

SECURITY CLASSIFICATION OF THIS PAGE

REPORT DOCUMENTATION PAGE				Form Approved OMB No. 0704-0188	
1a. REPORT SECURITY CLASSIFICATION UNCLASSIFIED			1b. RESTRICTIVE MARKINGS		
2a. SECURITY CLASSIFICATION AUTHORITY			3. DISTRIBUTION/AVAILABILITY OF REPORT Approved for public release; distribution unlimited.		
2b. DECLASSIFICATION/DOWNGRADING SCHEDULE					
4. PERFORMING ORGANIZATION REPORT NUMBER(S) NRL Memorandum Report 6138			5. MONITORING ORGANIZATION REPORT NUMBER(S)		
6a. NAME OF PERFORMING ORGANIZATION Naval Research Laboratory		6b. OFFICE SYMBOL (If applicable) Code 8312	7a. NAME OF MONITORING ORGANIZATION		
6c. ADDRESS (City, State, and ZIP Code) Washington, DC 20375-5000			7b. ADDRESS (City, State, and ZIP Code)		
8a. NAME OF FUNDING/SPONSORING ORGANIZATION ONR		8b. OFFICE SYMBOL (If applicable)	9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER		
8c. ADDRESS (City, State, and ZIP Code) Arlington, VA 22217			10. SOURCE OF FUNDING NUMBERS		
			PROGRAM ELEMENT NO. 61153N	PROJECT NO. RR0320841	TASK NO.
					WORK UNIT ACCESSION NO. 83-1137-00
11. TITLE (Include Security Classification) Inversion Algorithms for Geophysical Problems (U)					
12. PERSONAL AUTHOR(S) Lanzano, Paolo					
13a. TYPE OF REPORT Final		13b. TIME COVERED FROM 8/86 TO 7/87		14. DATE OF REPORT (Year, Month, Day) 1987 December 16	
15. PAGE COUNT 15					
16. SUPPLEMENTARY NOTATION					
17. COSATI CODES			18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number)		
FIELD	GROUP	SUB-GROUP	Fredholm integral equation Ocean spectra		
			Ocean bathymetry Calculus of variations		
			Radar measurements Minimization of functionals		
19. ABSTRACT (Continue on reverse if necessary and identify by block number)					
We discuss procedures for the iterative solution of an integral equation of the Fredholm type of generic kernel and provide details for two particular geophysical problems. We study the inversion of gravity data in order to retrieve the profile of the underlying topography by searching for solutions giving rise to the most compact configuration. We invert radar measurements to obtain ocean spectra: this is achieved according to a perturbation procedure by minimizing a functional of errors which also expresses the deviation of the solution from an expected spectral density.					
20. DISTRIBUTION/AVAILABILITY OF ABSTRACT ☒ UNCLASSIFIED/UNLIMITED ☐ SAME AS RPT ☐ DTIC USERS			21. ABSTRACT SECURITY CLASSIFICATION UNCLASSIFIED		
22a. NAME OF RESPONSIBLE INDIVIDUAL Paolo Lanzano			22b. TELEPHONE (Include Area Code) (202) 767-2778		22c. OFFICE SYMBOL Code 8312.8

DD Form 1473, JUN 86

Previous editions are obsolete.

S/N 0102-LF-014-6603

SECURITY CLASSIFICATION OF THIS PAGE

Naval Research Laboratory

Washington, DC 20375-5000

LIBRARY
RESEARCH REPORTS DIVISION
NAVAL POSTGRADUATE SCHOOL
MONTEREY, CALIFORNIA 93940



NRL Memorandum Report 6138

Inversion Algorithms for Geophysical Problems

P. LANZANO

*Space Sensing Branch
Space Systems Technology Department*

December 16, 1987

CONTENTS

I.	INTRODUCTION.....	1
II.	ANALYTICAL CONSIDERATIONS.....	1
III.	INVERSION OF GRAVITY DATA.....	4
IV.	INVERSION OF RADAR MEASUREMENTS.....	6
V.	REFERENCES.....	10
VI.	ACKNOWLEDGEMENTS.....	11

INVERSION ALGORITHMS FOR GEOPHYSICAL PROBLEMS

I. INTRODUCTION

The purpose of this note is to give a brief introduction to the concept of inversion and to discuss a few different approaches that have been considered. After some analytical preliminaries, we deal with two specific algorithms that have been used for the

(1) inversion of gravity data in order to retrieve the underlying topographic profile, and

(2) inversion of radar measurements in order to retrieve ocean spectra.

