mapping reduces to finding the intersection of the lines. We refer the reader to Ref. 5 for a complete discussion of the two-dimensional problem.

Let us consider the general three-dimensional problem of constructing the crack edge. For simplicity we consider the following example. We take as our crack edge the ellipse

$$x^2 + 4y^2 - 1 = 0, \quad z = 0. \quad (66)$$

There are two sources and 20 receivers situated on the plane $z = 10$. The two sources $S_i, i = 1, 2$ are at

$$S_1: (-10, 10, 10), \quad S_2: (-5, 10, 10),$$

while the 20 receivers are at

$$Q_j: (5, 15, 10) + (j/2)(1, -1, 0), \quad j = 1, 20.$$

Our first task will be to determine the plane of the crack approximately. Having found the crack plane, the crack edge is obtained as the polygon formed by the intersection of the crack plane and the planes of Eq. (64). For a given type of received signal (L or T) there are two sets of 40 planes corresponding to the first and second arriving signals. Let $\Omega(i, j), i = 1, 2, j = 1, 20$ be such a set of planes. The plane of the crack may be found as follows: Consider the intersection of the planes $\Omega(1, j), j = 1, 20$ with any one of the planes $\Omega(2, j), j = 1, 20$. The locus of intersection is a polygon, which we denote by $\Gamma(2, k)$, where $\Omega(2, k)$ is the plane which intersects the 20 planes corresponding to source S_1 . To obtain points which approximate the flash points, we test to see if any of the planes $\Omega(2, j), j \neq k$ intersect the polygon $\Gamma(2, k)$. For such a point of intersection to exist, the flash point of $\Omega(2, j)$ at $j = k$ must be interspersed with the flash points of $\Omega(2, j), j = 1, 20, j \neq k$. We also require that the segment of the crack edge containing the flash points of $\Omega(1, j), j = 1, 20$ has a section in common with the segment for the planes $\Omega(2, j), j = 1, 20$. Such is the case for our configuration of sources and receivers.

We require at least three points in order to specify the plane of the crack. Once this has been achieved, the remaining approximate flash points are easily determined as being on the polygon formed by the set of planes $\Omega(i, j), i = 1, 2, j = 1, 20$ and the crack plane.

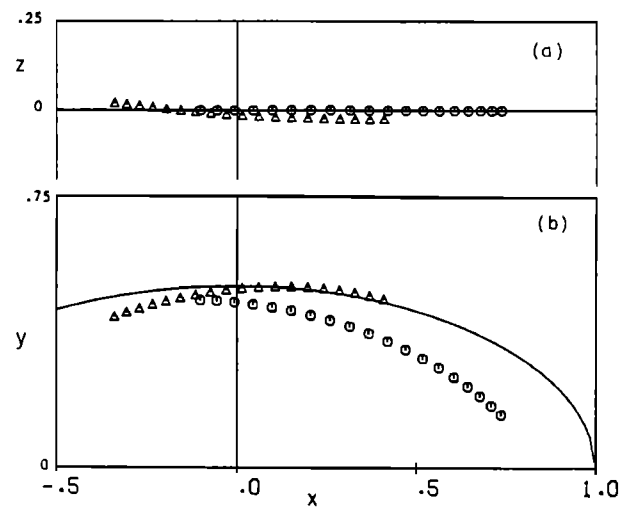


FIG. 2. Comparison of computed edge points and actual crack edge (solid line). Base point $\mathbf{x}_0 = (0.5, 1, 0.5)$. Using $L-L$ travel times (Δ), $L-T$ travel times ($\odot$). (a) Projection on plane $y = 0$, and (b) projection on plane $z = 0$.

In Figs. 2 and 3 are plotted approximate flash points for two different values of the initial point $\mathbf{x}_0$. In Fig. 2 we have $\mathbf{x}_0 = (0.5, 1, 0.5)$ while in Fig. 3, $\mathbf{x}_0 = (1, 2, 1)$. Only the first arriving signals of type (L, L) and (L, T) were considered. The indicated points are the vertices of the polygons formed by the intersection of the planes $\Omega(1, j), j = 1, 20$ with the calculated plane of the crack. In Figs. 2(b) and 3(b) the actual crack edge and the approximate flash points are viewed together from a position directly above the crack. The side-on view (i.e., $z = 0$), looking in the positive y direction, is given in Figs. 2(a) and 3(a).