Further details and developments on this topic will constitute the subject of a future report.

II. ANALYTICAL CONSIDERATIONS

The inversion problem is analytically akin to the solution of a Fredholm integral equation of the first kind:

$$y(\eta) = \int_a^b f(\xi)K(\xi,\eta)d\xi + e(\eta) \quad (1)$$

Here a, b are given constants and $f(\xi)$ is the unknown function which, when averaged over the closed interval (a, b) according to the kernel $K(\xi, \eta)$, must reproduce the observed function $y(\eta)$ allowing for a random error $e(\eta)$. In many applications the kernel of the integral equation is referred to as the "transfer function" of the physical problem.

Methods of solution for eq. (1) can be helpful in elucidating the procedure in cases of geophysical interest.

We can express the solution of eq. (1) as a linear combination of a set of functions according to the scheme:

$$f(\xi) = \sum_{j=1}^m x_j \phi_j(\xi) + \phi^*(\xi); \quad (2)$$

the choice of the set $\{\phi_j(\xi)\}$ and the number m of functions will be made according to the specific problem and can be left, at least for now, open.

Manuscript approved October 8, 1987.

Substituting eq. (2) into eq. (1) and assuming that we know the values $y(\eta_i) = y_i$ of $y(\eta)$ at n distinct points, we get a $(n \times m)$ linear system:

$$y_i = \sum_{j=1}^m A_{ij} x_j + e_i, \quad (i = 1, 2, \dots, n) \quad (3)$$

where

$$A_{ij} = \int_a^b \phi_j(\xi) K(\xi, \eta_i) d\xi$$

and the

$$e_i = e(\eta_i) + \int_a^b \phi^*(\xi) K(\xi, \eta_i) d\xi$$

are known as the "error" terms.

The suitability of the set of functions ϕ_j and the number m of the components of the set are two factors that vary according to the physics of the problem; also the nature of the kernel function is dictated by the specific physical problem.

Thus, according to this procedure, the solution of the integral equation has been reduced to a linear system. One question which is appropriate to clarify is the significance of the parameters x_j . We can show that the x_j represent the values of the integrals of the unknown function $f(\xi)$ according to the basis functions $\phi_j(\xi)$: in other words, we can say that the x_j are the moments of $f(\xi)$ with respect to the $\phi_j(\xi)$ provided that $\phi^*(\xi)$ be chosen to be orthogonal to the elements of the basis, i.e.

$$\int_a^b \phi^*(\xi) \phi_k(\xi) d\xi = 0. \quad (4)$$

One possible choice is to assume the basis to be of the "boxcar" type, i.e.

$$\phi_j(\xi) \quad \left\{ \begin{array}{ll} = 1 & \text{for } \xi_{j-1} \leq \xi \leq \xi_j, \\ = 0 & \text{outside the interval;} \end{array} \right.$$

in which case, we have

$$x_j = \int_{\xi_{j-1}}^{\xi_j} f(\xi) d\xi / \Delta \xi_j$$

and this is the average of $f(\xi)$ within the interval $\Delta \xi_j$.

The number of parameters required for an adequate representation will depend upon the smoothness of the unknown function in the given interval. Most of the existing methods are based upon constructing linear averages of the unknown function whose values are defined by empirical data. This is essentially the approach first introduced by Backus and Gilbert (1968, 1970).

It has long been recognized that inverse problems do not yield a unique solution. This is due to the fact that there exist subsets of solutions satisfying the homogeneous integral equation which might not correspond to geophysical reality but which do contribute to the general solution of the nonhomogeneous data equations. Steps must be taken to eliminate them.

One such step is to minimize some suitable functional of the variables and errors. Thus, to eq. (3) we could add the "variational" condition that an expression of the following kind

$$q = \sum_{r=1}^n [f(x_r) + h(e_r)] \quad (5)$$

be a minimum.

Another possible step is the incorporation within the original system of a priori information by treating the unknowns as Gaussian random variables having mean values $\vec{x}_0$ and C_x as covariance matrix. To elaborate this point, let us rewrite eq. (3) in vector notation

$$\vec{y} = A\vec{x} + \vec{e} \quad (6)$$

where A is a matrix, and let us suppose that we can build some averages of our data set which might have a useful geophysical interpretation and be stable functions of the data. Let us denote this averaging operation by

$$\hat{\vec{x}} = H\vec{y},$$

so that we can write

$$\hat{\vec{x}} = HA\vec{x} + H\vec{e}.$$

We can then construct the "estimation error" as

$$\hat{\vec{x}} - \vec{x} = (HA - I) \vec{x} + H\vec{e} \quad (7)$$

where I is the identity matrix.

If we assume a covariance matrix C_x for the parameters $\vec{x}$ and a covariance matrix C_e for the errors $\vec{e}$, then the covariance matrix C for the estimation error $\hat{\vec{x}} - \vec{x}$ is given by

$$C = (HA - I)C_x(HA - I)^T + HC_eH^T, \quad (8)$$

where the upper T denotes the transpose matrix operation. We are assuming here that the a priori estimates $\vec{x}_0$ are statistically independent of the data errors $\vec{e}$.

The matrix H which minimizes the diagonal terms in C can be shown to be

$$H = C_x A^T (A C_x A^T + C_e)^{-1} \quad (9)$$

by using procedures akin to least squares problems. This matrix is known as the "minimum variance estimator" and has been used for averaging purposes among others by Jackson (1979).

III. INVERSION OF GRAVITY DATA

We must devise an algorithm that will allow us to determine the density profile of various lower layers based upon certain gravity anomaly measurements taken at the upper surface. Such an inversion scheme can be applied to ascertain the topography of seamounts and the structuring of the uppermost layers of the ocean floor.

We assume a two-dimensional fixed geometry consisting of an array of identical rectangular blocks of variable densities. We denote by x_i the abscissa of the i -th data point and by (x_j, z_j) the coordinates of the centroid of the j -th block of dimensions (d, h) and of an as yet unknown density δ_j .

For a fixed depth z_j , we reach a linear system

$$\vec{g} = A \vec{\delta} + \vec{e} \quad (10)$$

where $\vec{g}$ is the vector of the given measurements g_i ($i=1, 2, \dots, n$), $\vec{\delta}$ is the vector of the unknown densities δ_j ($j=1, 2, \dots, m$) and $\vec{e}$ is the vector of the noise e_i associated with each of the data points. A is a matrix whose elements a_{ij} represent the influence of the j -th block on the i -th gravity value, and can be represented in terms of the dimensions of the block and its relative position with respect to the location of the measurement. It is essentially the gravity generated by a vertical slab. Its expression is available in Lanzano (1984).

It is a very ascertained fact, see, e.g. Fisher and Howard (1980), that the above linear system of equations when considered in its square form (i.e. when the number of data points equals the number of blocks) is a singular system. We must, therefore, take into account a situation whereby the number m of blocks is larger than the number n of data points which allows for an infinity of solutions.

In order to bracket our solution and exclude many density distributions that lack geophysical significance, we shall make recourse to a minimum principle. We impose to minimize a suitable functional of densities and errors. Specifically, we choose here to minimize the following quadratic form of densities and errors

$$f = W_\delta \vec{\delta}^2 + W_e \vec{e}^2 \quad (11)$$

with appropriate weighting matrices of diagonal form. This functional expression was first introduced by Last and Kubik (1983). The physical significance of this condition will become evident from what follows.

The area of the configuration is given by

$$\text{Area} = dh \lim_{\epsilon \rightarrow 0} \left(\sum_{j=1}^m \frac{\delta_j^2}{\epsilon + \delta_j^2} \right) \quad (12)$$

because

$$\lim_{\epsilon \rightarrow 0} \frac{\delta_j^2}{\epsilon + \delta_j^2} \begin{cases} = 1 & \text{when } \delta_j \neq 0 \\ = 0 & \text{when } \delta_j = 0 \end{cases} \quad (13)$$

Consequently, by choosing

$$W_\delta^{-1} = \text{Diag} (\epsilon + \delta_j^2)$$

we can minimize the area of the model, i.e., maximize its compactness, by adopting a very small value for ϵ . An algorithm leading to the most compact configuration presents many attractive features because it can be used to enhance the contrast between the denser seamounts and the lighter ocean water for each solution of our system.