We note the better approximation of the computed flash points to the actual edge when $\mathbf{x}_0$ is taken at $(0.5, 1, 0.5)$. The agreement is much poorer for the initial point at $(1, 2, 1)$ as shown in Fig. 3. Since the accuracy is never known, an iterative procedure should be used. In this procedure a new

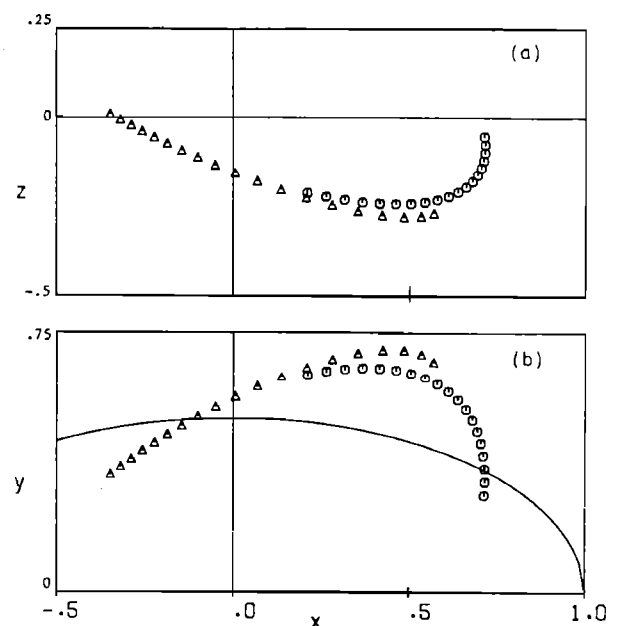


FIG. 3. Same as Fig. 2, except base point at $(1, 2, 1)$.

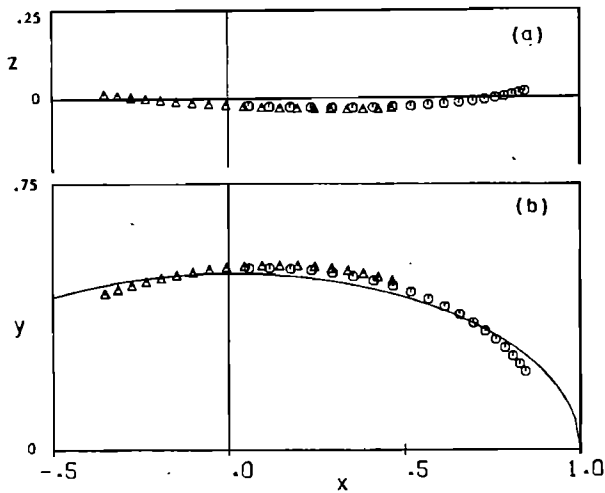


FIG. 4. Same as Fig. 2, except base point at $(0.2, 0.5, -0.5)$.

initial point is taken, say at one of the estimated points of Fig. 3. Then the mapping procedure is repeated. In general such an iteration could be performed several times until all the points converge to fixed points on the edge. In Fig. 4 we demonstrate the improved accuracy in the results of Fig. 2 for one iteration. We note that the (L, T) signals define a different segment of the edge than the (L, L) signals. Similarly, additional points follow by considering the (T, L) and (T, T) signals.