Once the weighting matrices are fixed, we can solve our linear system (for $m > n$) subject to the condition that the functional be a minimum by means of a least squares procedure. Specifically, we find that

$$\begin{aligned} \vec{\delta} &= W_\delta^{-1} A^T (A W_\delta^{-1} A^T + W_e^{-1})^{-1} \vec{g} \\ W_\delta^{-1} &= \text{Diag} (\epsilon + \delta_j^2) \end{aligned} \quad (14)$$

This turns out to be a nonlinear problem, because the weighting matrix W_δ which appears in the resolvent equation depends on the δ_j 's which are still unknown. This can be overcome by performing an iterative algorithm: we can choose the zero-order approximation to be

$$W_\delta^{(0)} = I \text{ (identity matrix)}$$

and subsequently we must have

$$\left[W_\delta^{(k-1)} \right]_{jj}^{-1} = \epsilon + \left[\delta_j^{(k-1)} \right]^2 \quad \text{when } k > 1$$

The k -th order iteration can then be written as follows:

$$\vec{\delta}^{(k)} = \left[W_\delta^{(k-1)} \right]^{-1} A^T \left\{ A \left[W_\delta^{(k-1)} \right]^{-1} A^T + \left[W_e^{(k-1)} \right]^{-1} \right\}^{-1} \vec{g} \quad (15)$$

In dealing with the noise terms, it has been found appropriate, see also Jackson (1979), to adopt a weighting matrix of the sort

$$W_e^{-1} = \ell_o^2 \text{Diag} (A W_\delta^{-1} A^T) \quad (16)$$

where ℓ_o is an apriori estimate of the noise to signal ratio. The signal being that part of the observed anomaly which is attributable to the structure being investigated.

The above matrix is independent of the errors themselves. Subsequent iterations will give rise to

$$\left[W_e^{(k-1)} \right]_{ii}^{-1} = \ell_o^2 D_{ii}^{(k-1)}$$

where

$$D = A W_\delta^{-1} A^T.$$

Convergence is deemed to have been achieved when further iterations do not appreciably alter the density distribution.

The constant ϵ should be chosen as small as possible but compatible with the computer so as not to cause any instability.

This algorithm can be upgraded to include an upperbound for the density values; i.e., we can impose the condition $\delta_j \leq \delta^*$ for every $j=1,2,\dots,m$ and for each depth. The algorithm must reset equal to δ^* the density of any block that goes beyond this preassigned limit and will downplay that block in the following step of the iterative procedure. This can be achieved by: (1) subtracting the gravity contribution for that particular block from the total gravity anomaly (data set), and (2) assigning a very small weight to that block so that its contribution will be negligible.

This is essentially a Penalty Function approach and can be performed by using a unit step function, see Lanzano and Myers (1985). Thus a small correction is applied to those blocks that had reached the limiting density δ^* ; should this correction be negative, bringing the total density of those blocks below the limit, the normal weighting functions should be used then, allowing those blocks to be used freely in the minimization procedure.

This algorithm has been tested at NRL on synthetic data using a small computer (HP7000). Stability and convergence required at least 12 iterations. The length of the computations did suggest that use of a faster computer is certainly desirable.

IV. INVERSION OF RADAR MEASUREMENTS

Another problem of interest concerns the determination of ocean spectra obtainable by means of experiments which depend on the direction of the wave propagation. If θ is the direction of propagation for a wave of frequency

f, the spectrum $S(\theta, f)$ is related to the data vector $\vec{d}$ via an integral equation of the sort:

$$\vec{d} = \int_0^{2\pi} S(\theta, f) \vec{k}(\theta) d\theta + \vec{e} \quad (17)$$

Here $\vec{k}$ is the vector kernel of the physical process and $\vec{e}$ is a random error vector.

The inverse problem consists of determining S by solving from eq. (17) once the $\vec{d}$ and $\vec{e}$ vectors are given. This inversion cannot be unique, since the directional distribution is a continuous function, whereas the data set $\vec{d}$ is of finite dimension. The most effective method of solution is to add to the integral equation a functional of related variables to be minimized.

We choose the functional to depend on the errors, furthermore it should establish a minimum variance of the spectrum from a particular preferred spectrum S_0 (i.e. the model to be expected) and must also express the positiveness of the spectrum. This outlook has been adopted and elaborated upon primarily by Long and Hasselmann (1979).

The functional to be minimized will be written as

$$\eta = \vec{e}^T Q \vec{e} + \alpha \int_0^{2\pi} (S - |S|)^2 d\theta + \beta \int_0^{2\pi} (S - S_0)^2 d\theta, \quad (18)$$

where Q is a symmetric, positive definite square matrix such that its inverse $V = Q^{-1}$ is the covariance matrix of the errors

$$V = \langle \vec{e} \vec{e}^T \rangle.$$

This covariance matrix V can be determined using standard techniques of time series analysis once some assumptions have been made concerning the distribution of errors.

In eq. (18), α and β are two weighting parameters whose values represent the relative importance of the conditions to which they are attached. They can be considered as Lagrangian multipliers in the sense of the classical calculus of variations approach, which is aimed at minimizing a functional to which auxiliary conditions have been appended.

The selection of the Q matrix establishes a probability region around the domain $\vec{d}$ such that

$$p(\vec{e}, \vec{d}) = \frac{1}{\sqrt{2\pi}} |V|^{-1/2} \exp \left(-\frac{1}{2} \vec{e}^T V^{-1} \vec{e} \right) \quad (19)$$

represents the probability of obtaining the values $\vec{d} \pm \vec{e}$.

We must therefore study the variation exhibited by the functional

$$\eta(S, |S|, \vec{e})$$

caused by all possible variations δS of the spectrum $S(\theta, f)$ within a class of allowable solutions to eq. (17).

For this purpose, let us recall that $\vec{e}$ and S are related because, from eq. (17), one has

$$\vec{e} = \vec{d} - \int_0^{2\pi} S(\theta, f) \vec{k}(\theta) d\theta .$$

Since neither the vector $\vec{k}$ (i.e., the physical process) nor the vector $\vec{d}$ (i.e., the data points) are supposed to undergo any variation, the variations δS and $\delta \vec{e}$ can be related according to

$$\delta \vec{e} = - \int_0^{2\pi} [\vec{k}(\theta) \delta S] d\theta . \quad (20)$$

Thus, the variation $\delta \eta$ can be written in terms of partial derivatives and we reach an expression of the sort

$$\delta \eta = \left(\frac{\partial \eta}{\partial S} + \frac{\partial \eta}{\partial \vec{e}} \cdot \frac{\partial \vec{e}}{\partial S} \right) \delta S .$$

We get

$$\begin{aligned} \delta \eta = & \alpha \int_0^{2\pi} [2(S - |S|) \left(\frac{dS}{dS} - \frac{d|S|}{dS} \right) \delta S] d\theta \\ & + \beta \int_0^{2\pi} [2(S - S_0) \delta S] d\theta + 2\vec{e}^T Q \delta \vec{e} . \end{aligned}$$

When $S > 0$, the integrand in the first term will vanish identically; on the other hand when $S < 0$, we have

$$\frac{d|S|}{dS} = -1 .$$

We can therefore write in both cases that $\delta \eta = 0$ is equivalent to

$$\int_0^{2\pi} \{ [4\alpha(S-|S|) + 2\beta(S-S_0) - 2\vec{e}^T Q\vec{k}] \delta S \} d\theta = 0,$$

where use has been made of eq. (20) to eliminate the $\delta\vec{e}$.

Due to the arbitrariness of δS , the integrand must vanish, leading to

$$4\alpha(S-|S|) + 2\beta(S-S_0) = 2\vec{e}^T Q\vec{k}.$$

When $S>0$, we get

$$S = S_0 + \frac{1}{\beta} \vec{e}^T Q \vec{k}. \quad (21)$$

Whereas, when $S<0$, we have

$$(8\alpha + 2\beta)S = 2\beta S_0 + 2\vec{e}^T Q\vec{k},$$

$$\text{or} \quad S = \frac{\beta}{4\alpha + \beta} \left(S_0 + \frac{1}{\beta} \vec{e}^T Q \vec{k} \right). \quad (22)$$

Knowing $\vec{k}$ and assuming a certain $\vec{e}$, we can thus determine the function S which minimizes the η functional provided that we choose the two parameters (α, β) .