In addition to the first arriving signals, which are those diffracted from the near edge of the crack, there are also signals diffracted from the far edge. We refer to these as the second arriving signals of type (L, L) , (L, T) , etc. The travel times for these signals may be inferred from the spacing of the peaks in the high-frequency spectrum (see Ref. 8 for a

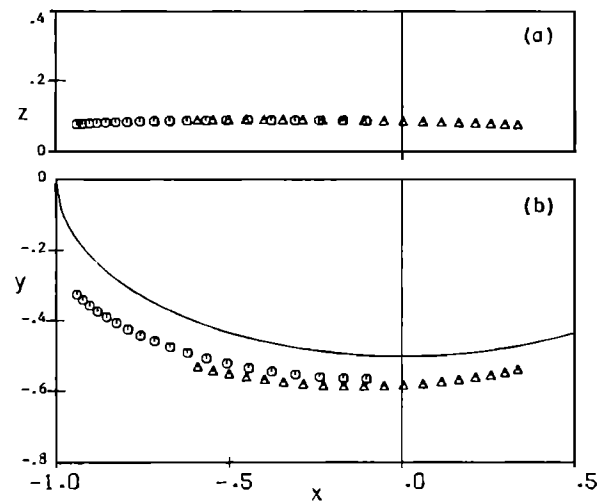


FIG. 6. Same as Fig. 5, except base point at $(-0.7, -0.75, 0.3)$.

further discussion on the use of the high-frequency spectrum). The flash points corresponding to the second arriving signals of type (L, L) and (L, T) are considered in Figs. 5 and 6. In Fig. 5 the same value of x_0 was used as in Fig. 2. The plane of the crack was taken the same as in Figs. 2–4. The result of a single iteration, in which the initial point is taken at one of the points of Fig. 5, is illustrated in Fig. 6. By combining Figs. 4 and 6 it is evident that a good estimate of the size of the crack can be obtained from the diffracted signals.

ACKNOWLEDGMENTS

This work was carried out in the course of research sponsored by the Center for Advanced NDE operated by the Science Center, Rockwell International for the Advanced Research Project Agency and the Air Force Materials Laboratory under Contract F33615-80-C-5004.

APPENDIX: APPLICATION OF FERMAT'S PRINCIPLE

The connection between the spatial gradient of the travel time for a stationary ray path and the ray direction is established in this Appendix. For generality, we consider an inhomogeneous anisotropic medium, although the result is required only for the homogeneous isotropic case.

Consider a curve $\mathbf{x}(t)$ between the two points $\mathbf{x}_S$ and $\mathbf{x}_Q$ such that $\mathbf{x}(t_0) = \mathbf{x}_S$ and $\mathbf{x}(t_1) = \mathbf{x}_Q$, $t_0 < t_1$. Suppose a signal is propagated from $\mathbf{x}_S$ to $\mathbf{x}_Q$ along this curve such that the signal propagation speed at any point is a function of position and of the direction of propagation. If c is the speed, then $c(t) = c[\mathbf{x}(t), \mathbf{p}(t)]$, where $\mathbf{p}(t)$ is the unit tangent vector to the curve,

$$\mathbf{p}(t) = \dot{\mathbf{x}}(t)/|\dot{\mathbf{x}}(t)|, \quad (\text{A1})$$

and the dot denotes differentiation. Fermat's Principle states that the signal describes a stationary ray path if the curve $\mathbf{x}(t)$ makes the travel time

$$T = \int_{t_0}^{t_1} c^{-1}(t) |\dot{\mathbf{x}}(t)| dt \quad (\text{A2})$$

stationary. The simplest stationary ray path is the curve that minimizes T . Other curves may minimize or maximize T

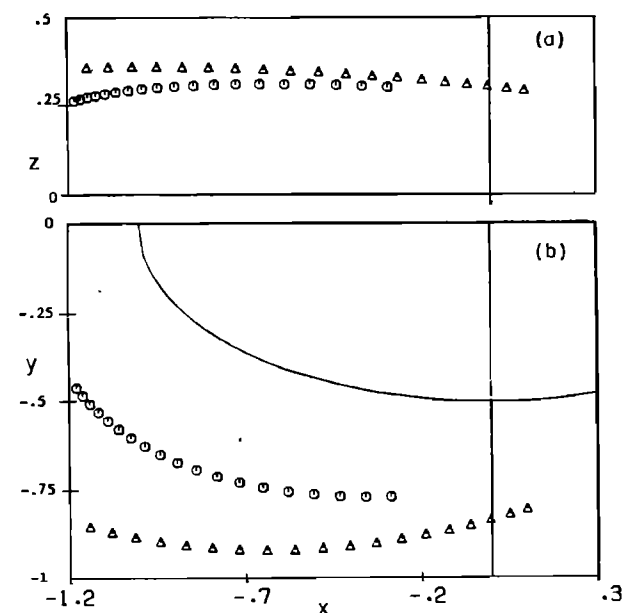


FIG. 5. Computed flash points corresponding to arrival of signals from second flash point of type $L-L$ (Δ) and $L-T$ ($\circ$). Base point at $(0.5, 1, 0.5)$.