Let us rewrite the last term within parenthesis in an equivalent but more convenient way. Since

$$\vec{e}^T Q\vec{k}$$

is a scalar, it must coincide with its transpose. Applying the property that the reverse order of the transposed factors is used in performing the transpose of a product of matrices, we can write

$$\vec{e}^T Q\vec{k} = (\vec{e}^T Q\vec{k})^T = \vec{k}^T Q^T \vec{e} = \vec{k}^T Q\vec{e}; \quad (23)$$

the last step is due to the fact that Q is a symmetric matrix.

We must now impose that the condition pertaining to the positiveness of S be verified exactly: this necessitates to let α approach infinity within eq. (18).

Eq. (22) will then vanish; whereas eq. (21) can be rewritten, through the use of eq. (23), as

$$S = S_0 + \frac{1}{\beta} (\vec{k}^T Q\vec{e}). \quad (24)$$

This function S , however, should also be a solution of eq. (17); we must therefore substitute S from eq. (24) into eq. (17). In doing so, we get

$$\vec{d} - \int_0^{2\pi} (S_0 \vec{k}) d\theta = \vec{e} + \frac{1}{\beta} \int_0^{2\pi} (\vec{k}^T V^{-1} \vec{e}) \vec{k} d\theta . \quad (25)$$

The left-hand side of eq. (25) is a known quantity; both $\vec{e}$ and β , however, are still unknown. The above equation turns out to be a Fredholm integral equation of the second kind with respect to $\vec{e}$ with variable parameter β ; in it the unknown function $\vec{e}$ appears both inside and outside the integral sign.

Equation (25) can be solved by iterative procedures, noting however that any initial choice for $\vec{e}$ or subsequent iterations should necessarily imply the determination of the covariance matrix V . Once $\vec{e}$ is determined by solving eq. (25), eq. (24) will provide the desired spectrum S .

Various possibilities exist for solving eq. (25) and they will be discussed extensively in a future report. Our present analysis has taken us from the original eq. (17) to this final eq. (25) through the process of assuming: (1) only positive values for S , and (2) a minimum variation of S from an expected function S_0 .

By imposing these physical limitations which are of plausible acceptance, we have gained a computational advantage because eq. (25) is easier to solve than the original eq. (17).

An equation similar to our eq. (25) above was used by Long and Hasselmann (1979) to measure the directional properties of swells in shallow waters for the purpose of comparing various swell decay models available.

V. REFERENCES

1. G. E. Backus and Gilbert, F. (1968), Geophys. J. R. Astron. Soc. 16, 169.
2. G. E. Backus, and Gilbert, F. (1970), Phil. Trans. R. Soc. Lond. 266, 123.
3. N. J. Fisher and Howard, L. E. (1980), Geophysics 45, 186.
4. D. D. Jackson (1979), Geophys. J. R. Astron. Soc. 57, 137.
5. P. Lanzano (1984), Trans. A.G.U. 65, 16, 185.
6. P. Lanzano and Myers, C. G. (1985), Trans. A.G.U. 66, 46, 844.
7. B. J. Last and Kubik, K. (1983), Geophysics 48, 713.
8. R. B. Long and Hasselmann, K. (1979), J. Phys. Oceanogr. 9, 373.

VI. ACKNOWLEDGEMENTS

The author wishes to express his sincere appreciation to Dr. Vincent E. Noble for sponsoring the development of geophysical inversion algorithms and to Ms. Linda D. Ridgely for typing the manuscript.

U 235378

DEPARTMENT OF THE NAVY

NAVAL RESEARCH LABORATORY
Washington, D. C. 20375-5000

OFFICIAL BUSINESS

PENALTY FOR PRIVATE USE, \$300

THIRD-CLASS MAIL
POSTAGE & FEES PAID
USN
PERMIT No. G-9