among a certain class of curves. For example, suppose a compact inhomogeneity is present in an elastic medium, then there is a curve that minimizes T with the constraint that the curve includes a point on the inhomogeneity. If the inhomogeneity is a crack with a smooth edge, this ray path might correspond to the first observed diffracted ray. Other curves may minimize T among the class of multiply diffracted rays. Similarly, there is a curve that maximizes T among the class of singly diffracted rays. For a further discussion about classes of stationary ray paths, we refer to the paper by Keller.¹⁰

Suppose that $\mathbf{x}(t)$ is a stationary ray path. Let $L(\mathbf{x}, \dot{\mathbf{x}}; t)$ denote the integrand of Eq. (A2), i.e.,

$$L \equiv c^{-1}(t)|\dot{\mathbf{x}}(t)|, \quad (\text{A3})$$

and define the "momentum" $\mathbf{m}(t)$ corresponding to the "Lagrangian" L as

$$\mathbf{m} \equiv \frac{\partial L}{\partial \dot{\mathbf{x}}}. \quad (\text{A4})$$

The Euler-Lagrange equations of the variational problem then are

$$\dot{\mathbf{m}}(t) = -\nabla L(\mathbf{x}, \dot{\mathbf{x}}; t), \quad (\text{A5})$$

and the associated "Hamiltonian" $H(\mathbf{x}, \mathbf{m}; t)$ is

$$H \equiv \dot{\mathbf{x}} \cdot \mathbf{m} - L. \quad (\text{A6})$$

Substitution of the explicit form for L of Eq. (A3) into the definition of $\mathbf{m}$ yields

$$\mathbf{m} = (\mathbf{p} - \bar{\nabla} \ln c)/c, \quad (\text{A7})$$

where the operator $\bar{\nabla}$ is equal to $\partial/\partial \mathbf{p}$. Since the vector $\mathbf{p}$ is of unit magnitude, it follows that $\bar{\nabla}$ is the angular gradient operator and that

$$\mathbf{p} \cdot \bar{\nabla} c = 0. \quad (\text{A8})$$

This result, combined with the definition of H , implies that the "Hamiltonian" is identically zero.

Now consider the variation in the stationary value of T due to a variation in the position of the endpoint $\mathbf{x}_S$. Keeping the other endpoint $\mathbf{x}_Q$ fixed, let the variation in $\mathbf{x}_S$ be of the most general type:

$$\delta \mathbf{x}_S = [\delta \mathbf{x}(t) + \dot{\mathbf{x}} \delta t]_{t=t_0}. \quad (\text{A9})$$

The variation of T is

$$\begin{aligned} \delta T &= -L(\mathbf{x}, \dot{\mathbf{x}}; t_0) \delta t_0 + \int_{t_0}^{t_1} (\delta \mathbf{x} \cdot \dot{\mathbf{m}} + \delta \dot{\mathbf{x}} \cdot \mathbf{m}) dt, \\ &= -(L \delta t + \mathbf{m} \cdot \delta \mathbf{x})|_{t=t_0}. \end{aligned} \quad (\text{A10})$$

Here we have used the Euler-Lagrange equations (A9), and the fact that $\mathbf{x}_Q$ is fixed. Therefore, by Eqs. (A9) and (A10), we have

$$\frac{\partial T}{\partial t_0} = H(\mathbf{x}, \dot{\mathbf{x}}; t_0) \equiv 0 \quad (\text{A11})$$

and

$$\nabla_S T = -\mathbf{m}(t_0) = -(\mathbf{p} - \bar{\nabla} \ln c)/c|_{t=t_0}, \quad (\text{A12})$$

where the subscript S indicates that the gradient is evaluated at the endpoint $\mathbf{x}_S$. This is our basic result relating the spatial gradient of the time delay T to the ray direction $\mathbf{p}$ and the (anisotropic) speed c . In the case of isotropy, $\bar{\nabla} c = 0$, and we have the simple result

$$c \nabla_S T + \mathbf{p} = 0. \quad (\text{A13})$$

Finally, we note that in homogeneous linear anisotropic elastic materials, the velocity of signal propagation, which is the velocity with which energy is propagated, is not equal to the phase velocity. Let $v\mathbf{n}$, $|\mathbf{n}| = 1$, be the phase velocity corresponding to the energy propagation or group velocity $c\mathbf{p}$, $|\mathbf{p}| = 1$. Define the phase slowness s in the direction $\mathbf{n}$ as the inverse of the phase speed v . Then it may be shown that the group and phase velocities are related by

$$s\mathbf{n} = (\mathbf{p} - \bar{\nabla} \ln c)/c, \quad (\text{A14})$$

see, for example, Musgrave.¹¹ Therefore, combining Eqs. (A12) and (A14), we see that

$$v \nabla_S T + \mathbf{n} = 0 \quad (\text{A15})$$

in homogeneous linear elastic media. Equation (A14) states that the points $c\mathbf{p}$ and $s\mathbf{n}$ are polar reciprocal (see Ref. 11) to each other. The inverse relationship is

$$c\mathbf{p} = (\mathbf{n} - \bar{\nabla} \ln s)/s, \quad (\text{A16})$$

where $\bar{\nabla} \equiv \partial/\partial \mathbf{n}$ is the angular gradient operator with respect to $\mathbf{n}$. By combining Eqs. (A15) and (A16) we see that there is a one-to-one relationship between the spatial gradient of the travel time $\nabla_S T$ and the ray direction $\mathbf{p}$.

¹N. N. Bojarski, *Three Dimensional Electromagnetic Short Pulse Inverse Scattering* (Syracuse University Res. Corp., Syracuse, NY, 1967).

²R. M. Lewis, "Physical Optics Inverse Scattering," *IEEE Trans. Antennas Propag.* AP-17, 308-314 (1969).

³W. M. Boerner and C. M. Ho, "Analysis of Physical Optics Far Field Inverse Scattering for the Limited Data Case using Radon Theory and Polarization Information," *Wave Motion* 3, 311-334 (1981).

⁴J. K. Cohen and N. Bleistein, "The Singular Function of a Surface and Physical Optics Inverse Scattering," *Wave Motion* 1, 153-161 (1979).

⁵J. D. Achenbach, K. Viswanathan, and A. Norris, "An Inversion Integral for Crack Scattering Data," *Wave Motion* 1, 299-316 (1979).

⁶S. Teitel, "Determination of Crack Characteristics from the Quasistatic Approximation for the Scattering of Elastic Waves," *J. Appl. Phys.* 49, 5763-5767 (1978).

⁷J. D. Achenbach and A. K. Gautesen, "Geometrical Theory of Diffraction for Three-D Elastodynamics," *J. Acoust. Soc. Am.* 61, 413-421 (1977).

⁸J. D. Achenbach and A. Norris, "Crack Characterization by the Combined Use of Time-domain and Frequency-domain Scattering Data," 1981 AF/DARPA—Review of Quantitative NDE (in press).

⁹A. K. Gautesen, J. D. Achenbach, and H. McMaken, "Surface-wave Rays in Elastodynamic Diffraction by Cracks," *J. Acoust. Soc. Am.* 63, 1828-1831 (1978).

¹⁰J. B. Keller, "A Geometrical Theory of Diffraction," in *Calculus of Variations and Its Applications*, edited by L. M. Graves (Am. Math. Soc., Providence, RI, 1958).

¹¹M. J. P. Musgrave, *Crystal Acoustics* (Holden-Day, San Francisco, 1970).